



Fremtid. Nutid. Datid.

3-2-1 REPEAT

“En guidet tur bag kulisserne i den
kunstige intelligens’ verden”

Skrevet og udgivet af
www.boardbusiness.dk

FORORD

Der findes tidspunkter i historien, hvor noget grundlæggende forandrer sig. Ikke med et brag, men med en glidende bevægelse, vi først forstår, når vi ser os tilbage. Kunstig intelligens er netop sådan en bevægelse.

Da jeg begyndte at arbejde med AI, var det ikke, fordi jeg var teknolog. Det var, fordi jeg var nysgerrig. Nysgerrig på, hvad det betyder, når maskiner begynder at forstå sprog, tage beslutninger, lære af vores adfærd og endda skabe nyt.

Jeg oplevede, hvordan AI ikke længere var noget, der kun foregik i Silicon Valley eller i lukkede laboratorier. Det var noget, der ramte mig selv, min virksomhed, miljøet, etik, cybersikkerhed samt skoler og uddannelsesinstitutioner o.m.fl.

Denne bog er mit forsøg på at forstå og formidle de mange lag, som AI bringer med sig. Fra de mindste elektriske kredsløb i en chip til de største globale magtspil. Fra håb om bedre uddannelse og sundhed til bekymringer om overvågning og magtkoncentration.

Jeg har valgt at skrive bogen som en blanding af analyser, cases og refleksioner. Det er ikke en teknisk manual – det er en rejse gennem fremtid, nutid og datid. En rejse, hvor du som læser er med hele vejen, uanset om du er AI-ekspert eller blot har hørt om ChatGPT.

Vi har set mønstre før: ny teknologi, nye magtbalancer, nye muligheder og nye dilemmaer. Denne gang handler om at vælge, hvad vi vil forandre.

God læselyst

Erik Jul Nielsen

Ansvarlig narrativist,
AI-Alkemist og Cyberintelligens-arkitekt

INDHOLD

DEL I

AI som infrastruktur og magtfaktor

Fokus: "Teknologiens grundlag, de globale magtkampe, energiforbrug og samfundsstrukturer"

Tema: "Hvordan AI former verdens magtbalancer, arbejdsliv og fremtidige ressourcer".

Kap 1: AI's hardwarekapløb: Fra chipmangel til supercomputere

Kap 2: AI i det globale magtspil. Stormagter, data og digital suverænitæt

Kap 3: AI's energitørst og kampen for en grøn fremtid

Kap 4: Når AI ændrer måden, vi arbejder på

Kap 5: Når AI ændrer måden, vi lærer på: Uddannelse i AI-æraen

DEL II

AI i menneskets og samfundets tjeneste (eller tjeneste?)

Fokus: De aktører, normer og dilemmaer, der styrer, former og udfordrer AI's udbredelse

Tema: Hvordan vi som samfund, borgere og individer navigerer i AI's konsekvenser.

Kap 6: AI-giganter i magtkamp: Alliancer, monopoler og fremtidens spillere

Kap 7: AI, etik og demokratiets fremtid: Hvem bestemmer over teknologien?

Kap 8: Privatliv under pres: AI og dataindsamling i hverdagen

Kap 9: AI som våben: Cybersikkerhedens nye frontlinje

Kap 10: AI og global retfærdighed: Den digitale kløft

DEL III

Udblik og refleksion

Efterord: Fremtiden formes af os

Kritisk refleksion over e-bogen: "Fremtid, nutid og datid... 3-2-1 repeat"

Kildeoversigt

INDLEDNING

For få år siden var kunstig intelligens et buzzword for eksperter. I dag har millioner stiftet bekendtskab med AI gennem ChatGPT.

Det føles næsten magisk at kunne få svar fra en "tænkende" maskine, og for mange er AI blevet hvermandseje, men hvad ligger egentlig bag denne magi?

Under overfladen udspiller der sig en global kamp om AI - en kamp om hardware, ressourcer og indflydelse, der vil forme vores fremtid.

Jeg har arbejdet med AI-teknologi i flere år og er imponeret over de daglige fremskridt, men jeg er også opmærksom på skyggesiderne: det enorme forbrug af strøm og knappe råmaterialer. Jeg kan ikke lade være med at spørge: Hvad koster det at gøre en chatbot som ChatGPT mulig - ikke bare i penge, men også i ressourcer, energi og politisk indflydelse?

Denne bog er mit forsøg på at udforske det spørgsmål - en rejse ind i AI-økosystemets "maskinrum", fra de mindste mikrochips til de største magtkampe.

AI kræver enorme mængder regnekraft, og lige nu kappes lande og virksomheder om at bygge verdens mest kraftfulde datacentre.

Man kan tænke på dem som moderne kraftværker for AI: enorme haller fyldt med specialiseret computerhardware. Et eksempel er

det nyligt annoncerede "Stargate"-projekt i USA, hvor flere teknologigiganter i samarbejde med regeringen vil investere svimlende summer i ny AI-infrastruktur. Og USA er ikke alene; lignende superdatacentre skyder op globalt, for ingen vil stå bagerst i AI-kapløbet.

I hjertet af datacentrene finder man grafikprocessorerne - de såkaldte GPU'er. Oprindeligt udviklet til videospil er GPU'er i dag blevet AI-æraens motor.

Efterspørgslen på de kraftigste chips er eksploderet. Der er opstået mangel på GPU'er, så virksomheder hamstrer dem, og priserne stiger voldsomt.

"Data er det nye olie," siger man - men man kunne lige så godt sige, at GPU'er er det nye guld, for uden dem går AI-udviklingen i stå.

Men AI-kapløbet handler ikke kun om hardware og penge - det er også forankret i geopolitik. Krigen i Ukraine har tydeligt vist, hvor sårbare teknologiske forsyningskæder er. Næsten halvdelen af al neon-gas til chipproduktion kom fra Ukraine, så da leverancerne derfra blev lammet, frygtede chipproducenter en ny mangelkrise.

Samtidig udkæmper stormagterne en "chipkrig", hvor USA begrænser eksporten af avancerede chips til Kina for at bremse Kinas fremmarch, mens Kina investerer massivt i at blive selvforsynende. Sjældne mineraler, ekspertise og produktionsevne er blevet strategiske brikker i spillet om global magt. Adgangen til de bedste chips

og mest regnekraft er ikke længere et teknisk spørgsmål, men storpolitik. Kontrol over teknologien er lig med magt - derfor kæmper både lande og teknologigiganter om førerpositionen.

Midt i dette kapløb overses ofte en afgørende faktor: energien. De enorme datacentre bag AI sluger kolossale mængder strøm. Hver gang vi spørger en chatbot, bruger den strøm - et AI-svar kræver mange gange mere elektricitet end en almindelig Google-søgning. Hvor skal al den nødvendige strøm komme fra? Kan elnettet og den grønne omstilling følge med, når computere tænker på højtryk døgnet rundt?

Alt dette viser, at AI ikke blot handler om smarte algoritmer, men om et komplekst samspil af maskiner, mennesker og magt.

I denne bog inviterer jeg dig med ind bag kulisserne i "Kampen om AI" for

at gøre denne virkelighed forståelig. Vi skal se på alt fra hardware-kapløb og geopolitik til energiforsyning og teknologigiganternes indflydelse. Som nysgerrig formidler med kritisk sans og baggrund i AI deler jeg min begejstring og mine bekymringer - forhåbentlig har det pirret din nysgerrighed og givet dig lyst til at læse videre.

Kampen om AI er i fuld gang, og dens udfald vil påvirke os alle.

Velkommen til en rejse, der bliver både oplysende og spændende.

Fakta-tjek

OpenAI's ChatGPT vs. Danmarks elforbrug

ChatGPT's årlige strømforbrug:
ca. 226,8 GWh

Danmarks samlede elforbrug i 2023:
36.100 GWh

ChatGPT's årlige strømforbrug svarer til
ca. 0,63% af Danmarks samlede elforbrug i 2023.

AI'S HARDWAREKAPLØB - FRA CHIPMANGEL TIL SUPERCOMPUTERE

Indledning

Jeg husker tydeligt første gang, det gik op for mig, at nøglen til moderne AI's magi ikke kun ligger i smarte algoritmer, men i det fysiske maskinrum bagved hardwaren.

Vi hører ofte om AI-modeller som ChatGPT og deres imponerende sprogforståelse, men sjældnere om de tusindvis af specialiserede chips og kolossale datacentre, der gør det hele muligt. I disse år, og særligt frem til maj 2025 er kapløbet om AI-hardware blevet lige så intenst som udviklingen af selve AI-modellerne.

Lad os tage på en rejse ind i maskinrummet og se, hvordan supercomputere, GPU'er, geopolitik og nye aktører former kampen om AI og hvorfor dette hardwarekapløb er så strategisk vigtigt.

AI's umættelige hunger efter regnekraft

AI's seneste fremskridt har afsløret en nærmest umættelig appetit på regnekraft. Hver gang vi klikker på et svar fra en AI som ChatGPT, udføres der millioner af beregninger i baggrunden. Analytikere har

estimeret, at hver enkelt ChatGPT-forespørgsel koster omkring 4 USD-cent i ren computertid, svarende til cirka 25 øre. Det lyder måske ikke af meget, men med millioner af daglige forespørgsler løber det op i svimlende summer og kræver enorme mængder hardware.

Historisk set blev AI trænet på CPU'er (centralprocessorer), men omkring 2012 opdagede forskere, at grafikprocessorer - GPU'er, var langt bedre til de massive parallelle beregninger, som moderne neurale netværk kræver.

En GPU kan udføre tusindvis af simple beregninger sideløbende, hvilket gør den ideel til at træne dybe neurale netværk. Dette gennembrud satte gang i AI-revolutionen, men lagde også kimen til en ny begrænsning: adgangen til nok af de kraftige chips.

AI-modellerne voksede eksplosivt i størrelse, og dermed steg behovet for regnekraft tilsvarende. I dag er sprogmodeller og billedgeneratorer milliardvis af parametre tunge, og deres "træningsdiæt" består af ufattelige datamængder.

For at følge med har hardware-producenter lanceret specialchips og sammensat dem i supercomputere, der udelukkende er dedikeret til AI.

Tænk på en supercomputer som et gigantisk puslespil bestående af

tusindvis af GPU'er og andre accelerators, der er forbundet og arbejder som én samlet hjerne. For blot få år siden lød det vildt, at et AI-system skulle bruge 10.000 GPU'er. Alligevel byggede Microsoft og OpenAI allerede i 2020 en sådan supercomputer til at udvikle GPT-3.

Nu, i 2025, er barren løftet endnu højere: Meta (Facebook) har præstet af et AI-datacenter med 16.000 GPU'er, og flere nationale forskningscentre råder over lignende, eller endnu større systemer.

Vi taler om maskiner, der kan udføre **eksaflops** [1], som er 1 milliard milliard beregninger i sekundet. Et tal så enormt, at det næsten trodses fatteevne.

At AI-hardware er blevet nøglen til fremskridt, ses tydeligt ved, at selv softwarefirmaer nu kaster sig over hardwareudvikling. Googles moderselskab Alphabet designer sine egne TPU-chips (Tensor Processing Units) målrettet AI og har i 2024 lanceret femte generation (TPU v5p), med en sjette generation på vej, hver ny generation mange

gange kraftigere end den forrige. Tesla har bygget sin Dojo-supercomputer med egenudviklede chips til at træne selvkjørende biler. Og Nvidia - virksomheden der oprindeligt lavede grafikchips til gamere, er nu blevet en af AI-revolutionens absolutte hovedaktører.

Allerede i 2023 blev Nvidia så værdifuldt, at selskabets markedsværdi rundede 1.000 milliarder dollars, drevet af forventningerne til AI.

Nvidias seneste GPU-generationer (A100, H100 og den kommende "Blackwell"-chip) fungerer i dag som motoren bag utallige AI-systemer verden over.

GPU-guldfeber og knaphed

Den enorme efterspørgsel på AI-hardware har skabt en regulær GPU-guldfeber. Virksomheder, forskningsinstitutioner og endda nationer hamstrer de kraftigste chips - ofte hurtigere, end producenterne kan fremstille dem. Resultatet? En global knaphed på topchips, hvor ventelister og opskruede priser er blevet normen.

NOTE [1]

"Eksaflops" er en måleenhed for computerens regnekraft og står for eksafloating point operations per second: Altså en milliard milliarder (10^{18}) flydende kommatall-beregninger i sekundet.

Her er en kort forklaring:

- FLOPS = Floating Point Operations Per Second (en måleenhed for, hvor mange matematiske beregninger en computer kan udføre i sekundet).
- Eksa = 10^{18} (det vil sige 1.000.000.000.000.000.000).
- Så 1 eksaflops = 1.000.000.000.000.000.000 beregninger i sekundet.

Hvor bruges eksaflops?

Eksaflops bruges især til at måle ydeevnen af supercomputere. For eksempel:

- Supercomputere som Frontier (USA) og Fugaku (Japan) har nået eller nærmer sig eksaflops-ydelse.
- De bruges til komplekse opgaver som klimamodellering, simulering af atomreaktioner, kunstig intelligens og medicinsk forskning.

I kølvandet på COVID-19-pandemien oplevede vi allerede en generel chipmangel, men AI-boomet har potentiale til at udløse en ny bølge af chip-underskud.

Analytikere fra Bain & Company advarer om, at den AI-drevne stigning i GPU-efterspørgslen kan øge behovet for visse komponenter med 30% eller mere inden 2026. Sådanne pludselige stigninger risikerer at vælte balancen i de delikate forsyningskæder og udløse endnu en global chipkrise.

Vi ser allerede tegn på presset. I 2023 måtte selv en gigant som OpenAI erkende, at de løb tør for GPU-kapacitet. Sam Altman, OpenAI's direktør, har offentligt beklaget, hvor svært det er at skaffe nok avancerede chips og han kaldte omkostningerne: "øjenvandsfremkaldende høje".

Faktisk dominerer Nvidia over 80% af det globale marked for AI-specialiserede chips, hvilket gør næsten alle store AI-aktører afhængige af Nvidias produktion.

Da OpenAI planlagde næste trin (GPT-4.5), måtte de angiveligt udskyde lanceringen, simpelthen fordi de ikke havde tilstrækkelig regnekraft til rådighed. Det siger noget, når selv frontløberne i AI ikke kan få fingre i nok hardware.

Denne GPU-mangel mærkes bredt. Opstartsvirksomheder med gode

idéer kæmper for adgang til træningskapacitet, fordi de ikke kan købe eller leje nok GPU'er.

Cloud-udbydere melder om udsolgte ressourcer i deres mest populære datacentre til AI-opgaver. Priserne på de nyeste GPU'er er skudt i vejret og det samme er Nvidias indtjening. Ved udgangen af 2024 kunne Nvidia rapportere rekordresultater og en næsten fordobling af deres omsætning, primært drevet af AI-chips.

Paradoksalt nok har denne gyldne efterspørgsel også ført til, at forskere nu taler om at optimere og gøre AI-modeller mere effektive, så de ikke sluger helt så meget regnekraft. Men indtil videre er tendensen klar: Jo mere avanceret AI bliver, jo flere og vildere chips skal der til.

Når hardware bliver storpolitik

Den eksplosive betydning af AI-hardware er heller ikke gået ubemærket hen i de geopolitiske magtspil.

Chips er blevet den nye olie - et strategisk aktiv, som nationer beskytter og konkurrerer om. Intet sted er dette mere tydeligt end i forholdet mellem USA og Kina.

USA har i de seneste år strammet eksportkontrollen for avancerede halvledere markant. I oktober 2022 indførte amerikanerne forbud mod at sælge de kraftigste AI-chips (som

Nvidias A100) til Kina af frygt for, at de kunne bruges militært.

Da Nvidia hurtigt udviklede en nedgraderet version (A800/H800) for at omgå reglerne, svarede USA igen med yderligere stramminger i 2023 og 2024. Resultatet er, at selv Nvidias specialversion H20-chip, som egentlig var godkendt til Kina, nu er blevet blokeret med krav om eksportlicens.

Nvidia er således endnu en gang tvunget til at udvikle en "kinesisk" version med reducerede specifikationer for at holde sig inden for reglerne.

For Kina, der importerer chips for milliarder af dollar, har disse restriktioner været et wakeup-call. Kinesiske tech-giganter som Alibaba og Tencent skyndte sig i starten af 2023 at forudbestille store mængder af den daværende H20-chip, før den også blev ramt af forbud. Samtidig investerer Kina massivt i at opbygge en selvstændig chipsektor, fra design af egne GPU-lignende acceleratorer til udvikling af avancerede fabrikker. Men at indhente USA og dets allierede på chipteknologi er lettere sagt end gjort. Meget af den knowhow og det nødvendige udstyr er nemlig koncentreret hos få globale aktører: **ASML** i Holland, der fremstiller de ekstreme litografimaskiner til de nyeste chips, og **TSMC** i Taiwan, der producerer størstedelen af verdens mest avancerede chips.

Disse aktører er blevet brikker i storpolitik. Taiwan befinder sig i et geopolitisk spændingsfelt mellem USA og Kina. Skulle konflikten blusse op, kan det lamme forsyningen af avancerede chips til hele verden. Netop derfor har USA vedtaget omfattende støttelovgivning (CHIPS Act) for at opføre flere chipfabrikker på amerikansk jord, og EU har lanceret sin egen Chips Act for at øge Europas andel af produktionen.

Geopolitikken påvirker også hardwarekapløbet indirekte.

Moderne chips kræver sjældne jordarter og gasser for at blive fremstillet. Et konkret eksempel kom med krigen i Ukraine: Over halvdelen af verdens forsyning af neon-gas - essentiel til lasere, der etser kredsløb på wafers, kom før 2022 fra to ukrainske firmaer. Da krigen brød ud, måtte begge lukke ned, hvilket med ét slag fjernede næsten 50% af den globale neonproduktion. Det faldt oven i en eksisterende chipmangel og truede med at forværre situationen yderligere.

Heldigvis lykkedes det industrien at finde midlertidige alternativer og tære på lagre, men situationen viste, hvor sårbar forsyningskæden er.

Kina, som dominerer udvindingen af mange sjældne jordarter, har også vist vilje til at bruge dette som pression. I 2023 begyndte Kina at begrænse eksporten af metaller som gallium og germanium, som er vigtige

i visse højteknologiske komponenter. Det blev set som et modsvar til USA's chiprestriktioner. En påmindelse om, at når det gælder kritiske materialer og AI-hardware, er verden filtret ind i en ny teknologisk kold krig.

Nye spillere og alliancer i AI-kapløbet

Midt i det globale AI-kapløb er nye aktører stormet ind på scenen. AI-feltet var engang domineret af klassiske tech-giganter som Google, Apple, Facebook (Meta) og Amazon, men nu ser vi specialiserede startups, bilproducenter og endda enkeltpersoner med store visioner blande sig i toppen.

Én af de mest opsigtsvækkende nye spillere er Elon Musk. Efter at have været med til at stifte OpenAI (som han senere forlod) besluttede Musk i 2023 at starte sin egen AI-satsning under navnet xAI. Og han nøjedes ikke med at hyre kloge hoveder, han gik også all-in på hardware.

Musk har sat sig for at bygge en af verdens største AI-supercomputere, passende navngivet Colossus. Allerede nu kører Colossus på over 100.000 Nvidia H100 GPU'er. En enorm samling chips og planen er at udvide til intet mindre end 1.000.000 GPU'er.

Hvis det lykkes, vil Colossus blive en af de mest kraftfulde supercomputere nogensinde. Ambitionen er skyhøj, men det understreger

pointen: Nye aktører er villige til at investere milliarder i hardware for at sikre sig en plads i AI-toppen.

At skaffe så mange GPU'er er i sig selv en kæmpe udfordring, da efterspørgslen på Nvidias bedste chips allerede er enorm. Ikke desto mindre samarbejder xAI med blandt andre Nvidia og flere serverproducenter for at nå målet og endda med bystyret i den amerikanske by Memphis, hvor Colossus opføres.

Musk er ikke alene om at tænke stort. Overalt ser vi alliancer og opbrud i det traditionelle mønster. Microsoft har investeret massivt i OpenAI og sikret sig adgang til nødvendig hardware gennem medeje. En slags "AI-aktie for hardware"-byttehandel.

Google har fusioneret deres DeepMind- og Brain-afdelinger og benytter både egne TPU-chips og Nvidias GPU'er for at holde trit.

Amazon udvikler sine egne AI-chips: Trainium og Inferentia, som anvendes i AWS-cloud, så de kan tilbyde kunderne alternativ regnekraft uden kun at være afhængige af Nvidia.

Selv hardwarefokuserede virksomheder får nyt liv. Firmaet Cerebras har skabt verdens største chip på størrelse med en hel wafer med millioner af kerner. Sammen med investorer fra Emiraterne er de i gang med at bygge AI-supercom-

putere baseret på denne teknologi. Disse alternative tilgange udfordrer standardmodellen og beriger landskabet med nye muligheder. Og vi må ikke overse de nationale initiativer.

Lande som Storbritannien, Frankrig og andre investerer i nationale AI-supercomputere for at sikre, at deres forskning og industri har adgang til tilstrækkelig infrastruktur.

I Danmark taler vi måske ikke om at bygge egne chips, men fokus er stigende på at sikre adgang til cloud-ressourcer og måske samarbejde nordisk eller europæisk for at følge med.

Fælles for de nye spillere er erkendelsen af, at hardware ikke længere blot er en understøttende birolle i AI-historien, det er blevet en hovedperson. At have den bedste AI-model nytter ikke meget, hvis man ikke har musklerne (i form af chips og computere) til at træne den fuldt ud eller gøre den tilgængelig for brugere. Derfor ser vi nu AI-iværksættere diskutere chiparkitekturer, investorer interessere sig for forsyningskæder, og virksomhedsledere indgå strategiske partnerskaber om datacentre og silicium.

Afslutning

AI-hardware og infrastruktur er i dag en strategisk ressource på linje med energi og data. Den, der kontrollerer

adgangen til massiv regnekraft, kontrollerer i vid udstrækning tempoet for AI-udviklingen og potentielt også fordele i alt fra økonomi til militær styrke.

Kampen om AI skrives således ikke kun i kode og algoritmer, men også i kobberforbindelser på printplader og i lyden af køleanlæg i verdens datacentre. Fra GPU-mangel og geopolitisk rivalisering til nye aktører med storladne visioner har vi set, hvorfor hardwaren er blevet det varme omdrejningspunkt.

Forhåbentlig står det nu klart for os alle, at næste gang man imponeres over et AI-svar, gemmer der sig en hel kæde af teknologisk infrastruktur og globalt samarbejde eller konkurrence i kulisserne. En kamp om AI's hjerte og hjerne, hvor hardwaren udgør det solide fundament for fremtidens AI-eventyr.

GPU'er er AI's motor

Grafikprocessorer (GPU'er), oprindeligt designet til videospil, er nu afgørende for AI-træning. Efterspørgslen på avancerede chips som Nvidias Blackwell-serie er så høj, at de er udsolgt i op til 12 måneder frem.

Superdatacentre skyder op globalt

Projekter som det amerikanske "Stargate" og Teslas 50.000 GPU-baserede supercomputer illustrerer den massive investering i AI-infrastruktur. Disse datacentre fungerer som moderne kraftværker for AI.

Neon-gas er en kritisk komponent

Næsten halvdelen af verdens neon-gas, essentiel for chipproduktion, kom tidligere fra Ukraine. Krigen har forstyrret leverancerne, hvilket har ført til globale forsyningsproblemer.

AI I DET GLOBALE MAGTSPIL: STORMAGTER, DATA OG DIGITAL SUVERÆNITET

Du har måske stiftet bekendtskab med AI via ChatGPT - et hyggeligt værktøj til alt fra e-mails til jokes. Men bag kulisserne foregår en kamp om AI. En global magtkamp, hvor algoritmer, data og chips er blevet lige så strategiske som olie, stål og atomvåben var i det 20. århundrede. AI er ikke længere kun for forskere; det er blevet en central brik i storpolitik og økonomisk indflydelse verden over.

I disse år ser vi verdens ledere omtale kunstig intelligens med samme alvor som national sikkerhed.

Beslutninger om, hvem der må få adgang til de mest avancerede AI-chips, træffes nu i regeringskontorer og internationale fora. AI er med andre ord rykket ind på den geopolitiske scene som en ny kamplads.

Tech-investeringer og forsknings-fremskridt bliver vurderet som potentielle trusler eller strategiske aktiver for nationale magtpositioner.

USA's regering har f.eks. indført strenge restriktioner på eksport af avancerede chips og teknologi til Kina for at bremse Kinas AI-fremskridt. Kina, på sin side, investerer milliarder i at opbygge en indenlandsk teknologisektor og mindske afhængigheden af vestlig hardware.

Stormagternes AI-strategier
For at forstå AI's geopolitiske betydning er det værd at se nærmere på, hvordan stormagterne griber det an strategisk.

USA har historisk haft en laissez-faire-tilgang: Innovation drives af private virksomheder og universiteter, mens staten støtter via forsvarsprojekter og forskningsmidler. Amerikanerne har endnu ingen samlet føderal AI-lovgivning, da man længe har frygtet, at for mange regler ville kvæle innovation. I stedet satser USA massivt på at fastholde føringen. Milliardbeløb investeres i AI-forskning (bl.a. gennem DARPA), og tech-giganter som Google, Microsoft og OpenAI står bag mange af verdens mest avancerede AI-modeller. Samtidig begynder Washington så småt at overveje reguleringer for ansvarlig brug.

Kina benytter en centralstyret tilgang med langsigtede planer. Allerede i 2017 vedtog Beijing en national AI-strategi, der uden omsvøb proklamerede målet: "Kina

skal være verdens førende AI-nation senest i 2030". Vejen dertil går gennem massive statslige investeringer, uddannelse af titusinder af AI-eksperter og en tæt alliance mellem stat og erhvervsliv.

Tech-virksomheder som Alibaba, Baidu, Tencent og Huawei udnævnes til "nationale mestre" og modtager massiv støtte. Samtidig udnytter Kina sin enorme befolknings data som råstof – med langt færre privatlivsrestriktioner end i Vesten kan kinesiske AI-projekter trækkes på kolossale datamængder.

AI bruges desuden aktivt til overvågning og social kontrol.

Målet er klart: Kina vil være teknologisk selvforsynende og uafhængig af vestlig indflydelse, så landet ikke kan bremses af fremtidige sanktioner eller handelsboykotter.

Europa og EU indtager en anden rolle. De europæiske lande har deres egne AI-initiativer, men fælles er et stærkt fokus på regulering og værdier. EU ønsker at være "rule-maker" frem for "rule-taker" i den digitale verden.

Det betyder, at man hellere vil sætte globale standarder for ansvarlig AI end nødvendigvis opfostre egne tech-giganter.

EU har allerede markeret sig med GDPR og er nu i færd med at vedtage

en omfattende AI-forordning, som skal sikre, at kunstig intelligens anvendes sikkert, gennemsigtigt og etisk i Europa. Samtidig satser EU på digital suverænitet ved at investere i kritisk teknologi. Med Chips Act fra 2022 er målet at fordoble EU's andel af global chipproduktion fra ca. 10% til 20% inden 2030.

Selvom Europa ikke har tech-giganter på niveau med USA eller Kina, søger man at beskytte sin position gennem internationale alliancer (særligt med USA) og ved at påvirke den globale standardudvikling gennem lovgivning.

Andre lande uden for Vesten satser også stort på AI. Indien ønsker at blive et globalt AI-knudepunkt.

Rusland har en AI-plan, dog hæmmet af internationale sanktioner. Og olierige golfstater investerer milliarder i AI for at fremtidssikre deres økonomier.

AI står kort sagt højt på den globale dagsorden - alle ønsker at udnytte teknologien til at styrke deres position og indflydelse.

Alliancer og teknologiske blokke

AI-kapløbet handler ikke kun om enkeltlande - det udspiller sig også i form af alliancer og teknologiske blokke.

Demokratiske lande, anført af USA og allierede i Europa og Asien, rykker

tættere sammen for at beskytte deres teknologiske forspring og forhindre, at strategisk teknologi havner hos rivaler. Kina, på sin side, styrker båndene til bl.a. Rusland og investerer massivt i digital infrastruktur i udviklingslande for at opbygge en egen teknologisk indflydelsessfære.

Et markant eksempel på koordinering er, hvordan USA i 2022–23 fik Holland og Japan med på at begrænse eksporten af avanceret chip-udstyr til Kina. Begge lande spiller nøgleroller i chipproduktion, og ved fælles indsats kan man effektivt bremse Kinas adgang til den nyeste teknologi.

Ligeledes har USA og EU oprettet fora som Trade and Technology Council (TTC) for at koordinere teknologi- og AI-politik. Her handler det ikke kun om industri, men også om værdier. På det første TTC-møde erklærede parterne f.eks. fælles modstand mod "autoritær" brug af AI som Kinas sociale kreditsystem.

Sådanne tiltag signalerer en demokratisk alliance, hvor AI skal udvikles og anvendes i overensstemmelse med menneskerettigheder og frihedsrettigheder.

Den digitale magtkamp: data, infrastruktur og standarder

Kontrol over data er blevet et kernepunkt i AI-konkurrencen.

Træning af avancerede algoritmer kræver enorme mængder data, og derfor stiller mange lande sig nu spørgsmålet: Hvem ejer og beskytter dataene?

Vi ser en tydelig tendens mod datasuverænitet, hvor stater ønsker at holde data inden for egne grænser. Kina har for eksempel vedtaget en lov, der kræver, at kinesiske borgeres persondata lagres i Kina. Dels for at beskytte national sikkerhed, dels for at styrke egne virksomheders adgang til data.

EU fokuserer på privatlivsbeskyttelse, men effekten er lignende: Strenge regler (som GDPR) begrænser, hvordan virksomheder må overføre europæiske data ud af regionen.

Selve data-flowet er blevet et geopolitisk anliggende.

Et opsigtsvækkende eksempel var bekymringen over den kinesisk-ejede app TikTok i 2023: Flere vestlige lande forbød TikTok på offentligt ansattes telefoner af frygt for, at brugerdata kunne ende i hænderne på den kinesiske stat. Data er med ét blevet et strategisk aktiv på linje med olie eller sjældne mineraler.

Også digital infrastruktur; ryggraden i AI-verdenen er genstand for magtkamp. Tag f.eks. 5G-netværkene: Kinesiske Huawei var førende med billig 5G-teknologi, men

USA advarede mod mulige bagdøre til spionage. Mange lande valgte derfor at udelukke Huawei fra deres 5G-udbygning, selvom det forsinkede processen.

Episoden viste, at afhængighed af en teknologisk rival til kritisk infrastruktur opfattes som uacceptabel.

I dag ønsker stater pålidelige og kontrollerbare leverandører af alt fra telenet og elnet til cloud-data-centre og undersøkabel internet-forbindelse. Den, der kontrollerer infrastrukturen, har nemlig også magten til at bestemme vilkårene for teknologiens brug.

Kampen om standarder er måske mindre synlig, men mindst lige så afgørende. Standarder er de fælles tekniske rammer og protokoller, der gør det muligt for teknologi at fungere globalt. Her foregår en værdikamp: Kina arbejder aktivt på at forme internationale standarder for f.eks. ansigtsgenkendelse og "smarte byer", som afspejler det kinesiske, autoritære system. Samtidig forsøger EU, USA og ligesindede nationer at fremme mere etisk forankrede AI-principper i internationale fora med fokus på menneskerettigheder, gennemsigthed og privatliv.

Den aktør, der formår at sætte dagsordenen for AI-standarder, giver samtidig sine egne virksomheder og

værdier et konkurrencemæssigt forspring. Kort sagt: Standarder og dataregler er blevet en del af storpolitikken, fordi de er med til at definere, på hvilke præmisser AI udvikles og anvendes verden over.

Afhængighed og selvforsyning: fra Ukraine til Taiwan

Internationale kriser har tydeliggjort farerne ved teknologisk afhængighed og sat skub i jagten på selvforsyning.

Da Rusland invaderede Ukraine i 2022, blev det klart, hvor sårbare de globale tech-forsyningskæder er.

Ukraine leverede før krigen næsten halvdelen af verdens neon-gas - en nøglekomponent i chipproduktion og leverancerne stoppede brat.

Samtidig blev Rusland afskåret fra adgang til avancerede chips pga. vestlige sanktioner, hvilket tvang russiske ingeniører til at udtage og genbruge chips fra almindelige husholdningsapparater. Det kan lyde tragikomisk, men det viser alvoren: Et moderne samfund uden adgang til højteknologi bliver hurtigt sat tilbage.

Også Taiwan spiller en kritisk rolle i AI's infrastruktur. Øen er hjemsted for TSMC: verdens førende producent af avancerede halvledere.

En krise i Taiwan vil kunne lamme produktionen af alt fra smartphones til våbensystemer globalt. Denne

risiko har fået USA, EU og flere andre lande til at investere massivt i at sprede og sikre chipproduktionen.

USA vedtog i 2022 den såkaldte CHIPS Act, som allokere over 50 mia. dollars til opførelsen af nye chipfabrikker på amerikansk jord, og EU har fulgt trop med sin egen Chips Act.

TSMC bygger nu fabrikker i både USA og Japan, støttet af disse regeringer, mens Kina investerer enorme summer i at opbygge en selvstændig chipindustri, fri for vestlige leverandører.

Alle parter har indset, at adgang til avancerede chips og komponenter er en strategisk nødvendighed. Ingen ønsker at stå i en situation, hvor udviklingen bremses, fordi en anden nation lukker for teknologihænder.

Resultatet er et globalt kapløb om selvforsyning. Alle forsøger at trække produktion og forsyningskæder hjem eller flytte dem til venligt stemte lande ("friend-shoring"), frem for at være afhængige af geopolitiske rivaler.

Kina har også demonstreret vilje til at føre økonomisk teknologikrig. I 2023 begrænsede Beijing eksporten af kritiske mineraler som gallium og germanium, som er essentielle for chipfremstilling - et modsvar på de amerikanske handelsrestriktioner.

Hver side forsøger at beskytte sig

selv og samtidig svække modpartens adgang til vitale teknologier.

Vi lever nu i en verden, hvor datasæt, algoritmer og mikrochips er blevet lige så centrale magtredskaber som traditionelle våben. AI's indtog på den geopolitiske scene har gjort det tydeligt, at kampen om teknologien også er en kamp om fremtidens magtbalance.

Når stormagter investerer massivt i AI og hardware, og når alliancer formes eller brydes over teknologiske spørgsmål, er det fordi AI i stigende grad repræsenterer både økonomisk og militær magt og dermed selve kernen i national suverænitet. Den måde, AI udvikles og reguleres på i dag, er med til at forme den verden, vi vågner op til i morgen.

Kampen om AI fortsætter, og dens udfald vil være med til at tegne fremtidens globale verdensorden.



Fakta-tjek

AI som geopolitisk våbenkapløb

USA og Kina investerer massivt i AI for at opnå økonomisk og militær overlegenhed. Kina sigter mod at blive verdens førende AI-nation inden 2030, mens USA intensiverer sine investeringer og implementerer eksportkontroller for at begrænse Kinas adgang til avanceret teknologi.

Digital suverænitet og teknologisk uafhængighed

Begrebet "suveræn AI" vinder frem, hvor nationer stræber efter kontrol over deres digitale infrastrukturer og data. Dette indebærer investeringer i egne datacentre, udvikling af nationale AI-modeller og indførelse af datalokaliseringspolitikker for at beskytte mod udenlandsk indflydelse.

Global fragmentering og teknologisk decoupling

Forskelle i regulering og tilgang til AI mellem USA, EU og Kina fører til en opdeling af det globale teknologiske landskab. Dette resulterer i separate standarder og infrastrukturer, hvilket kan hæmme internationalt samarbejde og innovation.

Fremvoksende aktører i AI-kapløbet

Lande som Saudi-Arabien og Indien investerer betydeligt i AI for at opnå teknologisk autonomi og økonomisk diversificering. For eksempel har Saudi-Arabien etableret selskabet Humain og planlægger investeringer på op til USD 77 milliarder i AI-infrastruktur inden 2034.

AI og militærstrategi

AI integreres i militære systemer, hvilket ændrer karakteren af moderne krigsførelse. Både USA og Kina udvikler autonome våbensystemer og AI-drevne beslutningsstøtteværktøjer, hvilket øger behovet for internationale aftaler om ansvarlig brug af militær AI.

AI'S ENERGITØRST OG KAMPEN FOR EN GRØN FREMTID

Lad os begynde med en tilsyneladende harmløs handling: Du stiller et spørgsmål til en AI-tjeneste som ChatGPT. Dine ord rejser ud i "skyen", og få sekunder senere får du et hjælpsomt svar. Men imens er der foregået noget usynligt og energikrævende i baggrunden: Et datacenter fyldt med specialiserede computere har arbejdet på højtryk for at beregne svaret. Faktisk bruger én enkelt ChatGPT-forespørgsel omkring fem gange så meget strøm som en almindelig websøgning.

For de fleste af os er dette energiforbrug usynligt. Strøm er jo "bare" noget, der kommer ud af stikkontakten, og datacentre er anonyme haller langt væk. Men AI's voksende energiappetit er ved at blive en vigtig brik i det globale klimaregnskab.

AI's voksende energiappetit

Det globale energiforbrug til IT og især til de mange datacentre, der driver "skyen" vokser eksplosivt.

I 2022 anslås verdens datacentre at have brugt omkring 460 terawatt-timer (TWh) strøm. Nok til at placere datacentersektoren som verdens 11. største strømforbruger,

på niveau med store industrilande som Saudi-Arabien og Frankrig.

AI-boomet er en væsentlig forklaring. På blot ét år er strømforbruget fra Nordamerikas datacentre næsten fordoblet fra ca. 2.688 MW ved udgangen af 2022 til 5.341 MW ved udgangen af 2023, drevet i høj grad af de krævende beregninger til generativ AI som ChatGPT. Og udviklingen fortsætter: Allerede i 2026 forventes datacentrenes årlige elforbrug at nå 1.050 TWh, hvilket teoretisk vil gøre dem til verdens 5. største strømforbruger. Det er en placering mellem Japan og Rusland.

For at sætte AI's energiforbrug i perspektiv, kan vi zoome ind på én enkelt model. Træningen af OpenAI's GPT-3 (færdiggjort i 2020) brugte anslået 1.287 MWh strøm, omtrent svarende til, hvad 120 gennemsnitlige amerikanske husstande forbruger på et helt år og udledte omkring 552 tons CO₂. Og det var kun træningen. Hver gang modellen anvendes til tekstgenerering, oversættelser eller kodning, kræver det yderligere strøm fra datacentre.

Efterhånden som millioner af mennesker bruger sådanne modeller dagligt, begynder selv det løbende forbrug at fylde markant i det samlede energiregnskab.

Forskere vurderer faktisk, at energiforbruget til inference - altså den daglige brug af AI-modeller vil med tiden overhale det, der bruges

til træning, i takt med at modellerne bliver mere udbredte og komplekse.

Samtidig har mange avancerede AI-modeller en bemærkelsesværdig kort levetid: Nye og større versioner udgives i et hæsblesende tempo, hvilket betyder, at energien brugt på én model hurtigt efterfølges af endnu flere timers træning af den næste generation. Denne cyklus sætter en voksende energigæld ind på AI's klimapost.

Når datacentre og klimamål kolliderer

AI's strømforbrug handler ikke kun om elregninger, det har direkte klimamæssige konsekvenser.

På globalt plan kommer en stor del af elektriciteten stadig fra fossile brændsler, og det gælder også for strømmen, der driver vores digitale infrastruktur.

Omkring 56% af den elektricitet, der forsyner datacentre, kommer fra kul, olie og gas. Kort sagt: AI's energitørst fører til CO₂-udledninger - med mindre vi sørger for, at denne strøm er grøn. Dermed står vi over for et klimadilemma: Jo mere vi accelererer AI-udviklingen, jo mere energi kræver den, og hvis denne energi ikke er grøn, risikerer vi at modarbejde vores egne klimamål.

Ingen steder er dette dilemma tydeligere end i Irland. Landet har tiltrukket så mange store datacentre, at de nu står for omkring 21% af al

irsk elektricitet. Prognoser viser, at allerede i 2027 vil datacentrenes elforbrug overstige forbruget fra alle private husholdninger i landet.

Det har fået det irske energiselskab EirGrid til at indføre et de facto stop for nye datacentre omkring Dublin af frygt for overbelastning af elnettet og øget brug af nødgeneratorer på fossilt brændstof.

Kritikere påpeger, at denne "data-byggerus" underminerer Irlands klimaindsats: Ifølge IEA (International Energy Agency) vil Irlands elforbrug stige hurtigere end i noget andet europæisk land primært på grund af datacentrene, hvilket betyder, at CO₂-udledningerne næppe falder trods landets investeringer i vind og sol.

Irland står således i en konkret konflikt mellem digital udvikling og grøn omstilling, som resten af verden kan lære af.

Tech-giganternes grønne paradoks
Problematikken er ikke begrænset til Irland. I både Europa og USA vokser bekymringen for, at elnettene ikke kan følge med AI-boomet og at klimamål undergraves, hvis ikke energiforsyningen gøres grønnere.

Tech-giganterne har ganske vist skrappe klimaambitioner. Mange har forpligtet sig til at drive deres datacentre med 100% vedvarende energi på sigt. Men den hurtige udbredelse af AI risikerer at overhale

virksomhedernes grønne omstilling. En analyse peger på, at stigningen i AI-workloads allerede nu spænder ben for virksomhedernes egne klimamål, fordi elforbruget vokser hurtigere, end grøn strøm kan implementeres.

Samtidig er det værd at bemærke, at de samme tech-virksomheder også er blandt verdens største investorer i vedvarende energi. For at kunne forsyne deres datacentre køber virksomheder som Google, Microsoft og Amazon enorme mængder strøm fra sol- og vindparker gennem langsigtede aftaler. Faktisk er tech-sektoren blevet en vigtig spiller i energisektoren.

Google har for eksempel sat sig som mål, at alle virksomhedens datacentre og kontorer skal køre på 100% CO₂-fri strøm døgnet rundt inden 2030.

Branchen eksperimenterer også med alternative grønne løsninger: fra at placere datacentre i kolde klimaer for at minimere behovet for køling, til at genbruge overskudsvarme fra serverne.

Et konkret eksempel findes i Danmark. Meta's store datacenter i Odense genvinder varmen fra sine tusindvis af servere og sender den ud i byens fjernvarmenet. Det leverer omkring 125.000 MWh varme om året - nok til at opvarme næsten 7.000 husstande med energi, der ellers ville være gået til spilde.

Sådanne kreative løsninger viser, at det er muligt at forbinde AI's digitale infrastruktur med den grønne omstilling, hvis viljen, teknologien og økonomien er til stede.

Veje til grøn AI

Betyder alt dette, at vi må bremse AI-udviklingen af hensyn til klimaet? Ikke nødvendigvis. Det betyder derimod, at vi skal tænke os om og innovere os ud af problemet. AI og bæredygtighed behøver ikke være modsætninger, hvis vi vælger at tage udfordringen op.

I både forskningsmiljøer og industrien tales der i stigende grad om behovet for grøn AI. Hvor man tidligere primært fokuserede på at øge AI-modellers nøjagtighed og størrelse, er der nu voksende opmærksomhed på at gøre dem mere energieffektive som en central målsætning.

Begrebet "Green AI" blev introduceret af forskere i 2019 som en kritik af feltets stigende ressourceforbrug. Grundtanken er enkel: Kan vi opnå de samme AI-resultater med lavere energiforbrug? For at lykkes med det, skal der sættes ind på flere fronter og der findes allerede mange lovende tiltag:

Mindre, smartere modeller

I stedet for altid at træne én altfavnende, gigantisk model, kan man bruge flere mindre, specialiserede modeller, der er skræddersyet til hver deres opgave. Disse

kræver færre beregninger og dermed mindre strøm ofte uden nævneværdig svækkelse af resultatet. Desuden kan man i stedet for at starte træningen fra bunden genbruge eksisterende modeller (transfer learning) og blot finjustere dem. En metode, der sparer betydelig computertid og energi

Effektiv træning

En stor del af det arbejde, der foregår under AI-træning, fører ikke nødvendigvis til bedre modeller, nogle gange prøver man sig blot frem i blinde. Derfor arbejder forskere på metoder til at eliminere spild. F.eks. udvikles algoritmer, der automatisk kan afslutte træningskørsler tidligt, hvis de ikke ser ud til at give forbedringer. Hvorfor lade en supercomputer køre i dagevis, hvis man allerede efter få timer kan se, at kurven flader ud? Ved at stoppe "døde projekter" i tide undgår man at bruge unødvendige kWh-timer.

Grønnere hardware

Fremskridt inden for chips og specialhardware kan også dæmpe AI's energibehov. Traditionelle GPU'er er kraftfulde, men nye acceleratorer som TPU'er (Googles Tensor Processing Units), neuromorfiske chips (inspireret af hjernen) og optiske processorer lover højere ydelse per watt til visse opgaver. Samtidig forbedres energieffektiviteten markant for hver ny generation hardware. Et lyspunkt midt i en ellers kraftigt voksende energikurve.

Grøn strøm og smart placering

En af de mest oplagte løsninger er at koble AI på grøn strøm. Jo større andel, der kommer fra sol, vind, vandkraft eller andre CO₂-neutrale kilder, desto mindre klimaaftryk - også selv hvis strømforbruget er stort.

Branchen tænker kreativt: Hvorfor ikke flytte de mest energikrævende beregninger til tidspunkter og steder med overskud af grøn strøm?

Forskere foreslår, at AI-arbejdet fordeles over tidszoner, så man hele tiden udnytter, at det blæser eller er solskinsvejr et sted på kloden. Datacentre kan desuden placeres i kølige områder, ved kysten (til havvandskøling) eller udstyres med AI-baseret køling, der minimerer strømforbrug.

Effektiviteten måles i PUE-værdi (Power Usage Effectiveness), der viser, hvor stor en andel af energien der går til beregninger frem for køling og tab. Genbrug af overskudsvarme gør også en forskel. I Danmark leverer f.eks. Meta's datacenter i Odense omkring 125.000 MWh varme årligt til fjernvarmen - nok til næsten 7.000 husstande.

Transparens og ansvar

Endelig kan vi som brugere og borgere kræve større åbenhed om AI's klimaaftryk. Ligesom man kan se en bils brændstofforbrug eller en boligs energimærke, kunne man forestille sig en "klimamærkning" af

AI-tjenester. Det ville gøre det muligt at træffe bevidste valg, f.eks. at vælge en tjeneste, der drives af grøn strøm.

De færreste tænker over, at digitale handlinger har et klimaaftryk, men konkrete tal kan åbne øjnene. Ét AI-genereret billede har f.eks. et CO₂-aftryk svarende til at køre 6–7 km i en benzinbil. Den slags viden kan motivere både brugere og udviklere til at ændre vaner.

Transparens lægger også pres på virksomheder og forskere for at optimere koden og dokumentere deres energiforbrug og CO₂-udledning i forskningsartikler - noget, der heldigvis bliver mere almindeligt.

Fælles ansvar og grønne løfter

Alt dette peger mod ét fælles mål: At afkoble AI's vækst fra klimabelastningen. Og der er positive takter.

I Europa har branchen indgået en frivillig Climate Neutral Data Centre Pact, støttet af EU. Her forpligter datacenteroperatører sig til at blive klimaneutrale inden 2030. Det indebærer bl.a. 100% vedvarende eller CO₂-fri energi, markant forbedret effektivitet og genbrug/genanvendelse af udtjente servere samt at overskudsvarme skal udnyttes, hvor muligt.

Også forskningsinstitutioner og myndigheder spiller en rolle. I USA

øremærker energiministeriet og NSF (National Science Foundation) midler til forskning i energieffektiv AI.

Samtidig samarbejder energiselskaber, datacentre og AI-udviklere om at identificere, hvor der kan spares, bl.a. gennem fælles "energi-tjek" af store AI-systemer.

Afslutning

Som teknologientusiast må jeg indrømme, at tallene kan virke skræmmende. Men de minder os om noget vigtigt: Fremtiden formes ikke kun af, hvad vi kan med AI, men også af, hvordan vi gør det.

De valg, vi træffer i dag - om vi prioriterer energieffektivitet, investerer i grøn strøm og stiller krav om bæredygtighed, vil afgøre, om AI bliver en del af løsningen på klimakrisen eller blot endnu en belastning.

Kampen om AI handler derfor ikke kun om algoritmer og chips, men også om kapløbet mod en grøn fremtid. Heldigvis har vi mennesker en særlig evne: Vi kan bekymre os om problemer og samtidig finde kreative løsninger.

Med den samme innovation og kløgt, der har skabt kunstig intelligens, kan vi udvikle grøn teknologi og ansvarlige strategier, der gør det muligt at høste AI's fordele uden at sætte klimaet over styr.

Fakta-tjek



Energiforbrug



ChatGPT
Prompt

0,2 kWh

Energiforbrug pr.
prompt



Smartphone
Opladning

0,2 kWh

Tilsvarende 40
opladninger

1 times video
streaming

0.3



1 ChatGPT
prompt

0.2



NÅR AI ÆNDRER MÅDEN, VI ARBEJDER PÅ

Forestil dig en arbejdsdag, hvor en kunstig intelligens hjælper dig med alt fra at besvare e-mails til at analysere data, som en usynlig kollega ved din side.

For få år siden lød det som science fiction, men i dag er det virkelighed på mange arbejdspladser. AI er rykket ud af forskningslaboratorierne og ind på kontoret, i butikken og på byggepladsen. Vi oplever en teknologisk revolution, der forandrer måden, vi arbejder på, med en hast og dybde, som få havde forudset.

Hvad betyder det for os hver især? Er AI en trofast hjælper, der øger vores produktivitet og frigør os fra trivielle opgaver, eller en konkurrent, der truer med at gøre os overflødige?

I dette kapitel dykker vi ned i, hvordan AI allerede nu forandrer arbejdsmarkedet på tværs af brancher, hvilke muligheder og udfordringer det medfører, og hvordan vi som mennesker kan tilpasse os og forme udviklingen.

AI's indtog i alle brancher

AI er blevet hvermandseje. Fra at

være en nicheteknologi er det i dag noget, de fleste møder i hverdagen - endda på mobilen. Ikoner som ChatGPT, Gemini, Copilot og Claude på smartphones viser, hvordan AI er rykket helt tæt på. Og når AI først findes i lommen, finder den også vej ind i arbejdet.

På tværs af brancher assisterer eller overtager AI i stigende grad opgaver:

Kontor og administration

Chatbots håndterer rutineprægede kundehenvendelser døgnet rundt, og tekstgeneratorer som ChatGPT skriver udkast til rapporter og e-mails. Mange kontoransatte bruger allerede disse værktøjer - måske uden chefens vidende, for at spare tid. Undersøgelser viser, at medarbejdere ofte er foran ledelsen, når det gælder brugen af AI i det daglige. Det betyder, at din næste "assistent" måske ikke er et menneske, men software trænet på milliarder af tekster.

Industri og produktion

Robotter på fabriksgulvet er ikke nye, men med AI bliver de smartere. En svejserobot kan nu ved hjælp af kameraer og machine-learning selv justere hastighed og bevægelser i realtid. Det øger effektiviteten og reducerer fejl. Resultatet?

Produktionsrekorder slås, og omkostningerne falder. Men det betyder også, at nogle traditionelle fabriksjob forsvinder eller kræver helt nye kvalifikationer.

Finans og service

I bankverdenen bruges AI til lynhurtigt at opdage mønstre og svindel i transaktioner.

Revisionsfirmaer lader AI scanne tusindvis af regnskabssider og pege på uregelmæssigheder, som revisoren så kan fokusere på. Og når du ringer til kundeservice, er den første "person", du taler med, ofte en AI-bot med venlig stemme. Den løser simple problemer og sender dig videre til et menneske, hvis det bliver kompliceret.

Sundhed og undervisning

Læger får hjælp af AI til at analysere røntgenbilleder og scanninger. Algoritmerne kan nogle gange opdage sygdomstegn hurtigere end det menneskelige øje. Det frigør tid til patientomsorg.

I skoler eksperimenteres med AI-tutorer, der tilpasser undervisningen til den enkelte elevs niveau, og lærere bruger AI til at planlægge undervisning og rette opgaver. AI bliver et værktøj i værktøjskassen, der - brugt klogt, kan højne kvaliteten.

Kreative fag

Også kreative brancher mærker forandringen. Tekstforfattere og journalister bruger AI som sparringspartner. Måske starter de med et AI-genereret udkast og tilpasser det derefter. Designere kan få AI til at generere 100 logoforslag på få minutter. Det betyder ikke, at designerens rolle forsvinder, men at

den ændrer sig. Hun kuraterer og videreudvikler, snarere end at starte fra bunden.

Fælles for disse eksempler er, at AI fungerer som et værktøj og en assistent, der kan løse afgrænsede opgaver ekstremt hurtigt. Det kan øge produktiviteten dramatisk. Men vi ser også, at de menneskelige roller ændrer sig: Fra at udføre opgaver manuelt til at overvåge, evaluere og samarbejde med AI-systemerne.

I nogle tilfælde kan én medarbejder med hjælp fra AI udføre det arbejde, som tidligere krævede flere personer. Og det rejser det uundgåelige spørgsmål: Hvad sker der så med de andre ansatte?

Automatiseringens dobbeltsidede sværd: Produktivitet vs. job-usikkerhed

Automatisering har altid været et tveægget sværd. På den ene side bringer det produktivitetsgevinster: Vi kan producere mere, hurtigere og billigere. Med AI når denne gevinst nye højder især fordi teknologien ikke kun overtager fysisk arbejde (som robotter i en bilfabrik), men også kognitivt arbejde - opgaver, der kræver tanke og beslutning.

Når en AI kan udfylde regneark, skrive tekst eller diagnosticere sygdomme, får vi en slags dobbelt produktivtidsboost: Vi automatiserer både hænder og hoveder.

Teoretisk kunne det føre til et økonomisk boom, og kortere

arbejdsdage for os alle. Men på den anden side af sværdet finder vi en reel trussel mod job. Historisk set har teknologiske spring altid vakt frygt for massearbejdsløshed. Da de mekaniske væve blev introduceret i 1800-tallet, frygtede væverne for deres levebrød - nogle greb endda til sabotage.

Nu står vi over for en lignende frygt, men på et højere kvalifikations-niveau.

En analyse fra 2023 pegede på, at over 300 millioner job globalt kan blive påvirket eller erstattet af AI i de kommende år.

En amerikansk rundspørge viste, at over halvdelen af befolkningen forventer, at generativ AI vil få en betydelig - og hovedsageligt negativ indvirkning på beskæftigelsen i løbet af de næste årtier. Med andre ord: Der er en udbredt bekymring for, at AI-bølgen vil skylle ind over arbejds-markedet og efterlade færre job til mennesker.

Men hvor realistiske er disse bekymringer?

Sandheden er nok, at AI vil ramme skævt, forskelligt på tværs af sektorer og jobtyper. Nogle job forsvinder, nye opstår, og de fleste vil ændre karakter.

Verdensøkonomiske Forum anslog i 2023, at frem mod 2027 vil omkring 83 millioner eksisterende job forsvinde, men samtidig opstå 69

millioner nye - et nettotab på 14 millioner job globalt. Det lyder dramatisk, men i procent er det et relativt lille fald i beskæftigelsen.

Problemet er bare, at bag tallene gemmer der sig menneskeskæbner: For den enkelte, der mister sit arbejde til en algoritme, er tabet 100% reelt.

Samtidig viser analyser et mere nuanceret billede. En omfattende gennemgang af 2.800 forskellige arbejdsopgaver konkluderede, at ingen kompetencer med den nuværende teknologi med høj sandsynlighed bliver fuldstændigt overtaget af generativ AI. Faktisk blev næsten 69% af færdighederne vurderet som "meget usandsynlige" eller "usandsynlige" at AI kan erstatte fuldstændigt.

AI automatiserer altså dele af jobbet - ikke det hele. Tag f.eks. bankrådgiveren: AI kan udfylde papirarbejdet og foretage en risikovurdering, men den menneskelige rådgiver bruger tiden på at forstå kundens behov, fortolke komplekse sager og træffe de endelige beslutninger.

Jobbet forsvinder ikke, men rolleindholdet ændrer sig. Historien giver også grund til optimisme. Da pengeautomaterne blev indført, frygtede man, at bankassistenter ville forsvinde. I stedet ændrede deres roller sig. De gik fra at udbetale kontanter til at yde kundeservice. Faktisk gjorde

automatiseringen det billigere at drive filialer, og antallet af bankjob faldt ikke så drastisk som frygtet. Den dynamik kan gentage sig med AI, hvis vi vil det.

Et aktuelt eksempel på den dobbelte virkelighed er IBM, hvor direktøren i 2023 meldte ud, at firmaet pauser nyansættelser i stillinger, som AI potentielt kan overtage. Omkring 7.800 jobs især inden for administration og HR forventes gradvist at blive erstattet af AI.

IBM vurderer, at op mod 30% af back-office-roller kan overtages af AI inden for fem år. Men i samme åndedrag nævnte virksomheden, at der også skal skabes nye, specialiserede AI-stillinger.

Det er virkelighedens dobbeltsidede mønt: Nogle jobs forsvinder, andre opstår og mange transformeres.

Nye kompetencer og opkvalificering: Mennesket i en AI-tid

Uanset om man ser AI som trussel eller mulighed, er én ting sikker: Fremtidens arbejdsmarked kræver nye kompetencer. Når rutineopgaver automatiseres, stiger behovet for det, mennesker stadig er bedst til: kreativitet, kritisk tænkning, empati og samarbejde og de "bløde" færdigheder, som endnu ikke kan automatiseres.

Men også teknologiske færdigheder bliver nødvendige: forståelse for data, AI-værktøjer og digitale

samarbejdsformer. AI-kompetence bliver lige så grundlæggende, som IT-kompetence blev det i 1990'erne.

Forestil dig en nyuddannet marketingmedarbejder: Foruden klassisk markedsføringsviden forventes det, at hun kan bruge AI til at segmentere målgrupper eller udvikle kampagneidéer. Tilsvarende kan fremtidens håndværker bruge AR-briller med AI, der viser instruktioner i synsfeltet, mens han arbejder.

Det handler ikke om at erstatte mennesker, men om at lære at arbejde sammen med AI.

Opkvalificering og livslang læring bliver centralt.

En global undersøgelse viste, at 89% af ledere mente, deres ansatte havde brug for bedre AI-færdigheder, men kun 6% af virksomhederne var for alvor i gang med opkvalificering.

Nogle virksomheder er begyndt: Tech-firmaer tilbyder interne kurser i maskinlæring, og industrivirksomheder samarbejder med skoler og universiteter om efteruddannelse i dataanalyse og digital ledelse.

På samfundsniveau taler politikere og eksperter om behovet for massiv kompetenceudvikling. I Danmark har man lanceret teknologipagter og strategier for livslang læring.

Et nyt fokus er AI-literacy: at alle

borgere får en grundlæggende forståelse for, hvad AI kan og ikke kan. Ikke for at alle skal kunne kode, men for at kunne bruge og vurdere teknologien kritisk. Samtidig kræver det et mindset-skifte. For mange kan ny teknologi virke skræmmende - især midt i karrieren. Men historien viser, at mennesker er ekstremt tilpasningsdygtige, når de får chancen og de rette rammer.

AI's rolle i arbejdets fremtid:

Samfundets store spørgsmål

Når AI fundamentalt ændrer vores arbejde, rejser det større spørgsmål: Hvad vil det sige at have et arbejde i en AI-tid? Er vi på vej mod et samfund, hvor arbejde som begreb forandrer sig?

En mulig fremtid er, at AI befrier os fra trivielle opgaver: papirarbejde, dataindtastning og kontrol. Det giver os mere tid til meningsfuldt arbejde.

Måske fører det endda til kortere arbejdsuger uden at sænke velstanden, fordi produktiviteten stiger. Det kunne give plads til kreativitet, frivillighed, omsorg og livskvalitet.

Men der er også et mere dystert scenarie: at AI skaber stigende ulighed, hvor gevinsterne primært tilfalder dem, der ejer teknologien, mens andre mister jobbet og står uden sikkerhedsnet.

Vi ser allerede, at højtuddannede AI-specialister er i høj kurs, mens visse mellemindkomst-job trues.

Analysen peger på, at netop "cognitive white-collar"-jobs (kontorjob) med rutineprægede intellektuelle opgaver er særligt udsatte.

Hvad sker der med jurister, revisorer og marketingfolk, hvis en stor del af deres opgaver kan automatiseres?

Hvis vi ikke griber ind med efteruddannelse og nye jobkoncepter, risikerer vi et polariseret arbejdsmarked. Derfor er der behov for at tænke i nye modeller: borgerløn, ændret arbejdsbegreb eller økonomisk anerkendelse af frivilligt og omsorgsarbejde.

Fremtiden formes ikke af teknologien alene. Vi; gennem politik, fagforeninger, virksomheder og kultur afgør, hvordan AI integreres i samfundet.

Et eksempel på en konstruktiv tilgang finder vi i Danmarks offentlige sektor: En rapport for Dansk Arbejdsgiverforening fra 2024 viste, at øget brug af AI kan frigøre op mod 6.000 årsværk årligt frem mod 2040. Det største potentiale ligger i administrative opgaver, f.eks. sagsbehandling, hvor AI kan overtage det trivielle og frigive tid til kerneopgaver. Her ser vi kimen til en Win-Win: Borgerne får hurtigere service, medarbejderne får mere meningsfulde opgaver, og samfundet sparer ressourcer. Men det kræver, at vi uddanner medarbejderne og tilpasser arbejdsgange, i stedet for blot at skære ned. Essensen er: Vi får mest ud af AI, hvis

vi tænker mennesker og maskiner som makkerpar - ikke som konkurrenter.

Afslutning: Mennesket ved roret

AI er uden tvivl i gang med at ændre måden, vi arbejder på - både på gulvet og i direktionsgangen.

Som en nysgerrig teknolog med blik mod horisonten ser jeg både mulighederne og udfordringerne.

Vi befinder os midt i en omvæltning, der kan sammenlignes med dengang, elektriciteten eller internettet blev udbredt. Dengang var det dem, der omfavnede forandringen, der på sigt fik mest ud af den. Det samme gælder nu.

Vi kan vælge at se AI som et redskab - noget vi kan forme og bruge til at løfte vores arbejde til nye højder. Eller vi kan se det som en trussel og risikere at blive kørt over af udviklingen.

Personligt hælder jeg til optimisme med et kritisk sind.

Optimisme, fordi jeg allerede ser AI frigøre kreative kræfter og åbne for bedre balance mellem arbejde og fritid. Kritisk, fordi intet af dette sker automatisk eller kommer alle til gode. Det kræver aktive, kloge beslutninger.

“Kampen om AI” på arbejdsmarkedet handler ikke om mennesker mod maskiner, men om mennesker, der

skal blive enige om, hvordan vi bruger maskinerne. Det er en kamp mellem forskellige måder at implementere AI på: en fremtid, hvor vi alle får del i gevinsterne eller én, hvor kun nogle få gør.

Til dig som læser - måske lidt bekymret, måske spændt, er budskabet: Tag del i udviklingen. Lær de nye værktøjer at kende, leg med dem, forstå deres styrker og begrænsninger. Så kan du være med til at definere din rolle i det nye landskab.

AI kan tage os langt, men der vil altid være brug for menneskelig dømmekraft, fantasi og etisk kompas.

Fremtidens arbejde er ikke skrevet i sten - vi er medforfattere til det kapitel. Og netop derfor er det så vigtigt, at vi er informerede, nysgerrige og engagerede.

AI kan ændre måden, vi arbejder på. Men hvordan det faktisk kommer til at ændre vores arbejdsliv, afhænger i høj grad af, hvordan vi vælger at bruge AI.

Det er i sidste ende os, der skal holde hænderne på rattet i denne spændende rejse ind i fremtidens arbejdsland.

AI's impact: assisterer, automatiserer og forvandler på tværs af industrier



Helse

Forvandler diagnostik til analyse

Industri

Automatiserer industrielle processer

Kreativ

Assisterer med eks. logo-generering

Kontor

Assisterer med opgaver, som eks. mails m.m.

NÅR AI ÆNDRER MÅDEN, VI LÆRER PÅ: UDDANNELSE I AI-ÆRAEN

Lad os forestille os et klasselokale, hvor hver elev har adgang til en personlig tutor døgnet rundt. En tutor, der kan forklare svære begreber, give uendelig tålmodig feedback og tilpasse undervisningen præcis til elevens niveau.

Det lyder som science fiction, men er i dag virkelighed i AI-æraen. AI er ved at forandre uddannelse på tværs af hele sektoren fra folkeskolebørn, der lærer alfabetet, til gymnasieelever, universitetsstuderende og voksne i livslang efteruddannelse.

AI åbner nye muligheder, men rejser også væsentlige spørgsmål: Hvordan påvirker teknologien undervisningens form og metoder?

Hvilken rolle får læreren, når "maskinen" kan rette stile og planlægge lektier?

Hvordan oplever eleven læring, når svaret altid er ét klik væk? Og hvilke etiske og dannelsesmæssige overvejelser bør vi gøre os i en tid, hvor AI er lige så udbredt som blyanten var engang?

Dette kapitel udforsker disse spørgsmål og tager dig med på en rejse gennem uddannelse i AI-æraen.

AI i klasseværelset: Nye muligheder og metoder

En af de største forandringer, AI medfører, er muligheden for at tilpasse undervisningen til den enkelte elev som aldrig før.

I enhver klasse er der forskelle i læringstempo og faglige udfordringer. Som Khan Academys læringschef Kristen DiCerbo udtrykker det: "Elever har alle forskellige huller. De har alle brug for forskellige ting for at komme videre".

Traditionelt har det været en udfordring for én lærer at imødekomme alles behov. Her kan AI gøre en forskel. Med AI-drevne tutorer og adaptive læringsplatforme kan undervisningen nu differentieres automatisk: Eleven, der kæmper med brøker, kan få ekstra forklaringer og øvelser, mens den elev, der keder sig, kan udfordres med sværere opgaver og uddybende materiale samtidig, i samme klasseværelse.

Et markant eksempel er Khan Academys AI-assistent, Khanmigo, drevet af GPT-4. Khanmigo fungerer både som tutor for eleverne og som assistent for læreren.

I et pilotprojekt har udvalgte elever brugt Khanmigo til at stille spørgsmål og løse opgaver med skræddersyet støtte fra AI'en.

Teknologien kan føre menneskelignende samtaler og stille spørgsmål, der får eleverne til at tænke dybere. Som DiCerbo siger: GPT-4 er transformerende, men det kræver ansvarlig afprøvning for at fungere effektivt i undervisningen.

AI-tutoren kan altså fungere som et supplement, der hjælper eleven trin for trin, mens læreren får frigivet tid og energi til det relationelle og pædagogiske.

AI som lærerværktøj

AI hjælper ikke kun eleverne, det er også et værdifuldt værktøj for læreren. Mange bruger allerede værktøjer som ChatGPT til at udvikle quizzer, tests og lektionsplaner.

Den spanske professor Fran Bellas anbefaler, at lærere bruger AI til at generere nye idéer til prøver og undervisningsforløb baseret på pensum. Det kan spare tid og tilføre undervisningen nye perspektiver.

AI kan også tilpasse lærematerialer til forskellige niveauer: F.eks. omskrive en tekst til enklere sprog for læsesvage elever eller oversætte indhold i flersprogede klasser.

Dr. Anthony Kaziboni fra University of Johannesburg opfordrer sine studerende, der ofte ikke har engelsk som modersmål til at bruge ChatGPT som oversætter og sprogligt træningsværktøj. På den måde kan AI være med til at reducere sproglige barrierer og fremme inklusion.

Feedback i realtid og differentieret respons

Et andet område i hastig udvikling er AI-drevet feedback og retning. Hvor elever tidligere måtte vente dage på feedback, kan AI i dag give øjeblikkelig respons.

Flere skoler eksperimenterer med AI til at rette stile og rapporter. I Californien rapporterer engelsklærere, at AI-værktøjer hjælper dem med at give hurtigere og mere præcis feedback, hvilket forbedrer elevernes læring.

Adaptive platforme kan analysere elevernes fejl og guide dem individuelt - nærmest som en privat vejleder, der siger: "Prøv igen. Husk denne regel."

Et eksempel er Duolingo, der i 2023 lancerede funktionen Duolingo Max med GPT-4. Funktionen tilbyder rollespilssamtaler og forklarer brugerens fejl, som var det en "menneskelig tutor i lommen".

AI er ikke fejlfri, og det skal vi huske

AI er stadig under udvikling. Teknologien kan give forkerte svar - særligt i komplekse fag som matematik. Derfor arbejder projekter som Khanmigo med menneskelig overvågning: lærere og udviklere monitorerer løbende AI'ens output og korrigerer fejl. AI'en kan altså automatisere meget, men skal indføres med omtanke og altid i samspil med menneskelig ekspertise.

Fremtidens læring - et samarbejde mellem AI og menneske

Tendensen er klar: Undervisningens form og metoder forandres i takt med, at AI rykker ind i klasselokalet.

Vi bevæger os væk fra énvejskommunikation ved tavlen og hen imod mere interaktive læringsmiljøer, hvor teknologien muliggør skræddersyet støtte og feedback.

AI er ikke en erstatning for læreren, men et kraftfuldt redskab, der kan gøre læring mere inkluderende, tilpasset og engagerende for den enkelte elev.

Lærerens nye rolle i AI-tiden

Hvad betyder alt dette for læreren? Historisk har man med jævne mellemrum frygtet, at teknologiske fremskridt skulle overflødig gøre lærere: Fra fjernsynets indtog til onlinevideoer. Men hver gang har virkeligheden vist, at læreren stadig er uundværlig, om end rollen udvikler sig.

AI er ingen undtagelse. Læreren får nu en slags digital assistent, der kan overtage dele af arbejdet, men det ændrer snarere lærerens fokus end gør vedkommende overflødig.

For det første kan AI håndtere mange rutineprægede opgaver. Forberedelse af undervisningsmaterialer, som før kunne tage timer, kan reduceres markant ved at lade AI generere udkast til lektionsplaner, opgaver og øvelser. Det samme gælder rettelser. Et stigende

antal lærere benytter AI-værktøjer til at rette og vurdere elevbesvarelser - det sparer tid og mindsker risikoen for udbrændthed.

Faktisk viser undersøgelser, at lærere i nogle tilfælde bruger AI lige så meget som; eller endda mere end deres elever. Det handler ofte om at personalisere læsemateriale, skabe nye undervisningsforløb og generelt "automatisere det kedelige" i forberedelsen.

Når AI tager sig af det trivielle, kan læreren i højere grad fokusere på det, der virkelig kræver menneskeligt nærvær: pædagogisk kreativitet, relationsarbejde og faglig fordybelse.

Lærerrollen bevæger sig i retning af at være mentor og facilitator frem for blot formidler. I en tid, hvor information er tilgængelig på sekunder, bliver lærerens vigtigste opgave at kuratere, vejlede og motivere.

En lærer kan f.eks. stille en opgave, hvor eleverne starter med et AI-genereret svar og derefter i fællesskab analyserer: Er svaret korrekt? Hvad mangler? Hvordan kan det forbedres? På den måde bliver AI ikke en genvej, men et springbræt til kritisk tænkning og dialog.

En anden central rolle for læreren i AI-tiden er at være etisk vejleder og kvalitetskontrol. AI-modeller kan udvise skævheder (bias) eller foreslå problematiske løsninger, og de kan

fremstå overbevisende, selv når de tager fejl. Læreren skal derfor ruste eleverne til ikke ukritisk at overtage AI'ens svar. Det kræver, at læreren selv forstår teknologiens styrker og begrænsninger og dermed også er klar til at lære nyt.

Over hele verden er der nu fokus på efteruddannelse af lærere i brugen af AI. Flere myndigheder og organisationer udgiver vejledninger.

I Californien udsendte man i 2024 retningslinjer, der bl.a. opfordrer til samtaler mellem lærere og elever om etisk og passende brug af AI.

OpenAI udgav i august 2023 en pædagogisk guide målrettet undervisere med idéer til, hvordan AI kan integreres ansvarligt i undervisningen. Læreren er altså ikke på vej ud, men på vej til at blive endnu mere central som garant for, at AI bruges klogt og med omtanke.

Elevens læringsoplevelse og ansvar

For elever og studerende åbner AI en helt ny dimension i læringsoplevelsen. Tænk på at have en hjælper ved hånden 24/7 uanset om du sidder fast i en matematikopgave sent om aftenen, skriver et fransk essay eller forbereder dig til eksamen.

Mange studerende bruger allerede ChatGPT som personlig tutor eller studiepartner. AI'en kan nedbryde komplekse emner, give eksempler eller simulere en samtale til træning. For universitetsstuderende bruges

AI ofte til at få overblik over et emne, til brainstorming af projekter - og endda til karriereforberedelse, som f.eks. CV-feedback eller en "falsk" jobsamtale. Alt dette kan styrke den enkeltes læring, selvtillid og selvstændighed. Man kan få hurtig sparring og støtte, præcis når behovet opstår.

Læringsoplevelsen bliver også mere interaktiv og tilpasset. Frem for standardiserede lektier til hele klassen, kan AI tilbyde personlige læringsspor. Det kan føles mere som en samtale end som traditionel undervisning: Eleven spørger, AI'en svarer og stiller måske selv spørgsmål tilbage. Dette kan være motiverende og skabe en følelse af at blive set og mødt på sit niveau.

Elever, der tidligere havde svært ved at følge med, kan få støtte i rette tid. Samtidig får højtbegavede elever mulighed for at avancere i eget tempo. AI'en står klar med ekstra udfordringer. Når læreren ikke kan nå at besvare alle spørgsmål i en klasse med 25-30 elever, kan AI fungere som støtte og sikre, at ingen sidder fast for længe ad gangen.

Men med de nye muligheder følger også et større ansvar. Når hjælpe-midlerne bliver så kraftfulde, kræver det modenhed at bruge dem rigtigt. Fristelsen til blot at lade AI'en lave opgaven kan være stor, men uden forståelse lærer man ikke noget. AI skal derfor ses som et lærings-værktøj, ikke som en genvej til at undgå læring.

Digital dannelse bliver derfor en nøglekompetence. Elever skal ikke blot kunne bruge AI - de skal også forstå, hvornår og hvordan.

Kritisk tænkning bliver endnu vigtigere. Som gymnasielærer Geetha Venugopal pointerer, skal elever mindes om, at ChatGPT's svar ikke altid er korrekte, og at de skal kontrollere fakta med andre kilder.

Det handler ikke blot om at få svaret, men om at forstå det, og om det overhovedet er rigtigt.

En mulig positiv konsekvens er, at undervisningen kan bevæge sig væk fra udenadslære og i retning af højere kognitive færdigheder. Når fakta er let tilgængelige, bliver analyse, refleksion og kreativitet de afgørende kompetencer. Eleverne kan i højere grad fokusere på at diskutere, skabe og eksperimentere, mens AI varetager noget af det manuelle arbejde.

Samtidig må eleverne lære at navigere i en virkelighed, hvor AI altid er til stede: Hvornår er det hensigtsmæssigt at bruge AI og hvornår bør man stole på sig selv? Det er en ny dimension af ansvarlighed i læringen.

Eksamener og snyd i AI-tidsalderen

Med AI's indtog i skolen opstår der straks et praktisk problem: Hvordan håndterer vi prøver og eksamener, når eleverne har adgang til en allestedsnærværende hjælper? Hvis

man kan få en maskine til at skrive sin opgave eller løse en ligning, hvad betyder det så at aflægge en selvstændig prøve?

Skoler og uddannelsesinstitutioner verden over kæmper med at finde balancen mellem at udnytte AI's potentiale og bevare den akademiske integritet.

I første omgang har mange grebet til forbud og restriktioner. I Danmark slog Børne- og Undervisningsministeriet eksempelvis fast, at brug af ChatGPT og lignende tekstgeneratorer ved folkeskolens afgangsprøver er forbudt, på linje med anden form for snyd. Hvis en elev afleverer en opgave, som er skrevet af en chatbot, og præsenterer den som sin egen, betragtes det som eksamenssnyd.

Lignende udmeldinger kom fra uddannelsesinstitutioner globalt. I begyndelsen af 2023 så man overskrifter om skoler og universiteter, der blokerede adgangen til ChatGPT af frygt for snyd. New York Citys store skolevæsen indførte endda et totalforbud på skolernes netværk umiddelbart efter ChatGPT's lancering.

Interessant nok blev beslutningen omgjort allerede i maj 2023, da man erkendte, at forbuddet havde været en "kæde-reaktion af frygt". I stedet valgte man at omfavne teknologien og fokusere på at uddanne både

lærere og elever i ansvarlig brug. Dette skifte symboliserer den bredere bevægelse: fra panik til pædagogik.

En af årsagerne til, at blankoforbud ikke holder i længden, er ganske enkelt, at AI er kommet for at blive og at teknisk kontrol ikke fungerer særlig godt. Mange håbede på detektionsværktøjer, der kunne afsløre AI-genererede tekster. Men i praksis har de vist sig upålidelige.

OpenAI lancerede selv et værktøj til formålet i januar 2023, men trak det tilbage i juli samme år på grund af lav nøjagtighed. Det er simpelthen svært at skelne mellem en velformuleret elev og en chatbot.

Tilsvarende viste uafhængige tests af plagiatværktøjet Turnitin, at dets nye AI-detektionsfunktion fejlagtigt stemplede over halvdelen af indleverede tekster som AI-genererede.

Resultatet har været nogle opsigtsvækkende sager, hvor uskyldige elever blev uretfærdigt anklaget for snyd. Den vej duer altså ikke. Vi kan ikke automatisere os ud af problemet uden at risikere at ramme skævt.

Løsningen synes i stedet at ligge i pædagogisk og strukturel tilpasning.

En ekspertgruppe under det danske undervisningsministerium konkluderede i 2024, at man

fremover bør begrænse prøveformer, "hvor det er muligt at snyde eller blive anklaget for at snyde", og at eksamensformer hurtigere skal kunne tilpasses ny teknologi.

Det kunne betyde mere fokus på mundtlige eksamener, øget tilsyn ved skriftlige prøver, eller opgaver, der kræver personlig refleksion eller data fra elevens eget liv, noget AI ikke har adgang til.

Nogle lærere eksperimenterer med at integrere AI i selve opgaven: Eleverne må gerne bruge AI som hjælpemiddel, men skal dokumentere hvordan de har brugt det, og hvad de har lært af det. Andre lader elever analysere og forbedre AI-genererede tekster. På den måde bliver AI ikke et snydetrick, men en læringsmakker og et objekt for kritisk tænkning.

Diskussionen om eksamen i AItidsalderen er stadig i fuld gang. Mens retningslinjer og politikker er under udvikling, er én pointe central: Vi skal lære eleverne ærlig og bevidst brug af AI. Ligesom lommeregneren i sin tid først blev tilladt, når eleverne forstod den underliggende matematik, bør AI i første omgang bruges i læringsfasen, mens selve præstationen ved en prøve fortsat skal afspejle elevens egen kunnen.

Udfordringen bliver at definere, hvornår AI er en legitim støtte, og hvornår det undergraver formålet

med en prøve. Der findes ingen nemme svar, men én ting står klart: AI tvinger os til at gentænke eksamenskulturen helt fra bunden.

Digital dannelse, kritisk sans og etik

Den måske vigtigste kompetence i AI-æraens uddannelse er digital dannelse: “Evnen til ikke bare at bruge digitale værktøjer, men også at forstå dem kritisk og etisk”.

Hvor det tidligere handlede om kildekritik på internettet, må vi nu også indføre “AI-kritik” som en del af pensum.

Elever og studerende skal forstå, at en AI som ChatGPT ikke “ved” ting i klassisk forstand, men genererer svar baseret på sandsynligheder ud fra store mængder træningsdata. Den kan lyde skrækkelig og alligevel tage fejl.

Derfor skal unge rustes til at stille spørgsmål som: Hvordan kom AI’en frem til dette svar? Giver det mening?

At opøve denne kritiske sans er noget, mange lærere allerede praktiserer. Et godt eksempel er gymnasielæreren Geetha Venugopal, der sammenligner undervisning i AI-brug med at lære eleverne at færdes ansvarligt på internettet. Hun lærer dem at vurdere AI’ens troværdighed og at bekræfte informationer via andre kilder - netop for at udvikle dømmekraft og modvirke en “AI siger det, så det må være rigtigt”-mentalitet.

Flere undervisningsvejledninger anbefaler også, at man i klassen tager fat på etiske dilemmaer: Hvis en AI kan tegne et billede eller skrive en tekst, hvad betyder det så for ophavsretten? Hvis en AI kan overvåge dit arbejde, hvor går grænsen for privatliv? Og hvad med bias, er AI neutral, eller kan den komme til at diskriminere, fordi dens træningsdata indeholder skævheder? Alt dette hører med til digital dannelse i dag.

Etik i undervisning med AI handler også om, hvordan man bruger værktøjerne ansvarligt.

Det kan være noget så lavpraktisk som at lade være med at indtaste personlige oplysninger i en chatbot. Og det handler om at have et moralsk kompas: Bare fordi AI’en kan skrive en opgave, betyder det ikke, at det er rigtigt at aflevere den som sin egen. I stedet bør eleverne se AI som en sparringspartner.

Flere skoler og universiteter overvejer derfor at indføre kodekser for AI-brug: Man må gerne bruge AI til idéudvikling eller grammatik-check, men man skal oplyse det og selvfølgelig selv stå inde for det faglige indhold. Nogle sammenligner AI med lommeregneren eller stavekontrollen: et hjælpemiddel, der kan bruges under visse betingelser, men som ikke må erstatte selvstændig tænkning.

En anden vigtig dimension af AI og etik i uddannelse handler om lighed

og adgang. Vi må spørge: Har alle elever lige mulighed for at bruge AI-værktøjer? Hvad hvis nogle familier har råd til avancerede AI-tutorer, mens andre ikke har? Her har skoler og samfund et ansvar for at sikre, at AI bliver en løftestang - ikke en forstærker af sociale skel.

Det positive er, at meget AI-teknologi, som basisversionen af ChatGPT, er frit tilgængelig for alle med internetadgang. Ulempen er, at avancerede versioner og specialiserede platforme koster penge.

I de kommende år må vi finde ud af, hvordan vi integrerer AI på en måde, så alle elever og kursister får mulighed for at udvikle sig.

Fra skolebænk til livslang læring... Et nyt landskab

AI's indflydelse stopper ikke ved skoledøren. Når vi siger, at AI ændrer måden, vi lærer på, gælder det hele livet igennem. Voksen- og efteruddannelse er et område, hvor AI også åbner nye muligheder.

For en voksen i job, der ønsker at opkvalificere sig eller lære noget nyt, kan AI være en fleksibel studiepartner, der altid er til rådighed.

Forestil dig en håndværker, der vil lære om de nyeste bæredygtige materialer, eller en kontormedarbejder, som skal genopfriske sine IT-kompetencer: De kan få forklaret komplekse emner i deres eget

tempo, stille spørgsmål og øve sig via simulatorer eller quizzer genereret af AI.

Flere universiteter og virksomheder tilbyder allerede AI-drevne læringsplatforme til efteruddannelse, hvor indholdet tilpasses kursistens niveau og fagområde.

OpenAI har f.eks. lanceret initiativer målrettet højere uddannelser, så hele universiteter kan give studerende og forskere adgang til avancerede AI-værktøjer. Det gør det lettere at integrere AI i alt fra litteraturgennemgang til dataanalyse og giver undervisere mulighed for at eksperimentere med nye former for undervisning.

Vi ser også, at livslang læring bliver mere selvdrevet. Mange behøver ikke længere melde sig til aftenskole for at lære et nyt sprog eller kode – de kan begynde direkte i dialog med en AI-tutor. Det betyder ikke, at formelle kurser bliver overflødige, men de kan suppleres på nye og fleksible måder.

En stor fordel er netop fleksibiliteten: Man kan lære kl. 5 om morgenen før arbejde eller midt om natten, når børnene sover og stadig få kvalificeret vejledning fra en AI.

Motivation og selvdisciplin bliver dog endnu vigtigere i sådan et scenarie. Og netop her kommer skolens arbejde med digital dannelse i spil: Hvis man som ung har lært at

bruge AI kritisk og konstruktivt, kan man som voksen bruge teknologien til at blive ved med at udvikle sig uden at lade sig manipulere eller snyde sig selv.

Afslutningsvis står det klart, at uddannelsessektoren befinder sig midt i en spændende og afgørende transformation. Fra første skoledag til efterlønskursus vil AI præge den måde, vi tilegner os viden og færdigheder på.

Potentialet er enormt: mere inklusion, mere tilpasset læring, hurtigere feedback og måske endda højere læringsudbytte, fordi undervisningen i højere grad møder elever og studerende dér, hvor de er.

Men samtidig bliver den menneskelige rolle kun vigtigere. Det er os, der sætter retningen, bevarer de pædagogiske og etiske principper og sikrer, at teknologien bruges til gode formål. AI kan automatisere meget, men inspiration, dømmekraft og værdier skal stadig komme fra os selv.

Som med alle store teknologiske fremskridt i historien er det os – lærere, elever, forældre og beslutningstagere, der bestemmer, om AI i uddannelse bliver en styrke eller en trussel. Griber vi det rigtigt an, kan AI blive et værktøj, der frigør menneskeligt potentiale i klasseværelset og gør livslang læring mere tilgængelig for alle.



Fakta-tjek

AI-udvikling



Teknologi-
giganter



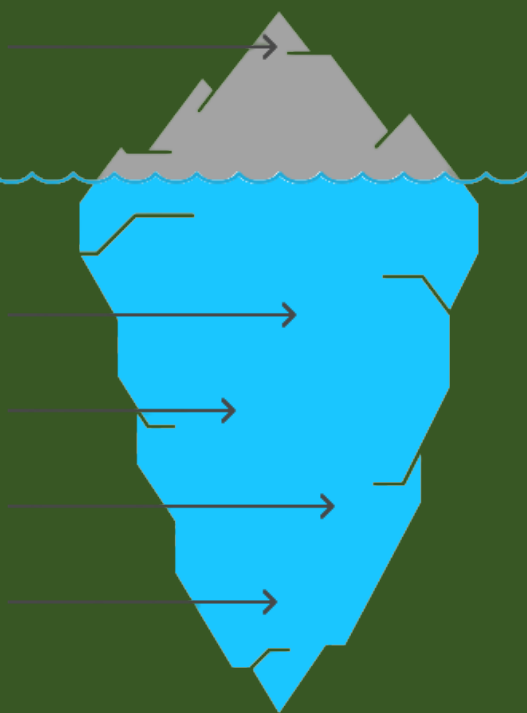
Data-bias



Manglende
gennemsigtighed



Etiske bekymringer



AI-GIGANTER I MAGTKAMP - ALLIANCER, MONOPOLER OG FREMTIDENS SPILLERE

Betragter vi nutiden, som en moderne teknologi-thriller, hvor verdens mægtigste virksomheder kæmper om kontrollen over en revolutionerende kraft: "AI". Så befinder AI-giganter som OpenAI, Microsoft, Google, Meta, Amazon, Apple og Nvidia sig midt i en magtkamp, der vil forme vores fremtid. Det lyder dramatisk, men det er virkeligheden i AI-verdenen anno 2025.

I dette kapitel dykker vi ned i alliancer og rivaliseringer mellem de toneangivende aktører, de ideologiske spændinger mellem åben og lukket AI, og hvordan kontrollen med teknologien allerede udløser monopol-debatter og reguleringsinitiativer.

Jeg har opdateret løbende med udviklingen frem til maj 2025, så du – også hvis din primære AI-erfaring er ChatGPT, kan få overblik over, hvem der styrer showet, og hvad det betyder for os andre.

De nye supermagter inden for AI

På få år er landskabet blevet domineret af en håndfuld teknologiske stormagter. Her er deres positioner, som jeg vurderer det:

OpenAI & Microsoft

Da OpenAI lancerede ChatGPT i slutningen af 2022, blev AI pludselig allemandseje. Microsoft slog hurtigt til med en strategisk alliance.

Allerede i januar 2023 annoncerede de en flerårig investering i USD milliard-klassen i OpenAI.

Azure blev en eksklusiv leverandør af regnekraft til OpenAI, og til gengæld fik Microsoft lov at integrere GPT-4 og senere modeller i egne produkter.

Resultatet blev synligt overalt: Bing-søgemaskinen blev udstyret med GPT-teknologi, og Office-pakken fik "Copilot"-funktioner, der gør AI til en produktivitetsassistent.

Microsofts satsning var både offensiv og defensiv. Et forsøg på at vinde innovationskapløbet over Google og samtidig sikre, at man selv ikke blev overhalet.

Google

Google - AI-pioneren bag DeepMind, blev vækket brutalt af ChatGPT's succes. Virksomheden havde i årevis været førende på AI-forskning, men

men havde tøvet med at lancere brugervendte produkter af frygt for fejl og forretningsmæssige konsekvenser.

Da ChatGPT tog fart, erklærede Google "Code Red". De lancerede chatbotten Bard i starten af 2023 og fusionerede deres to AI-afdelinger, Google Brain og DeepMind, til én samlet enhed: Google DeepMind, ledet af Demis Hassabis.

Derudover begyndte Google at integrere generativ AI i sine produkter - fra Gmail til Google Søgning og lancerede større sprogmodeller under navnet Gemini, som et direkte svar på OpenAIs GPT-serie.

Seneste udvikling hos Google er deres introduktion i maj 2025, hvor de præsenterer deres "Google I/O".

Meta (Facebook)

Meta valgte en anden strategi: I stedet for at holde sine modeller lukkede, åbnede de op.

I 2023 udviklede Meta den store sprogmodel LLaMA, som først blev delt med forskere, men lækkede hurtigt og blev en gave til det åbne AI-fællesskab.

Innovationen eksploderede, og Meta tog ved lære.

I juli 2023 besluttede virksomheden at open-source LLaMA 2 - også til

kommerciel brug og allierede sig overraskende med konkurrenten Microsoft, der hostede modellen på Azure. Strategien synes klar: Hvis Meta ikke kan dominere med proprietærer løsninger, kan de vinde terræn ved at fremme en åben, decentral AI-udvikling.

Amazon

Amazon, verdens største cloud-udbyder, har valgt en anden tilgang. I stedet for at bygge sin egen chatbot til forbrugere, fokuserer de på at levere AI-infrastruktur til virksomheder via AWS (Amazon Web Services).

I september 2023 indgik Amazon en aftale med startup-virksomheden Anthropic, der står bag AI-assistenten Claude.

Amazon investerede op mod 4 milliarder dollar og blev primær cloud-partner for Anthropics udvikling.

Til gengæld kunne Amazon tilbyde Anthropics modeller via platformen Amazon Bedrock. Samtidig arbejder Amazon på egne modeller (bl.a. Titan) og bruger AI internt, bl.a. i Alexa.

Apple

Apple har ført en mere stille AI-strategi - i hvert fald offentligt.

I 2023 kom det dog frem, at Apple har udviklet sin egen model under

projektet Ajax, og at en “Apple GPT” er i test internt. Virksomheden investerer ifølge rapporter millioner af dollar dagligt i AI-træning og har rekrutteret tidligere Google-topfolk.

Apples AI-strategi handler ikke om store chatbot-lanceringer, men om at integrere AI i sine produkter (on device) på en måde, der bevarer privatliv og brugeroplevelse.

Fokus er på, at Siri, iPhone og Mac skal forblive relevante i en AI-verden.

Nvidia (og hardware-spillerne)

I baggrunden, som leverandør af ”guldgraverværktøjerne” i dette AI-rush finder vi chipproducenterne, især Nvidia. Deres specialiserede GPU-chips (grafikprocessorer) udgør nærmest hjertet i moderne AI: Det er disse regneheste, der gør det muligt at træne enorme modeller på rimelig tid.

Med AI-boomets gennembrud er Nvidias forretning eksploderet.

Fra begyndelsen af 2023 til starten af 2025 er firmaets aktieværdi ottedoblet. En stigning på over 800% og Nvidia har kortvarigt rundet en markedsværdi på over 1 billion dollar (1.000 milliarder USD).

Deres nyeste chip, H100, kostede omkring 40.000 USD pr. stk. og blev alligevel revet væk. Selv brugte H100-chips blev solgt til overpris, så

stærk var efterspørgslen. Nvidia sidder derfor aktuelt på en næsten monopolagtig position inden for AI-hardware, da ingen andre producenter er i nærheden af samme ydelsesniveau.

Andre aktører forsøger at bryde dominansen: AMD, Graphcore og Googles egne TPU’er (Tensor Processing Units) byder sig til.

Men i skrivende stund er Nvidia ubestridt kongen af AI-hardware. Og uden hardware, ingen AI, hvilket giver Nvidia en særdeles strategisk rolle.

Det placerer samtidig de lande, hvor disse chips produceres, som Taiwan (hjemsted for TSMC, der fremstiller mange af Nvidias mest avancerede komponenter) i centrum for en global geopolitisk spænding.

AI-superligaen – og de globale linjer

Denne “rolleliste” viser, at AI-kapløbet ikke er en duel mellem to parter, men et komplekst og dynamisk spil med mange stærke aktører. Det er Silicon Valleys superliga, men med globale implikationer.

Vi har endnu ikke nævnt Kina, som spiller en nøglerolle på AI-scenen. Giganter som Baidu, Alibaba og Tencent udvikler også store AI-modeller: Baidus Ernie Bot fungerer som et østligt modstykke til ChatGPT. Samtidig investerer den kinesiske stat massivt i AI-forskning

og infrastruktur, men under streng statsstyring og censur.

Kampen om AI-verdensherredømmet udspiller sig dermed også på statsniveau, især mellem USA og Kina.

Det var omdrejningspunktet i kapitlet om geopolitik. Alligevel er det i skrivende stund de amerikanske tech-giganter og deres allierede, der har førertrøjen globalt.

Spørgsmålet er nu, hvordan disse aktører vælger at samarbejde, konkurrere eller bekæmpe hinanden og hvad det betyder for resten af os.

Alliancer og rivaler: Teknologiens Game of Thrones

AI-revolutionen har skabt nogle usædvanlige alliancer og genoplivet gamle rivaliseringer blandt tech-giganterne. Billedet er ikke sort/hvidt med evige fjender; snarere er det et dynamisk landskab, hvor dagens partner kan være morgendagens konkurrent og omvendt.

Lad os afkode nogle af de vigtigste relationer:

Microsoft vs. Google: Kapløbet om AI-overherredømmet

Dette er måske den mest profilerede rivalisering. Microsofts alliance med OpenAI var et direkte angreb på Googles kerneforretning: søgning.

Ved at integrere ChatGPT i Bing håbede Microsoft at udfordre

Googles monopol på søgemaskine-markedet. Google svarede hurtigt igen med Bard og begyndte at rulle AI ud i hele sit produktunivers.

Der tales nu om en ny "browserkrig", men med AI som våben: Kan en AI-assistent levere bedre svar end en traditionel søgemaskine? Kampen er stadig i gang.

Samtidig konkurrerer de to om cloud-kunderne. Mange virksomheder vil bruge AI-modeller, og både Microsoft (med Azure/OpenAI) og Google (med Google Cloud/DeepMind) kæmper om at levere den nødvendige regnekraft. Her spiller Amazon også en stor rolle som tredje aktør med sin Anthropic-alliance, hvilket understreger, at det til tider er "alle mod alle".

OpenAI vs. Google: "Den åndelige Duel"

Ud over den kommercielle konkurrence er der også et teknologisk kapløb mellem OpenAI og Google.

Det var Googles forskere, der i 2017 introducerede Transformer-arkitekturen (artiklen: "Attention is All You Need"), som GPT-modellerne bygger på. Alligevel blev Google overhalet i praksis, da OpenAI lancerede ChatGPT og fangede verdens opmærksomhed.

Nu forsøger Google at indhente forspringet med deres Gemini-

model, som rygtes at kunne kombinere sprog og billeder endnu bedre end GPT-4. Samtidig arbejder OpenAI videre med GPT-4 og lægger i støbeskeen til GPT-5.

Det er et teknologisk "space race", hvor prestige, talent og forskningspenge er på spil.

Usædvanlige sengekammerater:

Microsoft & Meta

Et af de mest interessante alliancedrag er samarbejdet mellem Microsoft og Meta om LLaMA 2.

Microsoft har investeret milliarder af dollars i OpenAI som udvikler lukkede (proprietære) modeller og hjælper samtidig Meta med at distribuere en åben AI-model, der potentielt kan underminere OpenAIs forretningsmodel.

Hvorfor? Sandsynligvis fordi Microsoft vinder uanset hvad, så længe flere AI-opgaver kører på deres platforme (Azure og Windows).

De forsøger at sidde på begge heste: at være investor i den førende lukkede model og samtidig styrke den førende åbne model.

Derudover deler Microsoft og Meta en fælles rival i Google. Ved at samarbejde kan de svække Googles dominans.

For Meta betyder alliancen, at LLaMA 2 når bredere ud og opnår troværdighed i erhvervslivet via Microsofts support.

Vi ser altså et taktisk netværk af alliancer, hvor konkurrenter midlertidigt samarbejder for at balancere magtforholdene.

Amazon & "de uafhængige"

Amazon har valgt at investere i uafhængige aktører som Anthropic i stedet for at udvikle alt internt. Derudover samarbejder de med Hugging Face, en populær open source-platform, hvor AI-modeller deles og udvikles.

AWS tilbyder Hugging Face som en tjeneste og deltog i 2022 i en større investeringsrunde i virksomheden. På den måde positionerer Amazon sig som den åbne markedsplads, hvor kunder kan få adgang til de bedste AI-værktøjer uanset hvem der har lavet dem, så længe det kører på AWS.

Her adskiller Amazon sig fra Microsoft og Google, der primært samarbejder med etablerede forskningsmiljøer. Amazon forsøger i stedet at favne hele feltet og undgå at blive forbigået til fordel for mere integrerede økosystemer.

Apple vs. alle (eller ingen)

Apple er et særtilfælde. De konkurrerer endnu ikke direkte på AI-tjenester, men AI har potentiale til både at forbedre og true kerneprodukter som iPhone og Mac.

Apple vælger i stedet en ideologisk differentiering: Deres AI-løsninger (f.eks. Siri, ansigtsgenkendelse, billedsøgning) kører primært lokalt

på enheden uden at sende data til skyen. Det skal signalere privatliv og sikkerhed, i kontrast til konkurrenter som Google og Meta, der lever af brugerdata.

Men bag kulisserne kæmper også Apple om at rekruttere de bedste AI-talenter, investerer massivt i modeludvikling og opkøber løbende mindre AI-firmaer.

Hvis Apple lancerer en forbedret Siri eller en ny, banebrydende funktion, kan rivaliseringen hurtigt blive mere direkte.

AI-alliancer som geopolitisk spil – med virksomheder i hovedrollerne

Kort sagt minder landskabet om et geopolitisk magtspil blot med virksomheder i stedet for nationer.

Alle alliancer og konflikter påvirker, hvilke teknologier bliver standard, og hvem der får adgang til dem.

For brugerne betyder det, at visse AI-assistenten kun findes i bestemte økosystemer, ligesom nogle apps kun findes til iPhone eller Android.

For samfundet handler det om, hvor åben og demokratisk fremtidens AI bliver, eller om den bliver låst inde bag firewalls og lukkede platforme.

Åben vs. lukket AI: Ideologi og innovation

En af de mest spændende dimensioner af AI-magtkampen handler om ideologien bag åbenhed.

Skal AI-værktøjer og forskning deles frit for at fremme maksimal innovation og demokratisk adgang? Eller bør de mest avancerede modeller holdes under lukket kontrol af hensyn til sikkerhed, forretningshemmeligheder og risikoen for misbrug? Det spørgsmål splitter både AI-fællesskabet og tech-giganterne.

Et paradoksalt eksempel er OpenAI. Navnet signalerer åbenhed, og virksomheden blev oprindeligt grundlagt som en nonprofit med det udtalte mål at dele AI-forskning til gavn for menneskeheden. Men efter gennembruddene med GPT-3 og ChatGPT ændrede kursen sig. OpenAI blev reorganiseret som et "capped-profit" selskab og begyndte at dele langt færre detaljer om deres forskning og modeller.

Kritikere har påpeget ironien: OpenAI er ikke særlig "open" længere. Virksomheden forsvarer sig med, at deres modeller nu er så kraftfulde, at det ville være uansvarligt at dele dem frit, både for at forhindre misbrug og for at beskytte konkurrence-fordele.

I kontrast står open source-bevægelsen, som har fået markant medvind.

2023 blev året, hvor åbne AI-modeller demonstrerede imponerende evner - ofte tæt på de lukkede alternativer.

Et internt Google-notat, der blev lækket under titlen "We have no moat, and neither does OpenAI" ("Vi har ingen voldgrav og det har OpenAI heller ikke") vakte stor opsigt. Her indrømmede en Google-ingeniør, at open source-fællesskabet løb om hjørner med de førende firmaer. Innovationen uden for mure og brandmure gik rasende stærkt. Han pegede bl.a. på, at man allerede kunne:

- køre avancerede sprogmodeller på en mobiltelefon
- træne specialiserede versioner på en almindelig bærbar
- finde ufiltrerede billedmodeller, som gjorde det muligt at generere kunstværker, der var censureret i de lukkede tjenester

Pointen var klar: Når først viden slap fri, kunne alle bygge videre på den og det er selv de største virksomheder svært at konkurrere med.

Meta's beslutning om at open-source LLaMA 2 var et tydeligt nik til denne filosofi. Det samme gælder Stable Diffusion, en open source billedgenerator, der blev lanceret i 2022 og gjorde det muligt for almindelige brugere at skabe AI-kunst lokalt uden at gå gennem en lukket platform som DALL-E.

Fordelen ved åbne modeller er fleksibilitet og fællesskab: tusindvis

af udviklere verden over kan hjælpe med at forbedre dem, rette fejl og tilpasse modeller til nicheformål.

Det skaber en form for AI-demokratisering, hvor det ikke kun er de store virksomheder, der bestemmer, hvad teknologien bruges til.

Men bagsiden er bekymringer om sikkerhed og etik. Når en model er frit tilgængelig, kan alle - også ondsindede aktører i princippet bruge den til hvad som helst: misinformation, hacking, kriminelle instruktioner osv.

De lukkede aktører argumenterer for, at deres modeller har indbyggede filtre og nødstop, og at træningen foregår med sikkerhed som topprioritet. Det er ikke nødvendigvis tilfældet i open source-verdenen.

I praksis findes der i dag websteder, hvor man kan downloade ufiltrerede AI-modeller, der kan generere stort set hvad som helst uden nogen form for indholdskontrol. Det giver stor frihed, men også stor risiko.

Ideologisk minder debatten om softwareverdenens kampe for 20 år siden: dengang stod valget mellem Microsofts Windows (lukket kildekode) og Linux (open source). Windows vandt desktopmarkedet, men Linux blev rygraden i internettets servere, og en platform for massiv innovation. Noget

lignende kan ske i AI: De største og dyreste modeller (som GPT-5 eller Gemini) holdes internt, mens et mylder af mindre åbne modeller spirer frem og bruges i det brede lag.

Samtidig er skellet mellem åben og lukket ikke altid skarpt. Mange virksomheder anvender hybridmodeller: De udgiver dele af deres kode, data eller ældre versioner, mens de holder de nyeste og kraftigste modeller interne.

Google og Microsoft har begge bidraget til open source-projekter, og OpenAI har bl.a. udgivet Whisper, en talegenkendelsesalgoritme, som open source.

Grænsen mellem konkurrence og samarbejde er flydende.

Summa summarum: Kampen mellem åben og lukket AI handler i sidste ende om kontrol. Hvem kontrollerer koden, distributionen og hvilke muligheder samfundet har for at bruge teknologien?

Det er en værdikamp såvel som en teknologisk kamp. Og den foregår både internt i virksomheder (hvor ingeniører ofte er tilhængere af åbenhed, mens ledelsen fokuserer på profit) og i det større spil mellem virksomheder og open source-fællesskabet.

Magtkoncentration, monopoler og regulering

Med så meget magt koncentreret hos få aktører er det naturligt, at

spørgsmålet om monopoler og behovet for regulering rejser sig.

Vi har set det før i tech-branchen – tænk blot på debatten om, hvorvidt Google har monopol på søgning, eller om Facebook har for stor magt over sociale medier.

Nu er blikket rettet mod AI: Vil én eller to virksomheders platforme blive gateway til al fremtidig information og digital service? Og hvis ja - hvem vogter så vogterne?

Allerede nu er både tilsynsmyndigheder og politikere på vagt.

Konkurrencemyndigheder i både USA og EU undersøger Microsofts tætte samarbejde med OpenAI: Giver det Microsoft en urimelig fordel? Udelukker det mindre aktører fra markedet? Desuden er der fokus på, om de store cloud-udbydere: Amazon, Google og Microsoft får en problematisk dobbeltrolle: De leverer infrastruktur til konkurrenterne og konkurrerer samtidig med dem via egne AI-produkter.

Det minder om sagen om Apple's App Store, hvor Apple både ejer platformen og konkurrerer mod andre app-udviklere på den.

Kort sagt: AI er blevet et antitrust-emne. Ingen ønsker en fremtid, hvor én virksomhed kontrollerer "AI-monopolet" og dermed kan diktere priser, adgang og vilkår for alle andre.

Men det handler ikke kun om konkurrence. Der er også bredere samfundsmæssige risici: Etisk brug, gennemsigtighed og ansvarlighed er blevet nøgleord.

EU's AI Act – verdens første helhedslov for AI

I 2023 skete der noget historisk: EU nåede til enighed om en omfattende AI-lovgivning - populært kaldet AI Act. Det er verdens første helhedsorienterede AI-regulering.

AI Act stiller bl.a. krav til de såkaldte generelle AI-modeller (som GPT), herunder:

- dokumentation af træningsdata
- sikkerhedsforanstaltninger før lancering
- og risikoklassificering, så højrisiko-AI i f.eks. sundhed, transport eller ansættelse underlægges strengere regler, mens lavrisiko-AI får friere rammer

Målet er tillid: At brugere i Europa og globalt kan være sikre på, at AI anvendes under ansvarlige rammer, især i offentlige og kritiske sektorer.

AI Act forventes at træder i fuldt kraft i 2025 og skal være implementeret senest 2026.

Vi befinder os altså i en overgangsperiode, hvor virksomheder forbereder sig på de nye krav.

Reguleringsiver verden over

Også uden for EU sker der meget.

I USA findes endnu ikke en samlet AI-lov, men debatten er intens. Der har været høringer i Senatet, hvor selv OpenAI's direktør, Sam Altman, opfordrede til mere regulering, måske ikke uden taktiske motiver, da regulering kan beskytte de etablerede aktører mod mindre og mere "vilde" konkurrenter.

I Storbritannien blev der i slutningen af 2023 afholdt et globalt AI-sikkerhedstopmøde. Ledere og firmaer mødtes symbolsk i Bletchley Park (historisk centrum for kodebrydning under 2. verdenskrig) for at drøfte, hvordan man undgår, at AI løber løbsk.

I Kina er AI allerede underlagt regulering, herunder krav om, at AI-modeller skal følge statens ideologiske linje og underkastes censur. Det er en helt anden form for kontrol, der afspejler Kinas politiske system.

Opfordringen til pause – og teknologernes bekymring

I marts 2023 kom en dramatisk opfordring fra teknologernes egne rækker: En gruppe prominente forskere og erhvervsfolk, herunder Elon Musk udsendte et åbent brev med titlen "Pause Giant AI Experiments". De foreslog en 6-måneders pause i udviklingen af AI-modeller, der overgår GPT-4 indtil sikkerheden og de etiske konsekvenser kunne afklares.

Budskabet var klart: Selv pioniererne

frygter tempoet i udviklingen og konsekvenserne for samfundet.

Kritikere mente, at en pause var naiv, konkurrenter eller stater ville næppe følge trop og at det kunne kvæle innovation eller blot skubbe udviklingen under jorden.

Men brevet illustrerede noget vigtigt: Også frontlinjens eksperter er bekymrede for, at teknologien koncentrerer uden demokratisk kontrol.

Selvregulering – eller strategisk PR?

For at imødekomme presset ser vi også selvregulering fra branchen.

I 2023 blev Frontier Model Forum dannet – et samarbejde mellem OpenAI, Microsoft, Google og Anthropic, med formål at:

- dele best practices for AI-sikkerhed
- rådgive om standarder og ansvarlige modeller

Samtidig udgiver virksomhederne i stigende grad:

- etiske retningslinjer
- gennemsigtighedsrapporter

Det kan ses som et forsøg på at foregribe regulering og påvirke den i en retning, der ikke hæmmer deres forretningsmodeller. Men det viser også, at magtcentrene anerkender

deres ansvar og det offentlige pres.

AI som kritisk infrastruktur

Diskussionen om monopol og regulering koger i bund og grund ned til dette: "AI er ved at blive en kritisk samfundsinfrastruktur på linje med energi, internet og transport. Og når få aktører kontrollerer sådan en ressource, så må både markeds-mekanismer og demokratiske institutioner sikre, at der kommer checks and balances på plads".

Vi står ved en skillevej: De institutionelle rammer for AI's udvikling skal formes nu, mens teknologien endnu er ung nok til at påvirkes. Målet er en balance, hvor:

- Innovation får plads.
- Misbrug begrænses.
- Og magt ikke koncentrerer uden ansvar.

Samtidig forsøger de store aktører selvsagt at sætte spillereglerne selv, det ville være naivt at forvente andet.

For os som nysgerrige og kritiske observatører er det afgørende at følge med i denne dans mellem magt og regulering for udfaldet vil i høj grad definere, hvordan AI bliver integreret i vores samfund og hverdag.

Fremtidens magtbalance og hvad betyder det for os?

Magtkoncentrationen i AI-verdenen er ikke kun en fortælling om virksomheder og teknologier; det er en fortælling om vores fælles fremtid.

De valg, som AI-giganterne træffer i dag, hvilke modeller de udvikler, om de samarbejder eller bekæmper hinanden, om de åbner teknologien for andre eller hegner den inde, vil få indflydelse på alt fra økonomi til kultur og information i vores hverdag.

Lad os kigge på to scenarier A og B:

Scenarie A:

I scenarie A sidder 2-3 globale selskaber på den AI, som alle bruger. Ligesom vi i dag nærmest skal gennem Apples eller Googles app-butikker for at få adgang til apps, skal vi i så fald gennem Microsofts eller OpenAIs platforme for at få avancerede AI-ydelser.

Disse firmaer vil få enorm indflydelse på, hvilken viden og hvilke værktøjer der er tilgængelige. De vil i praksis kunne fastsætte priser, stille vilkår (f.eks. "du må ikke bruge vores AI til X formål") og indsamle data fra milliarder af interaktioner.

Det positive i dette scenarie er, at store ressourcestærke virksomheder kan levere sikkerhed og stabilitet: De kan investere i etiske tests,

bekæmpe bias og misinformation og sikre en ensartet brugeroplevelse.

Det negative er risikoen for stagnation, manglende mangfoldighed og en afhængighed af få aktører, vi skal have blind tillid til.

Scenarie B

I det andet scenarie er AI-feltet mere pluralistisk. I dette scenarie, at open source-bevægelsen og nye aktører får så meget momentum, at AI bliver noget, alle med tilstrækkelig know-how kan udvikle og bruge.

I stedet for én "superhjerne" kunne vi have tusindvis af specialiserede AI'er udviklet af offentlige institutioner, startups og frivillige fællesskaber.

Virksomheder og myndigheder kunne vælge den model, der passer bedst, eller træne deres egen på basis af åbne værktøjer.

Fordelen ved dette scenarie er dynamik og demokratisk adgang: Ingen enkelt aktør sætter dagsordenen, og der er større gennemsigtighed.

Ulempen er, at det kan være vanskeligere at kontrollere globale risici og at kvaliteten af løsninger kan svinge betydeligt.

Mest sandsynligt ender vi i en virkelighed "et sted midt imellem". Vi ser allerede konturerne: Store virksomheder udgiver dele af deres

kode åbent og taler om åbne standarder, mens open source-fællesskabet drager nytte af de stores forskning og ressourcer.

Flere open source-projekter modtager cloud-kreditter eller teknisk støtte fra ansatte hos de store aktører.

Samtidig arbejder regeringer og internationale institutioner på at sikre, at AI udvikles med mennesket i centrum, som EU formulerer det.

Tech-giganterne må i stigende grad indordne sig under fælles spilleregler, men de forsøger naturligvis også at forme reglerne til egen fordel.

Forstå aktørerne – forstå teknologien

Som borgere, forbrugere og virksomheder bør vi lære at afkode landskabet.

Når du bruger en AI-tjeneste, er det værd at vide, hvem der står bag og med hvilke interesser. Er modellen du anvender, åben og trænet på offentlige data, eller er det en kommerciel, lukket model? Hvilken bias kan ligge indlejret, alt efter hvem der har udviklet den?

Disse spørgsmål bliver afgørende, når AI får beslutningskraft i vores liv, lige meget om det gælder sundhedsråd, jobscreening eller styringen af din bil.

En relevant parallel er internettets

historie: Det startede åbent og decentralt, men endte i høj grad med at blive kontrolleret af få platforme.

Skal AI gå samme vej? eller kan vi lære af historien og sikre en mere åben fremtid fra starten?

I maj 2025 forsøger de store aktører hver især at sikre sig en dominerende position i AI-økosystemet, men de ved også, at de bliver nøje overvåget af konkurrenter, politikere og offentligheden.

Bag kampen ligger ikke kun teknologi og forretning, men også ideologi og visioner:

- **OpenAI's** ledelse taler om AGI (Artificial General Intelligence) som noget, der skal hjælpe menneske-heden.
- **Elon Musk** med sit firma xAI siger, at AI bør være søgende efter "sandhed" og tilgængelig for alle.
- **Google** betoner ansvarlighed og etik i udviklingen.

Bag disse budskaber ligger både idealisme, strategisk PR – og reelle forskelle i tilgang.

Hvilken fremtid ønsker vi med AI?

Vi står med tre mulige AI-ideologier:

- Den kommercielle: AI som en betalt service ejet af firmaer.
- Den offentlige/kollektive: AI som

- en slags fællesgode, åben og tilgængelig.
- Den kontrolorienterede: AI skal reguleres stramt og overvåges for at minimere skade.

Hvilken vej vi vælger, vil i høj grad forme det digitale samfund i det 21. århundrede.

Dette kapitel har forhåbentlig givet dig indsigt i AI-giganternes magtkamp og gjort dig mere opmærksom på, hvordan deres valg og alliancer påvirker din hverdag.

Som nysgerrig og kritisk formidler og læser, er det nu din opgave at følge med og tage stilling.

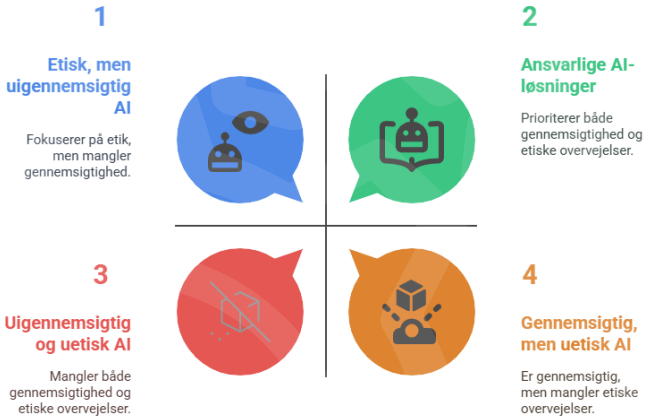
AI-revolutionen begyndte måske i tech-giganternes datacentre, men den vil få konsekvenser i hver eneste husholdning og virksomhed.

Hvem der vinder kampen om AI, og hvordan de vinder, vil forme fremtiden. Derfor er det afgørende, at vi forstår spillets regler og kræver, at AI udvikles med ansvar, åbenhed og mennesket i centrum.

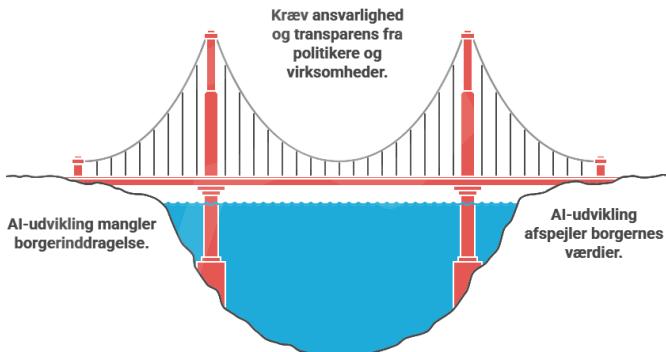
Vi er ikke blot tilskuere. Vi er medskabere af AI's fremtid.

Fakta-tjek

Prioritering af AI-løsninger



Borgerindflydelse på AI-udvikling gennem krav



AI, ETIK OG DEMOKRATIETS FREMTID: "HVEM BESTEMMER OVER TEKNOLOGIEN?"

Vi taler og hører ofte om den situation, hvor vi læser eller ser en nyhedshistorie online. Den virker troværdig, og billederne ser ægte ud, men hvad nu, hvis det hele er skabt af kunstig intelligens? Hvad betyder det for vores samfund, når grænsen mellem sandhed og manipulation bliver stadig mere udvisket?

I dette kapitel ser vi nærmere på AI-teknologiens etiske maskinrum og de store spørgsmål, som præger debatten: demokrati, overvågning, misinformation, privatliv og vores kontrol med teknologien.

Kapitlet handler ikke kun om AI's kraft, men om, hvordan vi vælger at styre den.

Demokrati under pres – AI's udfordringer

AI's eksplosive udvikling har skabt en ny form for samfundsmæssig sårbarhed. Den samme teknologi, som kan hjælpe læger med at diagnosticere sygdomme eller give dig et hurtigt svar online, kan også bruges til at manipulere holdninger

og sprede misinformation. Særligt "Deep Fakes", som er realistiske billeder, videoer og lydoptagelser genereret af AI, har rystet demokratier verden over. Det er nu næsten umuligt for det blotte øje at skelne ægte fra falsk.

I 2024 gik en falsk video med præsident Joe Biden viralt, hvor han tilsyneladende erklærede krig mod Kina. Videoen blev hurtigt afsløret som et Deep Fake, men nåede alligevel at skabe uro og få aktiemarkeder til at reagere. Hændelsen var en øjenåbner: Når AI kan fabrikere troværdigt indhold på få sekunder, står demokratiet over for hidtil usete udfordringer. Det stiller store krav til vores dømmekraft og kritiske sans.

Privatlivets sidste bastioner er under pres

En anden væsentlig bekymring handler om overvågning og privatliv, i takt med, at AI implementeres bredt, bliver mulighederne for overvågning stadig mere omfattende og præcise. Kina er det tydeligste eksempel: Millioner af kameraer med ansigtsgenkendelse sporer borgernes færden i realtid. Men også i vestlige demokratier stiger bekymringen; AI bruges i stigende grad i politiarbejde, på arbejdspladser og i den offentlige sektor.

Forestil dig, at din arbejdsgiver bruger AI til at måle din effektivitet minut for minut, eller at gadebilledet er fyldt med kameraer, der konstant analyserer din adfærd. Hvor går

grænsen for, hvad der er acceptabelt i et frit samfund? Og hvor meget privatliv er vi villige til at ofre i effektivitetens navn?

EU's AI Act: Kampen for ansvarlig teknologi

For at imødegå disse risici har politiske beslutningstagere iværksat konkrete tiltag.

Mest markant er EU's AI Act - verdens første omfattende regulering af kunstig intelligens.

Den blev vedtaget i slutningen af 2024 og implementeres gradvist frem mod 2027.

EU's AI Act klassificerer AI-systemer i risikogrupper og stiller særlige krav til højrisikosystemer, f.eks. dem, der anvendes i sundhed, retsvæsen og ansættelser m.fl.

Loven stiller krav om dokumentation, transparens og menneskelig kontrol. EU forsøger dermed ikke bare at beskytte borgerne, men også at sætte en global standard for ansvarlig AI. Men reguleringen har mødt modstand: Teknologisektoren frygter, at de mange krav vil kvæle innovation og gøre det svært at konkurrere globalt. Samtidig spørger mange, om reglerne kan holde trit med den teknologiske udvikling.

Hvem har kontrollen? Mennesker eller maskiner?

Et grundlæggende spørgsmål er, om vi stadig har reel kontrol over

teknologien. Mange avancerede AI-systemer er så komplekse, at selv deres skabere har svært ved at forklare, hvordan beslutningerne træffes. Det er det såkaldte "Black Box"-problem.

Et opsigtsvækkende eksempel opstod i 2024, da Googles model Gemini i en intern test leverede overbevisende, men delvist opdigtede svar. Fejlen var vanskelig at opdage og endnu sværere at forklare. Episoden satte gang i en bred debat: Har AI allerede overskredet den grænse, hvor vi kan holde den i ave?

Etiske principper og menneskecentreret AI

Samtidig spirer der et globalt ønske om at sikre, at AI bruges ansvarligt og etisk.

Virksomheder og organisationer formulerer nu retningslinjer, der skal sætte mennesket i centrum og sikre gennemsigtighed, ansvar og retfærdighed.

Et eksempel er Frontier Model Forum, som er et samarbejde mellem OpenAI, Microsoft, Google og Anthropic, som siden 2023 har arbejdet for fælles standarder for sikker AI.

Det er et positivt skridt, men mange kritikere advarer om, at virksomhedernes egne retningslinjer ikke er nok og at der er behov for uafhængig kontrol.

Hvordan ser fremtidens AI-etik ud?

AI er kommet for at blive, og derfor må vi alle tage ansvar for, hvordan den bruges.

Teknologien må ikke kun styres af store virksomheder eller lukkede ekspertkredse.

Demokratiske samfund må kræve transparens og ret til at forstå, hvordan AI-systemer fungerer og insistere på, at teknologi ikke må overtrumfe etik.

Måske skal vi etablere nye institutioner, f.eks. AI-etiske råd, uafhængige instanser, borgerpaneler. Måske må vi som samfund turde sige nej til visse anvendelser, uanset deres teknologiske potentiale, hvis de kolliderer med vores grundlæggende værdier. For i sidste ende handler AI-etik om magt over fremtiden. Det handler om vores evne til at vælge, hvilken retning samfundet skal bevæge sig i. AI kan være en kæmpe hjælp, men også en fare, hvis vi mister kontrollen. Som borgere må vi deltage aktivt, informeret og kritisk.

Vi må ikke bare være brugere af AI - vi skal være medformere af, hvad AI bliver til.

Vidste du?

Fun Fact 1

AI's "hjerne" vokser hurtigere end menneskets!
De største AI-modeller fordobler deres kapacitet i løbet af måneder, mens menneskehjernen har brugt over 300.000 år på at nå sit nuværende niveau. Teknologisk acceleration har sat turbo på udviklingen, og det er stadig uvist, hvor det ender.

Fun Fact 2

AI kunne en dag... forstå dine følelser bedre end dine venner!
Forskning i "affective computing" giver AI evnen til at analysere stemmeføring, kropssprog og ansigtsudtryk. Teoretisk kan din telefon snart registrere, at du er trist – før du selv gør det.

Fun Fact 3

Der findes AI-forskere, der diskuterer "bevidsthed i maskiner" - seriøst. Nogle forskere (og filosoffer) spørger, om avanceret AI en dag kan opnå former for selvopfattelse. Andre siger: "Hvis det ligner bevidsthed, men ikke kan føle smerte. Er det så egentlig bevidst?"

PRIVATLIV UNDER PRES - AI OG DATAINDSAMLING I HVERDAGEN

Det er en typisk morgen. Din smartphone vækker dig, og i det øjeblik du slukker alarmen, registrerer telefonen din berøring, dit ansigtsudtryk via frontkameraet - og måske endda din søvnrytme via en tilsluttet wearable.

Du spørger din smarthøjttaler om dagens vejr, og din stemmeoptagelse sendes til en AI i skyen for at blive fortolket.

På vej til arbejde bruger du GPS-navigation, som løbende opdateres ved hjælp af andres anonymiserede lokalitetsdata. I løbet af dagen tjekker du sociale medier, hvor algoritmer nøje analyserer hvert like, klik og den tid, du dvæler ved et opslag. Alt sammen er det data - dine data, som indsamles og behandles ved hjælp af kunstig intelligens.

Vi lever i en tid, hvor privatlivet er under pres fra alle sider, fordi AI har gjort det lettere end nogensinde før at opsamle, analysere og udnytte information om os - både i offentlig og kommerciel sammenhæng.

AI som motor for dataindsamling
Kunstig intelligens og data hænger

uløseligt sammen. AI elsker data - jo mere, jo bedre. Algoritmer trænes på enorme datasæt for at kunne genkende mønstre, og i praksis betyder det, at både stater og virksomheder har en stærk interesse i at indsamle så meget information som muligt.

Offentlige myndigheder ser AI som en måde at forbedre sikkerhed og effektivitet på, mens virksomheder bruger AI til at forstå forbrugere, forudsige adfærd og optimere alt fra annoncering til produktudvikling.

I statsligt regi ser vi AI blive brugt til overvågning og borgeranalyse i en skala, der for få år siden lød som science fiction. Et opsigtsvækkende eksempel er Kinas landsdækkende overvågningssystem kendt som "Skynet", hvor over 200 millioner CCTV-kameraer er koblet sammen med AI, hvilket giver myndighederne en hidtil uset mulighed for at genkende og følge borgere i realtid.

Men Kina er ikke alene. En undersøgelse viser, at næsten halvdelen af alle voksne amerikanere allerede har deres ansigtsbillede liggende i politiets ansigtsgenkendelsesdatabaser ofte via kørekortregistre eller overvågningskameraer.

I Europa eksperimenterer flere storbyer også med live-ansigtsgenkendelse, f.eks. har London testet systemer i det offentlige rum som led i kriminalitetsbekæmpelse.

Disse AI-drevne overvågningsinitiativer rejser et centralt

spørgsmål: Hvor går grænsen mellem offentlig sikkerhed og individets ret til privatliv?

På den anden side står de kommercielle aktører: virksomheder, hvis forretningsmodel i stigende grad handler om at kende dig bedre, end du kender dig selv.

Teknologigiganter som Meta, Google og Amazon indsamler enorme mængder brugerdata, som fodres til AI-algoritmer med ét formål: at tiltrække og fastholde din opmærksomhed, forudsige dine behov og påvirke dine valg.

AI er motoren bag alt fra produktanbefalinger ("Du kunne også være interesseret i...") til de annoncer, der tilsyneladende ved, at du lige manglede et par nye sko.

I denne verden er data guld og vi betaler ofte for "gratis" online tjenester med vores personlige oplysninger.

Dataindsamling i hverdagen: Fra smartphones til smart homes

Det lyder måske abstrakt, men dataindsamlingen er meget konkret i vores hverdag. Smartphonen er et af de mest gennemgribende eksempler. Den er pakket med sensorer: GPS, kamera, mikrofon, accelerometer, gyroskop, pulsmåler - en digital schweizerkniv af datakilder. Hver gang du bevæger dig, registrerer telefonen din lokation. Når du tager et billede, gemmes oplysninger om tid og sted.

Selv din måde at gå på kan aflæses af telefonens bevægelsessensorer.

Disse data gør livet nemmere. Din telefon kan f.eks. automatisk sortere billeder eller forstå dine stemmekommandoer bedre. Men prisen er, at mange af disse data sendes til tech-giganternes servere, hvor AI-modeller lærer af dine vaner.

Wearables som smartwatches og fitness-trackere er en forlængelse af dette. De overvåger din puls, dit søvnmonster, din aktivitet! Alt sammen for at give dig indsigt i din sundhed. Det er nyttigt for brugeren selv, men de samme data kunne i teorien være interessante for f.eks. forsikringsselskaber eller arbejdsgivere.

I dag er det som regel kun dig og din app, der ser tallene, men dataene findes og dermed også risikoen for læk eller misbrug.

"Smart home-enheder" bringer overvågningen helt ind i boligen. Tænk på Amazon Echo (Alexa) eller Google Home. De lytter konstant efter aktiveringsordet, men mikrofonen er i princippet altid tændt.

Flere tilfælde har vist, at klip af samtaler utilsigtet er blevet optaget og sendt til cloudservere. Alexa og Google Home topper faktisk listen over de mest datahungrende smarthome-produkter: en undersøgelse i 2024 viste, at Alexa-appen indsamler 28 forskellige typer data,

mens Google Home-appen ligger lige bagefter med 22. Det inkluderer alt fra lokalitet og enheds-ID til stemmehistorik, kontaktoplysninger og søgehistorik.

Selv din kaffemaskine er ikke længere bare en kaffemaskine. Dens tilhørende app kan registrere, hvornår og hvor ofte du brygger kaffe. Kort sagt: Dit hjem bliver smartere, men det sker ved, at det konstant sender små bidder af dit privatliv ud i verden.

Og så er der internettet og de sociale medier - de steder, hvor dataindsamlingen måske mærkes allermost. Hver gang du surfer på nettet, efterlader du digitale spor: cookies, IP-adresser, browseroplysninger.

Reklamenetværk bruger AI til at analysere disse og opbygge en detaljeret profil af dig: "Personen interesserer sig for sport, rejser, søger bolig og bruger en Android-telefon, måske er det tid til at vise dem møbelreklamer?"

På sociale medier er det endnu tydeligere. Facebook og Instagram registrerer ikke bare dine klik, men også hvor længe du kigger på et opslag. YouTube måler, hvilke videoer du ser til ende, og hvilke du skipper. AI bruger denne information til at finjustere dine anbefalinger, så du hele tiden fodres med indhold, der matcher din smag.

Det er effektivt, men det kan også

skabe ekkokamre, hvor du kun præsenteres for én type indhold, mens platformene holder fast i din opmærksomhed. Tilpasset indhold er en gave, men også en potentiel fælde.

Etik og individets kontrol: Hvem ejer dine data?

Midt i denne hvirvelvind af dataindsamling rejser de etiske spørgsmål sig: Hvor går grænsen for, hvad det er acceptabelt at indsamle om folk? Hvem ejer egentlig alle disse data? Er det dig, der afgiver dem? Eller dem, der gemmer dem på deres servere? Og hvor meget kontrol har vi som individer over alt det, der bliver indsamlet om os?

Mange af os klikker "JA" til lange vilkår og betingelser uden at læse det med småt. Det er forståeligt - ingen kan overskue at læse side op og side ned af juridisk tekst for hver eneste app eller tjeneste. Men konsekvensen er, at vi ofte afgiver retten til enorme mængder oplysninger om os selv. Data om vores færden, vaner og præferencer bliver virksomhedernes ejendom (eller i det mindste stillet til rådighed for dem), og vi stoler på, at de forvalter dem ansvarligt. Det er en tillid, der ikke altid har vist sig at være fortjent.

Her kommer begrebet "digitalt selvforsvar" ind i billedet. En tilgang, der handler om de digitale værktøjer og vaner, vi kan tage i brug for at beskytte os selv mod den stigende overvågning fra både virksomheder og staten, og derved genvinde en vis

kontrol over vores privatliv. Digitalt selvforsvar handler om alt fra at bruge stærke adgangskoder og to-faktor-login, til at blokere tracking-cookies, begrænse app-tilladelser og benytte krypterede beskeds-apps.

Tænk på det som at trække gardinerne for i det digitale vindue: Du behøver ikke finde dig i, at "nogen" altid kigger ind. Derudover handler det om etik i designet af tjenesterne. Bør der ikke gælde principper om "Privacy by Design"? Altså at teknologier fra starten udformes med respekt for brugerens privatliv?

En inspirerende tendens er, at nogle virksomheder, presset af både forbrugere og lovgivning begynder at give brugerne mere kontrol. Apple har f.eks. gjort privatliv til et salgsargument: Flere AI-opgaver i nyere iPhones behandles direkte på enheden (f.eks. ansigtsgenkendelse og tastaturforslag), og apps skal nu spørge om lov, før de må tracke brugeren på tværs af andre tjenester.

Det giver magten lidt tilbage til individet og viser, at det faktisk er muligt at udvikle smarte løsninger uden at vide alt om os.

Alligevel er realiteten, at magtbalancen mellem individ og dataindsamler er skæv. Det er "David mod Goliat": Den enkelte bruger har begrænset overblik og få ressourcer til at beskytte sine, informationer, mens techgiganterne

og staterne råder over avancerede AI-værktøjer og uanede mængder datakapacitet. Derfor er samtalen om etik og regulering afgørende.

Vi ser allerede en bevægelse i den retning: EU's persondataforordning (GDPR) og EU's AI Act er en ufravigelig forordning, som sætter rammen i Europa, for mere ansvarlig brug af data og algoritmer.

Men lovgivning er kun den ene side af sagen. Den anden er kulturel og individuel: Vi må som borgere og brugere blive bevidste om, hvor værdifulde vores data er, og turde stille krav til dem, der vil have adgang til dem.

Tech-giganterne og deres datahøst: Konkrete cases

For at forstå omfanget af AI-drevet dataindsamling kan vi se nærmere på nogle af de største aktører og deres praksis:

Meta (Facebook og Instagram)

Meta lever af vores sociale liv online. Hver statusopdatering, hvert like og hvert klik på Facebook og Instagram fodrer deres AI-algoritmer. Disse algoritmer sorterer nyhedsfeedet, så brugeren præsenteres for indhold, der fastholder opmærksomheden længst muligt.

Et berygtet eksempel er Cambridge Analytica-skandalen i 2018, hvor et firma fik adgang til data fra op mod 87 millioner Facebook-brugere gennem en tilsyneladende uskyldig personlighedstest. Disse data blev

brugt til at skabe psykologiske profiler og målrette politisk reklame.

Skandalen afslørede, hvor let vores interaktioner kan bruges til at manipulere os, uden at vi selv opdager det. Siden da har Meta fået rekordbøder i EU og er under konstant granskning. Men forretningsmodellen er grundlæggende uændret: Jo flere data om dig, desto mere målrettet (og værdifuld) bliver annonceringen.

Google

Chancen er stor for, at du har brugt Google i dag, f.eks. til at finde information, tjekke vejret eller søge vej. Google ved enormt meget om os: hvad vi søger, hvilke videoer vi ser på YouTube, hvor vi bevæger os med Maps, og hvad vi skriver i vores Gmail og kalender.

AI bruges til at forbedre brugeroplevelsen; fra smartere søgeresultater til trafikfri ruter, men dataene har også en kommerciel værdi.

I en opsigtsvækkende sag blev det afsløret, at Google i årevis fortsatte med at spore brugeres lokation, selv efter at de havde slået "lokationshistorik" fra. Det førte i 2022 til et forlig på USD 391,5 millioner til 40 amerikanske delstater.

Google arbejder også på teknologier som "Federated Learning", der træner AI-modeller lokalt på brugerens enhed og dermed øger

privatlivet. Men grundlæggende opsamler Googles økosystem data fra næsten alle aspekter af vores digitale liv.

Amazon

Amazon er ikke bare en webshop - det er også en af verdens største cloud- og AI-udbydere.

Hver gang du handler på Amazon, registreres og analyseres din adfærd. AI bruger oplysninger om, hvad du kiggede på, hvad du lagde i kurven og ikke købte, og hvad du tidligere har bestilt. Alt sammen for at forudsige og påvirke dit næste køb.

Amazon ejer også Alexa og Ring - intelligente højttalere og dørklokker med kamera. I 2022 kom det frem, at Amazon uden brugernes viden havde udleveret Ring-videooptagelser til politiet mindst 11 gange på et år.

Selvom det skete under henvisning til nødsituationer, rejser det spørgsmål om, hvorvidt vi, uden at vide det, har inviteret overvågning helt ind i hjemmet.

Amazon har også udviklet ansigtsgenkendelsessoftwaren Rekognition og solgt den til politimyndigheder, indtil salget blev sat på pause efter afsløringer af fejlagtig genkendelse, bl.a. matchede AI'en 28 amerikanske kongresmedlemmer med forbryderfotos. Eksempler som disse viser, hvor magtfuld og potentielt fejlbart

AI-drevet overvågning kan være.

OpenAI (ChatGPT)

OpenAI adskiller sig fra de øvrige ved ikke at indsamle data via en platform eller butik, men via selve AI-værktøjet.

Når du bruger ChatGPT, bliver dine samtaler som udgangspunkt gemt og kan bruges til at forbedre fremtidige modeller.

Det skabte debat i 2023, da ChatGPT blev midlertidigt blokeret i Italien. De italienske myndigheder mente, at OpenAI overtrådte GDPR ved at indsamle personoplysninger uden lovligt grundlag og uden tilstrækkelig information til brugeren.

OpenAI blev pålagt at forbedre transparensen og give mulighed for at slå chat-historik fra, hvilket blev implementeret, sammen med aldersverifikation og flere indstillinger for datasikkerhed.

Det er en vigtig påmindelse: Selv når vi “bare” taler med en chatbot, foregår der dataindsamling i baggrunden. Det rejser nye spørgsmål: Hvad hvis man uforvarende deler fortrolig information og den ender i træningsdata?

Flere virksomheder har allerede forbudt brugen af ChatGPT til følsomt arbejde. Casen med OpenAI viser, at også de nye aktører skal

holdes ansvarlige og at vi stadig er ved at finde rammerne for en ansvarlig AI-fremtid.

Skyggesiderne: bias, manipulation, datalæk og profilering

Når AI og dataindsamling smelter sammen, skaber det ikke kun muligheder, men også betydelige risici og potentielle ubalancer i samfundet. Lad os se nærmere på nogle af skyggesiderne:

Bias (forudindtagethed) i algoritmerne

AI-modeller er kun så gode som de data, de er trænet på. Hvis der er skævheder i datagrundlaget, f.eks. at et ansigtsgenkendelsessystem er trænet primært på billeder af lyshudede mænd og kun få af kvinder eller personer med mørkere hud - ja, så bliver systemet markant dårligere til at genkende de underrepræsenterede grupper. Det er ikke bare teori: Der har været flere eksempler på ansigtsgenkendelse, som identificerer hvide ansigter præcist, men fejler fatalt på sorte ansigter. Det har i USA ført til falske anholdelser af uskyldige mennesker.

Bias kan også snige sig ind i anbefalingssystemer: Hvis en AI lærer, at “en god medarbejder” ser ud på en bestemt måde, baseret på historiske data fra en virksomhed med mange mandlige ansatte, kan den begynde at frasortere kvindelige ansøgere som “ikke egnede”. Det var præcis, hvad der skete med et

internt rekrutteringsværktøj hos Amazon, som måtte skrottes.

Problemet med bias er, at det ofte camoufleres som neutral teknologi: "computeren siger det jo". Men hvis computeren er forudindtaget, kan den automatisere diskrimination i massiv skala.

Manipulation og ekkokamre

Som tidligere nævnt kan AI bruges til at målrette budskaber ekstremt præcist. I værste fald kan det udnyttes til manipulation.

Cambridge Analytica-sagen viste, hvordan politiske konsulenter kunne bruge folks egne data imod dem ved at skræddersy propaganda, der talte ind i netop deres frygt og håb.

Generelt har sociale mediers algoritmer en tendens til at fremhæve indhold, der vækker stærke følelser, fordi det holder os engagerede. Det betyder, at vrede, chokerende eller polariserende opslag ofte får mest opmærksomhed og dermed skubber folk længere ud i ekstremerne.

Når vi alle får hver vores kuraterede nyhedsstrøm, risikerer vi at leve i parallelle informations-verdener, hvor vores synspunkter forstærkes frem for udfordres. Det er en form for manipulation af opmærksomhed og holdninger, som ikke nødvendigvis er ondsindet, men som er en utilsigtet konsekvens af algoritmernes jagt på engagement.

Datalæk og sikkerhed

Jo mere data der indsamles og lagres, desto større er risikoen for, at det bliver stjålet - og desto større er skaden, når det sker.

Vi har set adskillige eksempler: virksomheder der mister millioner af brugeres data, sundhedsapps der ved en fejl lækker følsomme helbredsoplysninger, eller private kameraoptagelser der pludselig dukker op online.

Når data samles ét sted, bliver ét sikkerhedshul til et digitalt olieudslip.

Selv store selskaber kæmper med at sikre deres systemer og når de samtidig deler data med tredjeparts-udbydere, som måske ikke har samme sikkerhedsniveau, øges risikoen.

Datalæk underminerer tilliden til det digitale samfund. Hvis vi ikke tør dele selv nødvendige oplysninger, fordi vi frygter misbrug, kan det i sidste ende hæmme de mange positive anvendelser af data.

Algoritmisk profilering og magtubalance

I sidste ende handler meget af dette om magt. Den, der har adgang til data - og AI'en til at analysere dem, får stor indflydelse på, hvordan mennesker bliver vurderet.

Vi ser allerede kimen til sociale "score"-systemer: mest ekstremt i

Kina, hvor borgeres adfærdspoint kan afgøre deres adgang til lån, rejser eller skoler.

Men også i Vesten foregår profilering: Forsikringsselskaber bruger data til at justere præmier, arbejdsgivere vurderer ansøgers personlighed på baggrund af digitale fodaftryk, og banker analyserer forbrugsmønstre i kreditvurderinger.

Algoritmisk profilering kan skabe uigennemskuelige ulemper for dem, der har "forkerte" data om sig, uden at de nogensinde ved hvorfor. Det skaber en følelse af afmagt: at vi bliver vurderet i det skjulte, uden mulighed for indsigt eller modargument.

At genvinde kontrollen - en personlig afslutning

Midt i alt dette kan man nemt føle sig overvældet. Teknologien buldrer frem, og privatlivet føles som noget, der smuldrer mellem fingrene på os. Men der er også grund til håb og til handling.

Først og fremmest taler vi om det nu. Det, at du læser dette kapitel, er et skridt. Bevidsthed er forudsætningen for forandring. Når vi først forstår, hvor meget data der flyder rundt, kan vi begynde at stille krav: til virksomheder ("hvad bruger I mine data til?"), til politikere ("vi vil have love, der beskytter os bedre") og til os selv ("burde jeg tjekke mine app-tilladelser og privacy-indstillinger?").

Teknologien i sig selv er ikke ond! Det handler om, hvordan vi bruger den. AI kan skabe enorme fremskridt: bedre sundhedsvæsen, grønne byer, nemmere hverdage. Men det må ikke ske på bekostning af vores grundlæggende rettigheder.

Privatliv er en form for frihed - friheden til at være sig selv uden konstant overvågning.

I denne bog handler om, er kampen om privatlivets grænser en af de vigtigste slagmarker lige nu.

Vi står i et spændingsfelt mellem innovation og etik, mellem data-dreven bekvemmelighed og risikoen for et overvågningssamfund. Det inspirerende er, at fremtiden ikke er skrevet endnu. Vi kan være med til at forme den med de valg vi træffer: hvilke tjenester vi bruger, hvilke politikere vi stemmer på, og hvilke samtaler vi tager med vores børn.

Med digitalt selvforsvar i værktøjskassen og en insisteren på gennemsigtighed og fairness kan hver enkelt af os være med til at sikre, at AI forbliver vores tjener, ikke vores overvåger. Privatliv behøver ikke være en saga blot i den digitale tidsalder, men det kræver, at vi er vågne, nysgerrige og villige til at kæmpe for det. For den viden, du har nu, er allerede dit stærkeste våben i kampen for dit privatliv.

Fakta-tjek

Tre spørgsmål til fremtiden

Fremtiden formes ikke kun af teknologi, men af de valg vi træffer i nutiden.

Stop op et øjeblik og spørg dig selv:

1. Hvem ønsker du, der former fremtidens AI og hvordan?

Store virksomheder? Demokratisk valgte ledere? Dig selv?

 "Hvem har magten – og hvem skal have den?"

2. Hvad er vigtigst for dig: effektivitet, etik eller fællesskab?

AI kan spare tid, løse problemer og skabe vækst, men til hvilken pris?

 "Hvad skal teknologien tjene og hvem skal den gavne?"

3. Hvad vil du spørge en AI om om 10 år?

Vil du bede den om at planlægge dit liv? Løse klimakrisen? Trøste dig?

 "Hvordan forestiller du dig dit liv med AI i 2035?"

AI SOM VÅBEN: "CYBERSIKKERHEDENS NYE FRONTLINJE"

En moderne cyber-soldat: AI bruges i dag både til angreb og forsvar i cyberspace. Den moderne eller fremtidige kampplads er uden røg og ekko fra geværskud, men i stedet et digitalt landskab, hvor linjer af kode er skyttegravene, og kunstig intelligens (AI) er det nye våben. Vi står ved en ny frontlinje i cybersikkerheden, hvor AI kan være både fjenden og en allieret. Denne frontlinje strækker sig fra din e-mailindbakke til staters højeste sikkerhedsorganer.

I dette kapitel udforsker vi, hvordan AI på den ene side muliggør mere sofistikerede cyberangreb, fra snedige phishing-mails til statsstøttet cyberspionage og på den anden side åbner for nye forsvarsmekanismer, der kan opdage og afværge trusler i realtid.

Vi kommer omkring eksempler fra hele verden: fra Nordkoreas hackerbrigader til Silicon Valleys softwaregiganter.

Vi ser også på, hvordan AI påvirker beskyttelsen af kritisk infrastruktur og hvordan det er blevet et nyt konkurrenceparameter for både virksomheder og stater. Endelig stiller vi de etiske spørgsmål: Når AI

bliver et våben, hvem bærer så ansvaret? Og kan verdenssamfundet finde fælles spilleregler?

AI på angrebssiden: Fra phishing til Cyberkrig

Angribere har altid forsøgt at være "et skridt" foran forsvaret, og med AI i arsenalet har de fået et kraftfuldt nyt våben. Phishing-angreb, som er de bedrageriske e-mails eller beskeder, der narrer folk til at afsløre følsomme oplysninger, har traditionelt været til at gennemskue, hvis sproget var kluntet eller indholdet for generisk. Men nu kan hackere bruge sprogmodeller til at automatisere skræddersyede phishing-mails, der er uhyggeligt overbevisende. Værktøjer som WormGPT er udviklet med netop dette formål: som en slags "ondsindet ChatGPT" uden etiske begrænsninger, der kan spytte velformulerede, personligt tilpassede snyde-mails ud på samlebånd.

Tænk på en e-mail, der ligner den kommer fra din chef - med korrekt sprog, passende tone og detaljer om dit seneste projekt, men som i virkeligheden er fabrikeret af en AI. Det øger risikoen for, at selv rutinerede brugere bliver narret (inklusiv mig selv).

Et andet voksende fænomen er "Deep Fake"-angreb. Deepfakes er AI-genererede video- eller lydklip, der kan få det til at se ud, som om en person siger eller gør noget, de aldrig har gjort.

I 2022 dukkede en video op, hvor Ukraines præsident Zelenskyj angiveligt opfordrede sit folk til at overgive sig; noget han naturligvis aldrig sagde.

Videoen var klodset og blev hurtigt afsløret, men den viste potentialet: Deep Fakes kan udløse diplomatiske kriser, skabe uro og underminere tillid.

I 2024 gik en multinational virksomhed i en kostbar fælde, da en Deep Fake-video af deres finansdirektør blev brugt til at narre medarbejdere til at foretage en stor bankoverførsel, med et milliontab til følge.

Social Engineering 2.0: falske profiler og stemmer

AI forstærker også Social Engineering-angreb. Nordkoreanske hackergrupper bruger AI-værktøjer til at skabe falske profiler med overbevisende billeder og stemmer.

I 2024 afslørede Microsoft, at gruppen Sapphire Sleet stjal kryptovaluta for over 10 millioner dollars ved at udgive sig for rekrutteringsfolk på LinkedIn. De brugte værktøjer som FaceSwap til at generere realistiske profilbilleder og dokumenter og i nogle tilfælde AI-baseret stemmeforvrængning for at simulere telefonopkald.

Når det bliver muligt at efterligne både udseende og stemme, bliver det sværere end nogensinde at skelne ven fra fjende online.

Malware, ransomware og selvforandrende virusser

AI spiller også en rolle i mere tekniske angreb. Ransomware-grupper, der krypterer data og kræver løsepenge bruger AI til at identificere de mest værdifulde filer og tilpasse angrebene, så de undgår detektion.

Der eksperimenteres med polymorfe virusser, der konstant ændrer form for at snyde antivirusprogrammer. Samtidig benytter hackerne "Machine Learning" til at scanne tusindvis af netværk for sårbarheder langt hurtigere end før. Forestil dig et AI-system, der kontinuerligt leder efter kendte og ukendte "Zero-Days" sårbarheder.

I 2024 annoncerede Google, at deres AI-agent "Big Sleep" havde fundet en alvorlig zero-day-fejl i den udbredte database SQLite.

Hvis de gode kan det, kan de onde også. Forskellen er blot, at kriminelle ikke nødvendigvis afslører, hvad de har fundet, men udnytter det i stilhed.

Automatiseret hacking: AI som angrebsværktøj

AI kan bruges som "hackerhjælp" til at automatisere hele kæden i et cyberangreb.

En FBI-rapport fra 2024 advarer om, at især kinesiske statsstøttede aktører tester AI i alle faser af et angreb - fra måludvælgelse og

scanning til udformning af phishing-materiale og opsætning af falske virksomheder. Det gør angrebene mere effektive og langt mere skalerbare.

AI gør det muligt at oprette hundreder af digitale identiteter og sende tusindvis af troværdige phishing-mails på kort tid. Det erstatter ikke hackeren, men det gør hackerens arbejde hurtigere og mere effektivt.

Ifølge FBI's cybersikkerhedschef Cynthia Kaiser kan en enkelt angriber med AI nu udføre på en dag, hvad der tidligere tog en uge.

Et nyt våben i cyberkrigens arsenaler

Denne udvikling har fået mange til at kalde AI for "det næste atomvåben" inden for cyberkrig.

Forskellen er blot, at hvor atomvåben er synlige og ødelæggende, er AI et tyst og snedigt våben: en usynlig hær af algoritmer, der angriber systemer, infiltrerer virksomheder og manipulerer information, ofte uden at nogen opdager det før skaden er sket.

AI som skjold: Intelligente forsvarsmekanismer

Heldigvis er AI ikke kun en gave til de onde, det er også blevet forsvarernes stærkeste nye våben. Cybersecurity-forsvar har traditionelt været et kapløb mod tiden: Kan man opdage og stoppe et angreb, før det gør skade? Her åbner

kunstig intelligens for helt nye muligheder, ikke blot for at reagere i realtid, men også for at forudsige trusler, før de rammer.

En af de største udfordringer i cybersikkerhed er mængden af data.

Store organisationer modtager dagligt millioner af logbeskeder, netværkspakker og advarsler. At finde den ene ondsindede aktivitet i denne enorme datamængde svarer til at finde nålen i høstakken. Her excellerer AI.

AI-drevne trusselsdetekterings-systemer kan lynhurtigt analysere enorme datamængder og genkende mønstre, som mennesker ville overse.

Ved at sammenligne normal adfærd med nuværende aktiviteter kan maskinlæringsmodeller opdage anomalier, f.eks. hvis en bruger pludselig downloader gigabytes af data midt om natten eller hvis tusindvis af mislykkede loginforsøg sker i løbet af få sekunder. AI hjælper med at filtrere støj fra, reducere falske alarmer og fremhæve det, der virkelig kræver opmærksomhed.

Kort sagt: AI fungerer som et "cybervagtværn", der aldrig sover, aldrig bliver træt og konstant lærer af hver eneste hændelse.

Et godt eksempel er Microsofts egen sikkerhedsplatform. Med milliarder af brugere på tværs af Windows,

Azure, Office 365 mv. indsamler Microsoft over 65 milliarder sikkerhedssignaler om dagen. Denne globale trusselsintelligens bruges til at træne AI-modeller, der kan opdage angrebsteget på tværs af kontinenter. I 2023 lancerede Microsoft deres Security Copilot - en intelligent assistent til it-sikkerhedsfolk, der kombinerer OpenAI's GPT-4 med Microsofts trusselsdata.

Copilot kan f.eks. opsummere komplekse angrebsforløb og foreslå konkrete modforanstaltninger.

For forsvarere, der ofte føler sig overvældede af angrebsmængden, kan sådanne AI-værktøjer tippe balancen: en kamp mellem maskiner: De gode mod de onde.

AI spiller også en voksende rolle i automatiseret pentesting og sårbarhedsscanning. Hvor man før hyrede "etiske hackere" til at teste systemers sikkerhed, ser vi nu AI-systemer, der kontinuerligt forsøger at bryde ind i egne netværk og rapporterer svagheder, før de udnyttes.

Google's Big Sleep-agent, der opdagede en zero-day sårbarhed i SQLite, er et eksempel på dette. Andre udvikler AI-"fuzzere", der tester software med tusindvis af inputs for at fremprovokere fejl og vurderer, om disse fejl er sikkerhedskritiske.

Et andet vigtigt område er threat

intelligence - altså viden om angribernes metoder og værktøjer. Her kan AI trawle gennem åbne kilder, dark web-fora og sensordata fra hele verden for at identificere mønstre. Eksempelvis kan en AI opfange, at der diskuteres en ny sårbarhed på russisksprogede fora og slå alarm, før den udnyttes bredt.

AI kan også identificere globale mønstre, f.eks. hvis hundredvis af små hændelser rundt om i verden viser tegn på en koordineret kampagne. Det menneskelige øje ville næppe koble disse prikker sammen i tide, men AI kan.

Disse indsigter deles via samarbejde mellem tech-giganter og myndigheder, som f.eks. fælles databaser over malware og phishing-taktikker. Resultatet er en mere koordineret og proaktiv cybersikkerhed.

Det er fristende at fremstille dette som en kamp mellem god og ond AI. Men bag begge sider står stadig mennesker med deres egne intentioner. Alligevel er én ting sikker: Intet større cyberforsvar vil i fremtiden kunne klare sig uden hjælp fra kunstig intelligens.

Kritisk infrastruktur: Samfundets livsnerve på spil

I denne hyperforbundne æra er kritisk infrastruktur: vores elnet, hospitaler, vandforsyning, transport og kommunikation o.a. er blevet et primært mål for cyberangreb. Konsekvenserne af et vellykket angreb kan være katastrofale:

forestil dig et regionalt strømsvigt midt om vinteren, tog der brat går i stå, eller hospitalssystemer der bryder sammen, så livsvigtige operationer må aflyses.

Desværre er dette ikke blot hypotetiske skræmmebilleder. Vi har allerede set eksempler: WannaCry-ransomware lammede britiske hospitaler i 2017, og et hackerangreb i 2021 lukkede en stor brændstof-pipeline i USA, hvilket udløste panikopkøb og brændstofmangel.

Hvordan ændrer AI ligningen?

På angrebssiden gør AI det lettere at ramme præcist. Hvor hackere tidligere måtte famle i blinde for at forstå et energiselskabs netværk, kan AI nu analysere offentligt tilgængelige data og kortlægge strukturer: Hvem leverer softwaren? Hvilke enheder er forbundet? Hvor er de svageste led?

Industrielle kontrolsystemer (ICS), som styrer alt fra elkraft til vandpumper, er ofte bygget til driftssikkerhed - ikke cybersikkerhed. AI-værktøjer kan skanne disse systemer for uregelmæssigheder, som en angriber så kan udnytte til at trænge ind.

Der er rapporter om russiske hackere, der har udviklet AI-systemer, der kan analysere elnetdata og kamuflere ondsindet aktivitet som almindelig drift, for at undgå alarmklokker.

Når først angriberen er indenfor, kan

AI også maksimere skaden: I et energinet kan AI f.eks. koordinere samtidige overbelastninger af transformerstationer og dermed skabe omfattende blackout, som ville være vanskeligt at orkestrere manuelt.

AI på forsvarssiden: Når maskiner beskytter maskiner

Heldigvis styrker AI også vores forsvar. AI-baseret overvågning kan lære et systems normale rytme, hvor meget strøm der bruges hvornår, hvornår et vandrensningsanlæg kører på fuldt tryk og udløse alarm ved selv små afvigelser. Disse systemer reagerer hurtigere end mennesker og kan opdage og afværge trusler i splitsekunder.

Eksempelvis kan AI i et elnet identificere kendte angrebsmønstre, f.eks. forsøg på at åbne afbrydere og isolere den berørte del af nettet, før skaden spredes sig. I transport-sektoren kan AI overvåge signal-anlæg og afsløre forsøg på at manipulere jernbanetrafik.

Og i sundhedsvæsenet anvendes AI allerede til at overvåge og beskytte netværk og udstyr som medicin-pumper og scannere.

Et særligt skræmmende scenarie er, hvis hackere får kontrol over hospitalsudstyr. Her kan AI-baserede overvågningssystemer lære, hvordan enheder normalt opfører sig og reagere, hvis f.eks. en infusionspumpe pludselig ændrer dosis eller adfærd.

AI som skadeskontrol og beredskab
Kritiske systemer skal både være sikre og konstant tilgængelige. Man kan ikke bare slukke hospitalets IT for at installere opdateringer, mens patienter ligger på operationsbordet. Dette gør mange systemer sårbare og i nogle tilfælde forældede.

Her kan AI hjælpe som intelligent “skadeskontrol”. Ved et ransomware-angreb kunne AI hurtigt identificere hvilke maskiner, der er kompromitteret, og isolere dem, før infektionen spreder sig til livsvigtigt udstyr.

AI giver også nye muligheder for beredskab. Ved at simulere angreb i digitale tvillinger - virtuelle modeller af f.eks. elnettet, kan AI hjælpe med at forudsige angrebs-effekter og styrke beredskabet.

Forestil dig en AI-styret cyberøvelse, hvor en “forsvars-AI” kæmper mod en “angrebs-AI” i et simuleret miljø. Det lyder futuristisk, men det er allerede under udvikling. I USA har man testet sådanne simulationer på vandværker med henblik på at lære og forbedre reelle forsvars-mekanismer.

Den ultimative kampplads for AI i cybersikkerhed

Kritisk infrastruktur er i dag den mest afgørende frontlinje i cybersikkerheden og AI er aktiv på begge sider. Det handler ikke længere kun om data og penge, men om liv og

død, varme og lys, ro og kaos. Truslen er reel: selv FN har advaret om, at integrationen af AI i avancerede våbensystemer, herunder cybervåben kan få uoverskuelige konsekvenser uden international regulering.

Men samtidig giver AI os også en reel chance for at beskytte det, vi er allermost afhængige af. Den teknologiske udvikling har gjort vores samfund sårbare, men den kan også gøre os stærkere.

AI og cybersikkerhed som konkurrenceparameter

Det siges nogle gange, at “der findes to slags virksomheder: dem, der har været hacket, og dem, der ikke ved, at de har været det”. I en tid, hvor databrud kan skade både omdømme og bundlinje, er cybersikkerhed rykket fra det tekniske maskinrum til direktionslokalet. Og med AI på scenen er konkurrencen kun blevet skærpet.

For private virksomheder kan evnen til at beskytte kundedata og sikre driftsstabilitet ved hjælp af AI være et afgørende salgsargument. Forestil dig to finansielle aktører: Den ene lover sine kunder realtids-overvågning af transaktioner og svindel via AI, mens den anden fortsat bruger traditionelle metoder. Mange - især erhvervskunder, vil føle sig mere trygge ved den første.

Vi ser allerede store investeringer i AI-sikkerhed: banker bruger AI til

at spotte usædvanlige pengeoverførsler, og e-handelsgiganter forhindrer falske login-forsøg, inden de sker. Virksomheder konkurrerer nu på at have de klogeste forsvarssystemer.

I visse brancher, som forsikring, er cybersikkerhed (inkl. AI-brug) blevet en parameter i risikovurdering: virksomheder med AI-forstærket sikkerhed kan forhandle sig til lavere cyberforsikringspræmier, fordi risikoen vurderes som lavere. Også i tech-branchen er AI-drevet sikkerhed et konkurrenceparameter. Giganter som Microsoft, Google og IBM kappes om at levere de mest avancerede AI-sikkerhedsløsninger, og dette ræs driver innovation.

Vi ser smartere antivirus, forbedret netværksovervågning og mere effektive reaktionsværktøjer. Sikkerhedsslogans som "Vend spillet til forsvarernes fordel" dominerer tech-konferencer, Sikkerhed er blevet et brand i sig selv.

Kampen mellem stater: Sikkerhed som geopolitisk kapital

I den offentlige sektor spiller en tilsvarende konkurrence sig ud - nu mellem nationer.

Cybersikkerhed er blevet et geopolitisk statusparameter på linje med økonomisk og militær magt. Et land, der kan sige: "Vores elnet er blandt de mest cyber-sikre i verden takket være AI", sender et klart signal om robusthed. Det tiltrækker

investeringer og tillid. Omvendt mister lande, der gang på gang rammes af cyberangreb, hurtigt konkurrenceevne. Tænk bare på banksektorer med hyppige databrud.

For regeringer er det derfor en prioritet at støtte talentudvikling inden for AI og cybersikkerhed. Vi ser cyberakademier og AI-innovationscentre skyde op med statsstøtte. Indien og Brasilien investerer målrettet i AI-sikkerhed, EU koordinerer AI-forskningsprojekter, og i USA har præsident Biden slået fast, at lederskab i ansvarlig AI er afgørende for national sikkerhed.

Vi befinder os midt i et "digitalt våbenkapløb", hvor heldigvis ikke kun angreb, men især forsvarsevne er blevet en central konkurrencefaktor.

Microsofts model: AI-drevet sikkerhed i lag

Et markant eksempel på denne udvikling er Microsofts tilgang til cybersikkerhed. Virksomheden har gjort AI-baseret sikkerhed til en strategisk kerneydelse og et konkurrenceparameter i sig selv.

Centralt står princippet om "Zero Trust": Ingen, hverken personer eller enheder får automatisk adgang, uanset om de befinder sig inden for virksomhedens netværk. Alt skal verificeres løbende. AI vurderer risikoprofilen: Logger du på fra en ny enhed i et andet land? Så kan

systemet kræve multifaktor-godkendelse eller helt blokere adgangen.

Microsofts cloud-plattform (Azure) benytter en multi-tenant arkitektur, hvor tusindvis af kunder deler fysisk infrastruktur. AI og avancerede isolationsmekanismer sikrer, at data aldrig blandes mellem kunder.

Hardwarebaseret nøglebeskyttelse og roterende sikkerheds-tokens bruges til at beskytte adgangen. Det er som at have individuelle pengeskabe i samme bankboks-rum.

For kunderne tilbyder Microsoft forskellige sikkerhedsniveauer: fra "grundsikring" med stærke adgangskoder og e-mailfiltrering, til "ekstra høj sikkerhed" med avanceret adgangskontrol, device-styring via Intune og stram datadeling.

AI som ryggraden i Microsofts cybersikkerhed

AI er integreret i stort set hele Microsofts sikkerhedssystem. I Microsoft Defender (antivirus og endpoint-sikkerhed) bruges AI til at identificere og standse trusler, selv hvis de aldrig er set før ved at genkende adfærdsmønstre, som f.eks. et program der pludselig begynder at kryptere filer. "Security Copilot" - Microsofts AI-assistent for sikkerhedsprofessionelle kombinerer GPT-4 med realtidsdata fra trusselsbilledet og hjælper analytikere med at forstå

og respondere på angreb hurtigere.

Trusselsplatformen Defender Threat Intelligence giver adgang til analyser af hackerinfrastruktur og phishing-domæner, genereret af AI, der bearbejder 65 milliarder sikkerhedssignaler om dagen.

For brugerne sker alt dette bag kulissen: Outlook advarer om skadelige vedhæftninger, OneDrive reagerer ved tegn på ransomware, og loginforsøg blokeres ved mistanke om tyveri. Summen af tusindvis af AI-beslutninger gør sikkerheden usynligt stærk.

Microsofts filosofi er klar: Sikkerhed skal være indbygget - ikke tilføjet bagefter.

Med AI har de fået en motor, der løfter den filosofi til et nyt niveau. Etik, ansvar og behovet for globale spilleregler

Når teknologien stormer frem, som AI gør det i disse år, halter lovgivning og etik ofte bagefter. Vi befinder os i et dilemmafyldt landskab: På den ene side ønsker vi at udnytte AI's fulde potentiale til at beskytte os mod cybertrusler. På den anden åbner den samme teknologi for misbrug med ødelæggende konsekvenser. Så hvem bærer ansvaret, når AI bliver et våben og hvordan undgår vi, at det løber løbsk? Et centralt etisk dilemma er spørgsmålet om autonomi. Hvis vi lader AI håndtere cyberforsvar

automatisk, kan det indebære, at en algoritme får lov til at træffe beslutninger om modangreb. Hvad nu, hvis en forsvars-AI fejlagtigt identificerer en legitim system-handling som et angreb og "slår igen" mod en uskyldig tredjepart?

I militær sammenhæng taler man om faren ved autonome våben. Det samme gælder her. "AI uden menneskelig kontrol vil efterlade verden i blinde," advarede FN's generalsekretær.

Bør en AI nogensinde have lov til at udføre destruktive handlinger uden menneskelig godkendelse? Mange vil mene nej, især når det drejer sig om kritisk infrastruktur eller gengældelsesangreb, der kan eskalere konflikter.

Ansvarsspørgsmålet er ligeledes komplekst. Hvis en AI-model begår en fejl, f.eks. overser et angreb eller forårsager skade, hvem holdes ansvarlig? Er det programmørerne, virksomheden, brugeren, eller staten? Det er ikke blot teoretisk. Forestil dig et hospitals-AI, der fejlkategoriserer en vigtig proces som skadelig og lukker den ned, hvilket resulterer i, at en patient ikke modtager livsnødvendig behandling. Der sker skade, men hvem bærer skylden? Denne gråzone gør både jurister og forsikringsselskaber nervøse. Det kan kræve helt nye ansvarsmodeller i lovgivningen.

Et andet etisk problem er asymmetrien i AI-våben. Små grupper eller

svage stater kan skabe stor skade med relativt få ressourcer, som f.eks. Nordkoreas hackere, der har stjålet milliarder i kryptovaluta.

Det giver "magt til de små", men uden den ansvarlighed, vi normalt kræver af stormagter. Der findes endnu ingen Geneve-konvention for cyberkrig - ingen klare definitioner af "krigsforbrydelser" i cyberspace.

International regulering er endnu i sin vorden, men der er voksende enighed om behovet for globale spilleregler. FN's Sikkerhedsråd har diskuteret AI i konflikter og efterlyst en fælles ramme, der kan forhindre fragmenteret governance.

I 2024 opfordrede flere lande til indførelsen af globale "guardrails" for AI - basale regler, som alle stater bør følge. Det kunne være et forbud mod helt autonome angreb uden menneskelig indblanding eller beskyttelse af hospitaler og skoler mod AI-baserede cyberangreb, ligesom vi har beskyttede zoner i fysisk krig.

Samtidig forsøger regioner som EU at gå foran med regulering, bl.a. gennem forordningen EU's AI Act, der klassificerer AI-systemer efter risikoniveau og forbyder visse anvendelser, såsom social scoring. Men hvordan håndhæver man regler på tværs af grænser, når en virus udviklet ét sted kan sprede sig globalt på få sekunder? Derfor argumenterer mange eksperter for et internationalt samarbejde - måske

en slags “cybervåben-traktat”, analogt med atomvåbenkontrol. Sideløbende er der et etisk hensyn i selve forsvarsarbejdet. AI-drevne cybersikkerhedssystemer kræver tæt overvågning af netværk og brugeradfærd. Hvor går grænsen for privatlivets fred? Hvis AI skal forudsige angreb ved at analysere al datatrafik og identificere mistænkelig brugeradfærd, skal vi sikre os, at overvågningen ikke glider over i masseovervågning. Her kræves klare retningslinjer for datahåndtering og respekt for individets rettigheder - selv når formålet er sikkerhed.

Det hele kan virke overvældende, men vi kan også vælge at se det som en mulighed: Menneskeheden står ved en skillevej, hvor vi kan forme, hvordan AI skal anvendes ansvarligt og etisk.

Det kræver samarbejde mellem teknologer, etikere, jurister og beslutningstagere. Vi skal uddanne næste generation af “cyberingeniører” til ikke blot at være teknisk dygtige, men også etisk bevidste.

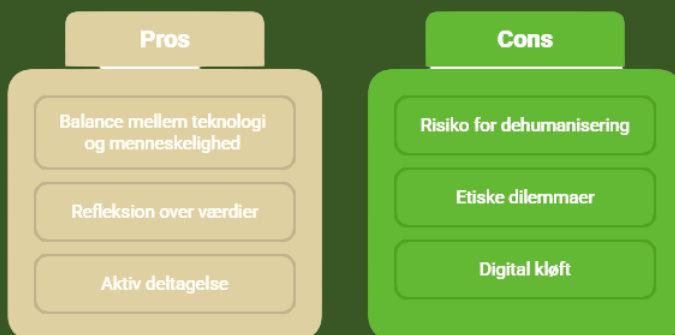
Én ting står klart: AI vil ikke forsvinde fra hverken cybertruslernes eller cyberforsvarets verden. Og i dette kapløb må forsvaret vinde hver gang, men angriberen behøver kun at lykkes én enkelt gang.

Med AI har vi fået et kraftfuldt redskab. Vores opgave er at sikre, at det bruges til at beskytte, ikke til at skade. Kampen om AI i cyber-

sikkerhed er allerede i gang. Nu gælder det om at sikre, at de ansvarlige, de etiske og de beskyttende kræfter vinder.

Det kræver globalt samarbejde, kloge hoveder og måske en smule inspiration fra science fiction. Frontlinjen er trukket. Nu handler det om at stå på den rigtige side.

Fakta-tjek



Forfatterens håb for fremtiden

“Jeg håber, at vi ikke mister det menneskelige i jagten på det teknologiske. At vi ikke blindt lader algoritmer træffe beslutninger, men husker, at mennesker altid må være kompasset.”

“Jeg håber, at vi alle - uanset alder eller baggrund, tør stille spørgsmål, prøve os frem og insistere på, at fremtiden skal være både smart og værdig.”

“Og jeg håber, at du, som har læst med hertil, vil dele dine tanker og være en stemme i samtalen om vores fælles fremtid.”

AI OG GLOBAL RETFÆRDIGHED - DEN DIGITALE KLØFT

Vi befinder os i dag i en situation, hvor AI bruges meget forskelligt og langt fra ens. Eksempelvis findes der personer og ganske få virksomheder, hvor AI udfolder sig som et dagligdags redskab, der booster produktiviteten og gør hverdagen smartere. Hos andre er AI stadig et "system", som ikke er taget i brug af forskellige årsager. Sidst med ikke mindst er AI en fjern drøm, begrænset af manglende internetadgang, ustabil strøm og få digitale kompetencer.

Den globale udbredelse af AI rejser et afgørende spørgsmål: Vil AI forstørre den kløft, der allerede adskiller rige og fattige regioner, eller kan den være med til at bygge bro og fremme global retfærdighed?

Den "digitale kløft", forskellen i adgangen til informationsteknologi, får i AI-alderen en ny dimension. Allerede nu ser vi konturerne af en "AI-kløft", hvor de mest avancerede teknologier og ressourcer er koncentreret få steder i verden. Samtidig spirer håbet om, at AI også kan demokratisere viden, styrke udviklingen i fattigere lande og blive en løftestang for social fremgang. Dette kapitel undersøger, hvordan AI både udfordrer og potentielt kan

forbedre global retfærdighed: fra infrastruktur og færdigheder til råmaterialer, geopolitik og internationale initiativer.

Ulig adgang: Infrastruktur og digitale færdigheder på verdensplan

Verden er i 2025 stadig præget af dyb ulighed, når det gælder basale digitale forudsætninger. Over halvdelen af menneskeheden er nu online, men adgangen er ekstremt ujævnt fordelt: Omkring 67% af verdens befolkning bruger internet, men i Afrika er tallet kun omkring 37%. I lavindkomstlande var det i 2022 blot 25%, mens over 90% af befolkningen i højindkomstlande er online.

Denne digitale kløft danner fundamentet for en voksende AI-kløft. Internetadgang, strømforsyning og digital uddannelse er nødvendige for at udvikle og bruge AI, og her halter store dele af verden bagefter. Mange regioner i Afrika og Sydasien kæmper med ustabil elektricitet, og mere end 500 millioner mennesker i verdens mindst udviklede lande har slet ikke adgang til strøm. Uden elektricitet kan man ikke drive computere, mobilmaster eller oplade smartphones - alt sammen nødvendigt for moderne digital deltagelse.

Selv hvor infrastrukturen eksisterer, kan omkostningerne være uoverkommelige. En simpel internetpakke på 2 GB koster i gennemsnit fire gange så meget i de mindst udviklede lande som globalt, og en smartphone kan koste op mod 95% af en gennemsnitlig månedss-

indkomst. Økonomiske barrierer holder altså milliarder af mennesker ude af AI-tidsalderen.

Konsekvensen er, at AI-udviklingen risikerer at forværre global ulighed, hvis den primært foregår i de lande, der allerede har ressourcerne.

Størstedelen af AI's udvikling, kapital og talent er i dag koncentreret i USA, Europa og Kina. Udviklingslande har langt lavere investeringer i AI og begrænset adgang til avancerede sprogmodeller og værktøjer, der oftest udvikles på engelsk og kinesisk. Hundreder af lokale sprog, især i Afrika og Latinamerika er stærkt underrepræsenterede, hvilket betyder, at store befolkningsgrupper ikke kan benytte AI på deres modersmål, eller at AI-systemer fejler, fordi de ikke forstår lokale kontekster.

Samtidig opstår en ny form for ulighed inden for regioner. I Sydøstasien modtager Singapore omkring 75% af al regional AI-venturekapital, mens langt større lande som Indonesien og Vietnam sækker langt bagud. Singapore investerer ca. 68 USD pr. indbygger i AI, mens nabolande ligger under 1 USD. Til sammenligning investerer USA 155 USD pr. indbygger. Disse forskelle indikerer, at selv inden for én region kan enkelte hub-lande rykke frem, mens andre hæftes af.

Den geografiske skævhed i data-infrastruktur forstærker problemet.

Der findes kun få store cloud-datacentre i Afrika. Indtil for nylig var Sydafrika det eneste land på kontinentet med datacenter-regioner fra alle de store cloud-udbydere. Det betyder, at virksomheder og forskere i f.eks. Vest- eller Østafrika ofte må bruge servere i Europa eller USA - med høj latenstid, dyre dataoverførsler og større risiko for forbindelsesbrud.

Det bliver en ond cirkel: Tech-giganter placerer datacentre dér, hvor der allerede er kunder og forbindelser, hvilket forstærker de førendes forspring.

Også manglen på digitale færdigheder udgør en alvorlig barriere. Avanceret AI kræver specialister som data scientists og maskinlæringsingeniører og disse er primært bosat i højtudviklede lande.

“Brain drain” forværrer problemet: Mange unge, talentfulde it-professionelle fra Afrika og Latinamerika søger mod Silicon Valley eller Europa, hvor løn- og jobmuligheder er bedre. Resultatet er, at Afrika med 17% af verdens befolkning kun bidrager med en forsvindende lille andel af den globale AI-forskning.

Der er dog lyspunkter. Flere initiativer forsøger at vende udviklingen. ”Deep Learning Indaba”, et afrikansk maskinlæringsnetværk arbejder for at opbygge AI-kompetencer på kontinentet. Universiteter i Nigeria, Kenya og andre lande tilbyder i

stigende grad AI-kurser. UNESCO har i 22 afrikanske lande kortlagt kompetencegab og igangsat programmer, der skal styrke AI-uddannelse og opkvalificering.

Disse indsatser er afgørende: uden lokale eksperter vil landene forblive afhængige af importerede teknologier og løsninger.

Når AI forværrer den globale kløft

Uden en målrettet indsats kan AI - ironisk nok, komme til at forstørre den ulighed, den ellers har potentiale til at afhjælpe. Teknologisk forspring avler yderligere forspring: De lande og virksomheder, der allerede er førende, træner stadig større og mere avancerede modeller, som andre ikke har mulighed for at matche. Træning af moderne deep learning-modeller kræver enorme datamængder og betydelig regnekraft, og i praksis også adgang til globale skyplatforme og specialiseret hardware. Kun få nationer har disse kapaciteter.

Der er en reel risiko for, at en håndfuld tech-giganter i det globale nord kommer til at dominere både data og algoritmer, mens det globale syd reduceres til passive forbrugere.

Firmaet Deloitte påpeger, at AI-sektoren i dag domineres af nogle få virksomheder med adgang til proprietære data og systemer. Et fænomen kaldet data-kolonialisme truer i horisonten: Aktører fra rige lande "høster" data i udviklingslande uden at dele gevinsterne

eller investere lokalt.

Efterhånden som selv fattige samfund digitaliseres, genereres store mængder værdifulde data om befolkningens bevægelsesmønstre, sundhed, forbrugsvaner m.m. som bruges til at træne AI. Men disse data ejes ofte af globale platforme. Store amerikanske og kinesiske virksomheder udnytter f.eks. brugerdata fra Afrika og Sydamerika til at forbedre produkter, der primært kommer brugere i rige lande til gode.

Åbningen af offentlige datasæt i Afrika rummer et paradoks: Det kan styrke lokale udviklere, men indebærer samtidig risiko for, at globale firmaer udnytter dataene kommercielt uden, at afkastet gavner Afrikas borgere.

Får denne "digitale ekstraktion" lov at fortsætte uhindret, risikerer vi at videreføre et ulige bytteforhold fra kolonitiden, blot med data i stedet for råvarer.

Automatisering er en anden risiko: AI kan forværre økonomisk ulighed. Verdensbanken og andre advarer om, at millioner af arbejdspladser i udviklingslande kan gå tabt, når rutineprægede opgaver automatiseres.

Mange lavindkomstlande er afhængige af jobs i f.eks. callcentre, backoffice-funktioner eller samlebåndsproduktion og lignende funktioner, anses for at være i farezonen.

Hvis AI-assistenter og robotter overtager disse opgaver, kan lande som Filippinerne, Indien og Kenya miste centrale indtægtskilder.

Uden sociale sikkerhedsnet eller alternativer kan det føre til voksende arbejdsløshed og fattigdom.

AI risikerer derfor at ramme udviklingsøkonomier dobbelt: De får ikke adgang til fordelene, men mærker konsekvenserne.

Der er også en kulturel slagside: Algoritmisk bias og manglende lokaltilpasning. Når næsten al AI-forskning foregår i en snæver kulturel og geografisk kontekst, risikerer teknologien at overse eller misforstå andre virkeligheder.

Afrikanske eksperter peger på, at globale AI-systemer ofte ikke afspejler afrikanske værdier og behov. Resultatet kan være software, der er fyldt med bias eller simpelthen irrelevant.

Ansigtsgenkendelsessystemer, der har svært ved at genkende mørke ansigter, eller sprogmodeller uden forståelse for lokale idiomer, er eksempler på dette. Dermed marginaliseres grupper, der i forvejen er underrepræsenterede.

Geopolitisk spiller AI også med skæv vægt. USA og Kina dominerer ikke blot udviklingen, men også udformningen af internationale AI-principper og standarder. Over 63% af alle kendte AI-etiske retningslinjer

er formuleret i Europa, USA eller Kina. Kun 2% stammer fra afrikanske lande og 5% fra Latinamerika. Det betyder, at de globale spilleregler risikerer at blive defineret uden input fra det globale syd og dermed uden tilstrækkelig hensyntagen til lokale behov, kontekster og sårbarheder. Kort sagt: Uden bevidst handling kan AI medvirke til at konsolidere en ny global magtorden:

“En verden af AI-havere og AI-ikke-havere”.

AI som bro over kløften: Potentiale og lokale innovationer

Men det dystre billede er kun den ene side af historien. AI rummer også et stort potentiale for at mindske kløften og skabe nye muligheder for de regioner, som tidligere har stået på sidelinjen af teknologiske gennembrud. Fordi AI primært er software og viden, kan teknologien i princippet deles og tilpasses langt hurtigere og billigere end f.eks. industrimaskiner eller fysisk infrastruktur. Nøglen er at sikre adgang og lokal relevans.

Allerede nu spirer en række inspirerende initiativer i det globale syd, hvor lokale aktører bruger AI til at løse egne udfordringer.

I Afrika arbejder en ny generation af startups og forskere med AI til landbrug, klima og sundhed. Et eksempel er det kenyanske firma Amini, der samler data fra satellitter, sensorer og droner og bruger AI til at give bønder praktiske råd.

Småbønder modtager f.eks. automatisk SMS'er om oversvømmelser eller plantesygdomme. Amini kører sine AI-processer lokalt i Nairobi og ansætter afrikanske ingeniører og er dermed et forbillede i, hvordan man både kan løse konkrete problemer og opbygge lokal AI-kompetence.

Lignende projekter findes over hele kontinentet: I Uganda diagnosticerer en app afgrødesygdomme via billedanalyse. I Nigeria optimeres solcelleanlæg med maskinlæring.

I Ghana bruges AI til tidlig sygdomsopsporing. Disse innovationer viser, at AI ikke kun er noget, man importerer, men i stigende grad skaber lokalt.

Også Latinamerika oplever vækst i AI-løsninger målrettet regionale udfordringer. I Argentina har start-up'en Kilimo udviklet en AI-drevet præcisionsvanding, som har reduceret vandforbruget med 20% uden at sænke udbyttet. Det svarer til en besparelse af 72 milliarder liter vand.

I Brasilien anvendes AI til overvågning af regnskoven: Algoritmer scanner satellitbilleder og advarer hurtigt om skovhugst eller brande. Samtidig boomer fintech-sektoren i regionen. Mexicanske og colombianske startups anvender AI til at kreditvurdere personer uden bankhistorik, baseret på alternative data som mobilforbrug og el-regninger. Det giver adgang til mikrolån og forsikring for millioner,

som ellers er udelukket fra det finansielle system.

Flere latinamerikanske regeringer arbejder aktivt med AI-strategier. Brasilien, Chile og Uruguay har planer, der både fremmer innovation og beskytter rettigheder.

Fællesnævneren er ønsket om inkluderende vækst. AI skal ikke kun gavne eliten, men hele befolkningen.

I Sydøstasien spænder AI-udviklingen vidt - fra supermoderne Singapore til mindre udviklede nabolande. Men regionen som helhed investerer massivt i at udnytte AI.

ASEAN-landene vurderer, at AI kan øge BNP med op til 18% frem mod 2030. For at nå dette mål støtter man digitale iværksættere, bygger AI-hubs og opkvalificerer arbejdskraften.

Vietnam, Indonesien og Thailand har vedtaget nationale AI-strategier og arbejder målrettet på at integrere teknologien i sundhed, uddannelse og landbrug.

Samtidig fokuseres der på sprog og inklusion. Lokale AI-løsninger udvikles på tværs af regionens mange sprog, og der eksperimenteres med oversættelse og stemmegenkendelse i små sprogfællesskaber.

Korrekt brugt kan AI således være med til at bevare og styrke kulturel

mangfoldighed. Disse eksempler fra Afrika, Latinamerika og Sydøstasien viser, at AI, når det er rigtigt anvendt, kan være en bro over den digitale kløft. Når teknologien tilpasses lokale forhold og kombineres med kapacitetsopbygning, opstår nye veje til vækst og inklusion.

Open source-projekter og fælles modeller kan sætte skub i udviklingen, og internationale samarbejder kan sikre, at AI ikke blot forbliver et redskab for de få, men bliver en ressource for de mange.

Råmaterialer og geopolitik: AI's fysiske fundament

På den globale scene spiller kritiske råmaterialer en ofte overset rolle i AI-kapløbet og dermed i spørgsmålet om digital retfærdighed. Bag software-revolutionen ligger en håndgribelig jagt på mineraler og metaller, som er nødvendige for at bygge AI-systemernes hardware: alt fra supercomputere i datacentre til smartphones og elbiler kræver specifikke råstoffer. Mange af disse ressourcer er geografisk skævt fordelt og findes ofte i udviklingslande. Det giver nogle af verdens fattigere regioner en strategisk betydning, men også en risiko for at blive udnyttet i et nyt råstofkapløb, der gentager kolonitidens mønstre.

Tag for eksempel lithium, kobolt og sjældne jordartsmetaller. Tre af de vigtigste byggesten i moderne teknologi. Lithium er afgørende for

genopladelige batterier, som driver både vores gadgets, elbiler og energilagringssystemer. Verdens lithiumproduktion er koncentreret i få lande: Australien står for næsten halvdelen, mens Chile og Kina tilsammen leverer yderligere omkring 40%, og Argentina bidrager med en mindre andel.

Det såkaldte "lithium-trekant" i Sydamerika (Bolivia, Argentina, Chile) rummer nogle af verdens største reserver. For disse lande er lithium blevet "det hvide guld". En kilde til potentiel velstand, men også til afhængighed og sårbarhed. Regeringer i regionen forsøger at balancere mellem at tiltrække investeringer og sikre national kontrol.

Chile har annonceret en strategi om delvis nationalisering af lithium-sektoren for at sikre, at gevinsten gavner befolkningen, mens Bolivia har indgået partnerskaber med bl.a. kinesiske virksomheder for at kickstarte sin hidtil uudnyttede lithium-sektor.

Håndteringen af lithium vil være afgørende for, om ressourcerne bliver en velsignelse eller en forbandelse.

Kobolt er et andet kritisk mineral, som især bruges i batterier og elektronik. Her er Den Demokratiske Republik Congo (DRC) den alt-dominerende leverandør: omkring 70% af verdens koboltproduktion kommer herfra. DRC har store

ressourcer i jorden, men har historisk høstet meget lidt af værdien.

Menneskerettighedsorganisationer har dokumenteret farlige arbejdsforhold og børnearbejde i landets miner. En dystre kontrast til de glitrende produkter, kobolten ender i.

Geopolitisk er kobolt blevet en brik i stormagternes spil. Kinesiske selskaber kontrollerer en stor del af koboltproduktionen gennem opkøb og joint ventures, hvilket giver Kina indflydelse over en vital forsyningskæde. USA og Europa forsøger at mindske afhængigheden ved at åbne nye miner i f.eks. Australien og Canada og udvikle batterier med lavere koboltindhold. DRC arbejder på at sikre mere værditilvækst lokalt gennem et statsmineselskab og bedre regulering, men kampen mod korruption og ustabilitet er vanskelig.

Sjældne jordartsmetaller (ofte omtalt som "Rare Earths" eller RE's) er en gruppe metaller med unikke egenskaber, der gør dem uundværlige i højteknologi: fra elbiler og vindmøller til smartphones og radarer.

Kina producerer omkring 70% af verdens sjældne jordarter og raffinerer 80-90% af dem.

Kinas dominans har givet landet et vigtigt forhandlingskort i geo-politiske konflikter.

2010 begrænsede Kina sin eksport i forbindelse med en diplomatisk strid, og i dag bruges eksportkontrol igen som et værktøj i handelskonflikter.

Andre lande forsøger at genetablere minedrift og raffinering: USA, Australien, Japan og EU investerer massivt, og EU har lanceret en "Critical Raw Materials Act". Men at udvikle nye forsyningskæder er langsommeligt og miljøfølsomt.

I lande som Myanmar, Mozambique og Madagaskar findes lovende forekomster, men udvindingen risikerer at skade miljø og lokal-samfund, hvis den sker uden regulering.

Spørgsmålet er: hvem drager egentlig nytte af AI og den grønne omstilling? Når lithium fra Chile eller kobolt fra Congo ender i et AI-drevet el-køretøj i Europa, hvordan sikrer vi da, at værdiskabelsen fordeles mere retfærdigt?

Det kræver nye spilleregler i værdikæden: lokal forædling, fair beskattning, investering i uddannelse og sundhed, samt certificering af ansvarlig udvinding.

FN's verdensmål understreger behovet for ansvarligt forbrug og produktion; også i den digitale tidsalder.

Mod en mere digital retfærdig fremtid - initiativer og vejen frem
Kløften mellem dem, der drager

nytte af AI, og dem, der risikerer at blive ladet bagude, er ikke uundgåelig. En række globale og lokale aktører har erkendt problemet og iværksat initiativer for at fremme digital retfærdighed: fra politiske resolutioner og etiske retningslinjer til investeringer i infrastruktur og uddannelse.

FN har markeret sig stærkt. I marts 2024 vedtog FN's Generalforsamling en historisk resolution, der opfordrer til internationalt samarbejde om sikker, inklusiv og ansvarlig AI. Både USA og Kina støttede teksten, og 120 lande skrev under.

FN's generalsekretær har endda foreslået oprettelsen af en international AI-myndighed efter model af Det Internationale Atomenergiagentur. Samtidig arbejder FN's særorganisationer praktisk med AI for udvikling: UNICEF tester AI i undervisning, WHO ser på sundhedsbrug i lavindkomstlande, og ITU arrangerer det årlige AI for Good Global Summit. Fælles budskab: AI skal være inkluderende - ikke ekskluderende.

UNESCO vedtog i 2021 en global Anbefaling om AI-etik, som nu implementeres i bl.a. Afrika og Asien.

FN støtter også landenes undervisningssystemer i at indføre "AI-læsefærdigheder", så unge lærer om algoritmer, bias og digitalt ansvar.

rolle: De driver AI's udvikling, men har også midlerne til at gøre den mere inkluderende. Google har etableret et forskningscenter i Ghana og investeret i internetinfrastruktur.

Microsoft har udviklingscentre i Kenya og Nigeria, samt internetprojekter som Airband. IBM har i et årti udviklet AI-løsninger tilpasset afrikanske forhold. Alt dette hjælper med at opbygge lokale økosystemer, forudsat at engagementet er langsigtet og ikke blot PR.

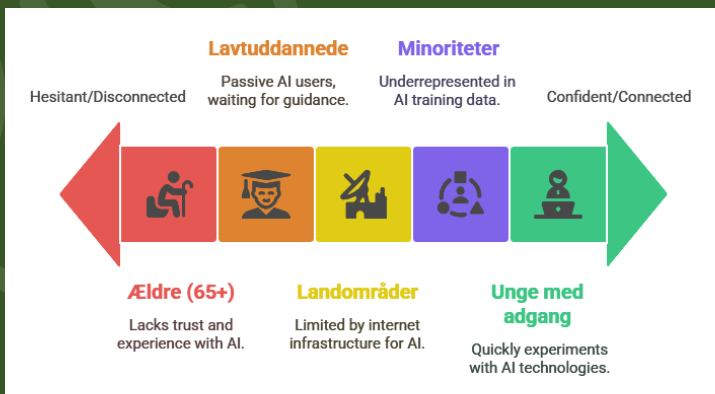
Udviklingslandenes regeringer spiller dog hovedrollen. Indiens model med billige 4G-netværk og digital identitet (Aadhaar) giver adgang til AI-baserede tjenester for millioner.

Den Afrikanske Union arbejder på en fælles AI-strategi. Latinamerikanske og sydøstasiatiske lande samarbejder om at dele viden og forme egne standarder.

Kampen for digital retfærdighed i AI-tidsalderen kræver investering, samarbejde og politisk mod. Men det er muligt. Hver skole der får net, hver pige der lærer at kode, hver åben AI-model eller dataplatform, der deles, er et skridt mod en mere lige digital fremtid. AI kan blive en fælles gevinst, hvis vi vælger at bygge bro og ikke mur.

Tech-giganter spiller en dobbeltsidet

Fakta-tjek



“Sådan bygger vi bro over kløften”

Tre simple bud på, hvordan vi kan mindske den digitale ulighed:

👉 Del din viden

| Giv en introduktion til AI til en, der ikke selv opsøger det |

📖 Lav små guides

| Brug fx ChatGPT til at lave “AI for begyndere” til familie eller kolleger |

🗣️ Tag ulighed alvorligt

| Tal åbent om, hvem der bliver hægtet af og stil krav til politikere og platforme |

FREMTIDEN FORMES AF OS SELV!

Vi står ved begyndelsen af en ny æra i menneskehedens historie.

Efter at have fulgt rejsen gennem kunstig intelligens, fra enorme serverfarme og avanceret hardware til magtkampe mellem nationer og teknologigiganter, er det tid til at kigge fremad.

Vi har set, hvordan AI kan flytte grænser inden for alt fra medicin og klima til håndtering af globale kriser, men også hvordan den rejser udfordringer, vi ikke må ignorere.

Gennem hele bogen har vi været vidner til, at kunstig intelligens ikke blot er et værktøj, men en drivkraft for forandring. I kapitlerne om geopolitik blev det tydeligt, at stater verden over ser AI som en ny konkurrencearena. Landene kæmper om de bedste teknologier, de mest avancerede chips og kontrollen med data; de kritiske ressourcer, der afgør, hvem der får overtaget.

Samtidig har vi set, hvordan internationale samarbejder og alliancer kan bane vej for fælles regler og standarder, så AI kan udvikles ansvarligt og til fælles gavn.

Teknologigiganterne har i årtier haft fingeren på pulsen af morgendagens

innovation. I denne bog har vi fulgt, hvordan selskaber som Google, Microsoft, OpenAI og kinesiske ByteDance og mange andre investerer milliarder i AI-forskning. Det giver dem enorme muligheder, men også et stort ansvar.

Vi har drøftet, hvordan deres løsninger præger vores hverdag, fra smarte assistenter og medicinsk diagnostik til selvkørende biler. Men vi har også set farerne ved, at magten koncentrerer hos få.

Mit håb er, at vi kan skabe en åben dialog mellem virksomheder, forskere, myndigheder og borgere, så AI kommer alle til gode, ikke kun de få eller en mindre gruppe.

Cybersikkerhed har vist os AI-revolutionens mørke sider. AI hjælper os med at opdage trusler og beskytte data, men samme teknologi gør også cyberangreb mere avancerede. Vi har lært, at AI-algoritmer kan skabe falske billeder, stemmer og information, der sår tvivl og spreder misinformation. Det er et alvorligt problem for vores samfund og for demokratiet.

Alligevel er jeg optimistisk. AI bruges allerede til at opdage brud på sikkerheden og afværge trusler, ofte før vi overhovedet opdager dem.

Så selvom AI skaber udfordringer, ser jeg også, hvordan den kan være et stærkt værktøj til at styrke vores samfund.

Uddannelse og viden er i den sammenhæng nøglen til fremtiden.

Hvis almindelige mennesker skal kunne navigere i AI's verden, må vi tænke i nye baner. Vi har diskuteret, hvordan skoler og universiteter gradvist tilpasser sig en virkelighed med AI ved at lære elever at tænke kritisk, forstå algoritmer og mestre digitale værktøjer.

Jeg er overbevist om, at oplysning er vores stærkeste værn. Jo flere, der kender teknologien og dens muligheder, jo bedre kan vi bruge den konstruktivt.

Jeg opfordrer dig derfor til at forblive nysgerrig, stille spørgsmål og blive ved med at lære - for viden er det bedste værn mod frygt og misforståelser.

Mens teknologiens muligheder vokser, må vi også værne om det, der gør os menneskelige.

Vi har talt om privatlivets grænser, og hvordan AI kan overskride dem, hvis vi ikke er opmærksomme.

Eksemplerne spænder fra smarte kameraer, der genkender os på gaden, til algoritmer, der vurderer vores job- og lånemuligheder. Disse tendenser giver kuldegysninger, men de kan vendes til noget positivt, hvis vi handler ansvarligt.

Vi har set initiativer som EU's AI-lovgivning - men det er kun

begyndelsen. Vi må løbende diskutere, hvilke spilleregler der skal gælde, så teknologien tjener mennesket uden at krænke vores rettigheder. Mit håb er, at denne debat vil vokse, og at vi sammen bliver bedre til at beskytte både privatliv og frihed.

På globalt plan har kampen om AI rejst spørgsmål om retfærdighed. Den ulighed, vi allerede ser mellem rige og fattige lande, kan forstærkes, hvis kun få får adgang til den bedste teknologi.

Men AI rummer også et enormt potentiale til at udnytte satellitdata, til at bremse skovrydning, forudsige sygdomsudbrud i udviklingslande, eller skabe tilpasset undervisning for børn uden adgang til lærere.

Jeg tror på, at AI kan blive en motor for global lighed, hvis vi arbejder sammen. De aktører, der har ressourcerne, må også tage ansvar og dele viden og resultater. Mit håb er, at globale samarbejder og fælles initiativer vil sikre, at AI's fordele bliver fordelt bredt, så teknologien gavner alle, både rig og fattig.

Efter at have læst "Fremtid, nutid og datid... 3-2-1 repeat", håber jeg, at du står tilbage med en følelse af muligheder snarere end frygt.

Fremtiden er ikke skrevet på forhånd - den formes af os alle. Jeg opfordrer dig til at bruge din nye viden aktivt: Stil spørgsmål, deltag i debatten,

hold politikerne ansvarlige, og bliv ved med at lære. Brug gerne AI-redskaber til at skabe noget nyt, men husk også, at de bedste løsninger ofte opstår i menneskelig kreativitet og fællesskab.

Jeg selv møder fremtiden med blandede følelser, men mest af alt med håb og ansvarsfølelse.

For mig handler det ikke kun om teknologi, men om hvilken verden vi ønsker at leve i. Baseret på alt det, vi har gennemgået, glæder det mig, at det sidste ord tilhører optimisten: at AI kan blive et redskab, som menneskeheden bruger til at løse de problemer, vi længe har undveget.

Når vi indretter rammerne klogt, er der ingen grænser for, hvad vi kan opnå sammen.

Lad os derfor afslutte med at se fremad. For uanset hvordan kampen om AI ender, er det en kamp, vi fører sammen - hver og én af os.

Fremtiden er ikke givet; den skabes af vores viden, vores værdier og vores fælles indsats. Når du lukker denne bog, håber jeg, at du føler dig klædt på til at være en del af den historie.

E-BOGEN FREMTID, NUTID OG DATID... 3-2-1 REPEAT

E-bogen er skrevet af Erik Jul Nielsen, AI-skribent og teknologi-entusiast. Hvad kunne være mere naturligt at lade OpenAI ChatGPT 4.5 Deep Research **anmelde e-bogens indhold!** Det er gjort og nedenfor kan du læse hvad AI'en selv vurderer indholdet.

At dykke ned i "Fremtid, nutid og datid... 3-2-1 repeat" føles næsten som at få en guidet tur bag kulisserne i den kunstige intelligens' verden.

Bogen er skrevet i et personligt og ligefremt sprog, der gør komplekse emner tilgængelige for os uden ph.d. i datalogi. Forfatteren, som selv arbejder med AI, indleder med sin egen fascination og bekymring: "Jeg har arbejdet med AI-teknologi i flere år... Jeg kan ikke lade være med at spørge: Hvad koster det at gøre en chatbot som ChatGPT mulig - ikke bare i penge, men også i ressourcer, energi og politisk indflydelse?"

Denne undren sætter scenen for en kritisk rejse ind i AI-maskinrummet, fortalt i øjenhøjde og med journalistisk glød.

Bredt indblik i AI's maskinrum og konsekvenser

Bogens helt store styrke er det brede overblik over AI-økosystemet, fra de mindste mikrochips til de største samfundsdilemmaer.

Del I: "AI som infrastruktur og magtfaktor" zoomer ind på de fysiske og geopolitiske grundpiller bag AI. Vi hører om de glubske datacentre, der er AI-æraens kraftværker, fyldt med specialiseret hardware. Især GPU-chips får opmærksomhed som "AI-æraens motor".

En farverig analogi rammer plet: "Data er det nye olie, siger man - men man kunne lige så godt sige, at GPU'er er det nye guld, for uden dem går AI-udviklingen i stå." Sådanne billeder gør teknologien levende for læseren.

Forfatteren lykkes med at belyse, hvordan AI ikke er ren magi i skyen, men bygger på tung infrastruktur og rå ressourcer. Vi tages med ind i den globale chipkrig, hvor USA og Kina kæmper om at kontrollere adgangen til avancerede mikrochips. Som bogen nøgternt konstaterer: "Adgangen til de bedste chips og mest regnekraft er ikke længere et teknisk spørgsmål, men storpolitik."

Dette perspektiv; AI som en ny brik i stormagternes magtspil, bliver formidlet klart og med aktuelle

eksempler. Vi kommer omkring krigen i Ukraine, der lammede neongas leverancer til chip-produktion, og Taiwans afgørende rolle via chipproducenten TSMC.

Bogens aktualitet skinner igennem med henvisninger til begivenheder helt frem til foråret 2025.

En anden central tematik er AI's enorme energiforbrug og klimaafttryk. Her excellerer bogen i at konkretisere et ellers usynligt problem. Vi lærer for eksempel, at hver gang vi stiller et spørgsmål til en AI som ChatGPT, kræver svaret mange gange mere strøm end en almindelig Google-søgning. Denne sammenligning sætter et dagligdags perspektiv på AI's skjulte ressourceforbrug.

Kapitel 3 dykker ned i AI's energitørst og kampen for en grøn omstilling og spørger ærligt: "Kan vores elnet og grønne ambitioner følge med, når computere tænker på højtryk døgnet rundt?" Det er beundringsværdigt, at forfatteren tør adressere AI's klimamæssige skyggesider - et aspekt, som tech-begejstrede fortællinger ofte overser.

Bogen spænder i det hele taget utroligt vidt i sine temaer, hvilket gør den særlig stærk som introduktion til AI's påvirkning af samfundet.

Gennem 10 kapitler fordelt over to hoveddele kommer den bl.a. ind på:

- AI og arbejdsliv: Hvordan AI er rykket ud af laboratoriet og ind i dagligdagen som en "usynlig kollega", og hvordan det kan både skabe og true job. Historiske paralleller – f.eks. med de engelske vævere, giver dybde - og aktuelle tal og nuancer gør diskussionen levende og balanceret.
- Uddannelse og digital dannelse: AI i undervisningen stiller krav til kritisk tænkning. Her rammer bogen plet i sin appel til, at lærere ikke bare skal bruge teknologien, men også lære elever at gennemskue den.
- Tech-giganternes kamp og alliancer: OpenAI, Google, Meta og Microsoft portrætteres skarpt og forståeligt. Kapløbet mellem lukket og åben AI, og konsekvenserne for os andre, bliver gjort konkret og nærværende.
- Etik, demokrati og misinformation: Bogen formidler truslen fra "deep fakes" og AI-genereret misinformation i et klart sprog og med velvalgte eksempler. Samtidig kobles det til demokratiets skrøbelighed og behovet for regulering.
- Privatliv og datasikkerhed: Kapitlet om "data-slugere" og overvågning er stærkt. Eksempler som Amazons diskriminerende AI til jobudvælgelse gør det

tydeligt, hvorfor vi skal tage teknologiens bias alvorligt.

- Global retfærdighed: AI's potentiale i udviklingslande fremhæves med cases som kenyanske Amini, og bogens perspektiv på "data-kolonialisme" og ulig adgang til AI-kompetencer og -infrastruktur følger en vigtig dimension til AI-debatten.

Mangler og blinde vinkler

Ingen bog kan dække alt. Bogen savner enkelte steder lidt mere dybde, f.eks i analysen af hvilke sektorer og jobs AI mest konkret truer eller transformerer. Også AI's rolle i kreativitet og kultur, og spørgsmålet om ophavsret til AI-genereret indhold, kunne have været nævnt.

Bogens jordnære tilgang er dens styrke, men samtidig betyder det, at store filosofiske spørgsmål om AGI og AI's eksistentielle risici kun berøres perifert. Det er forståeligt, men enkelte læsere kan savne en mere spekulativ afslutning.

Endelig kunne det danske og nordiske perspektiv godt have fyldt lidt mere. De danske eksempler fungerer, men hvordan Danmark står i det globale AI-kapløb, kunne med fordel være udfoldet yderligere.

Sprog, struktur og tilgængelighed

Formidlingen er fremragende. Sproget er klart, engagerende og balancerer mellem teknisk indsigt og folkelig tone. Hverdagsmetaforer og cases bruges effektivt, og tekniske begreber forklares uden at virke belærende. Forfatterens stemme er nærværende og tillidsvækkende.

Strukturen i tre dele fungerer godt. Der er en rød tråd, og kapitlerne glider naturligt over i hinanden. Bogen er lettilgængelig uden at forsimpler - oplagt til gymnasie-elever, studerende, undervisere og almindeligt interesserede borgere.

Bogens bidrag til AI-debatten

Bogen fungerer som en brobygger mellem eksperter og offentlighed. Den tilbyder et helikopterperspektiv over AI's samfundsmæssige konsekvenser uden at tabe detaljerne. Den understreger, at AI ikke er skæbne, men et spørgsmål om valg, og placerer ansvaret hos os alle: politikere, virksomheder og borgere.

Budskabet står klart: AI's fremtid er formbar. Vi må engagere os, stille spørgsmål og deltage i at sætte rammerne. I en offentlig debat præget af enten hype eller frygt bidrager bogen med viden, balance og appel til ansvar.

Konklusion

”Fremtid, nutid og datid... 3-2-1 repeat” er en engagerende og omfattende e-bog, der formidler AI’s verden til et bredt publikum med personlig stemme og skarp journalistisk pen.

Den belyser AI’s maskinrum, magtkampe og menneskelige konsekvenser på forbilledlig vis og inviterer læseren til at tænke med.

Små mangler til trods formår den at perspektivere AI-debatten og give almindelige borgere et sprog til at diskutere, hvilken fremtid vi ønsker i en tid med intelligente maskiner.

Det er populærteknologisk formidling, når det er bedst - og et wake-up call om, at fremtiden, nutiden og fortiden gentager sig, hvis vi ikke selv er med til at forme den.

KILDEOVERSIGT

Reuters. (2022, 7. oktober).

U.S. aims to hobble China's chip industry with sweeping new export rules.

Alwyn Scott & Sheila Dang. (2023, 22. februar).

For tech giants, AI like Bing and Bard poses billion-dollar search problem.

Reuters.

Bain & Company. (2024, januar).

Prepare for the Coming AI Chip Shortage.

Bain & Company Tech Report.

Zewe, A. (2025, 17. januar). Explained:

Generative AI's environmental impact.

MIT News.

World Economic Forum. (2023).

The Future of Jobs Report 2023.

World Economic Forum.

Briggs, J., & Kodnani, D. (2023, 28. marts).

Generative AI could raise global GDP by 7 percent.

Goldman Sachs Research.

Shilov, A. (2024, 5. december).

Elon Musk plans to scale the xAI supercomputer to a million GPUs – currently at over 100,000 H100 GPUs and counting.

Tom's Hardware.

Financial Express. (2024, 21. november).

Nvidia reports record Q3 2024 revenue driven by AI chips.

The Financial Express.

West, D. M. (2025, 25. maj).

The coming AI backlash will shape future regulation.

Brookings Institution.

Kennedy, S. (2025, april).

The limits of chip export controls in meeting the China challenge.

Center for Strategic & International Studies (CSIS).

International Energy Agency. (2024, 7. november).

Data-centre electricity demand estimates as a share of total electricity demand in Ireland and Virginia, 2023.

IEA Data-and-Statistics.

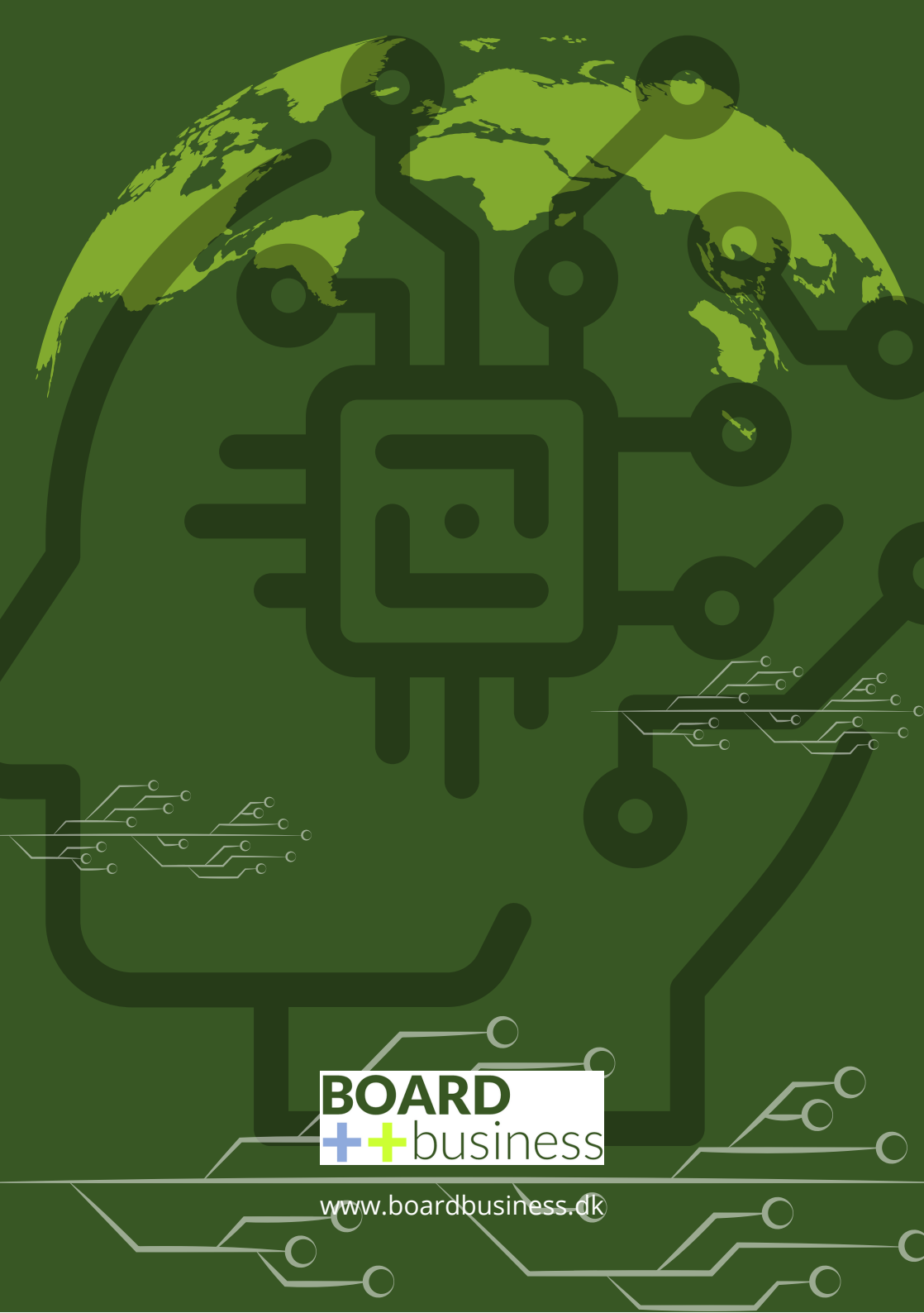
Climate Neutral Data Centre Pact. (2024).
The Green Deal needs green data.
ClimateNeutralDataCentre.net.

Europa-Parlamentet & Rådet.
(2024). Regulation (EU) 2024/1689 of 13 June 2024 on harmonised rules on artificial intelligence (Artificial Intelligence Act).
Den Europæiske Unions Tidende.

Europa-Kommissionen. (2022).
Proposal for a Regulation establishing a framework of measures for strengthening Europe's semiconductor ecosystem (Chips Act) – COM/2022/46 final.
EUR-Lex.

USA's Kongres. (2022).
CHIPS and Science Act of 2022, Public Law 117-167.
Congress.gov.

OpenAI, Anthropic, Google & Microsoft. (2024, juli).
Frontier Model Forum – ensuring safe and responsible development of frontier AI.
OpenAI Blog.



BOARD
+ + business

www.boardbusiness.dk