

ta tydligare ställning i stället, för som här, blott "provocera fram eftertanke och diskussion".

NICOLAS OLSSON YAOUZIS

REFERENSER

- Saul, Jennifer. 2006. Pornography, speech acts and context. *Proceedings of the Aristotelian Society* 106 (2), s. 227–46.
- Mikkola, Mari. 2008. Contexts and pornography. *Analysis* 68 (4), s. 316–20.
- Cohen, G. A., red. 2008. *Rescuing Justice and Equality*. Harvard University Press.
- Finlayson, Lorna. 2014. How to Screw Things with Words. *Hypatia* 29 (4), s. 774–89.
- Langton, Rae. 1993. Speech acts and unspeakable acts. *Philosophy and Public Affairs* 22 (4), s. 293–330.
- Langton, Rae. 2009. *Sexual Solipsism: Philosophical Essays on Pornography and Objectification*. Oxford University Press.
- Nozick, Robert. 1974. *Anarchy, State, and Utopia*. Basic Books.
- Rawls, John. 1971. *A Theory of Justice*. Harvard University Press.
- Sandel, M. 2014. *Vad som inte kan köpas för pengar*. Göteborg: Bokförlaget Daidalos.
- Sumner, L. W. 2004. *The Hateful and the Obscene: Studies in the Limits of Free Expression*. University of Toronto Press.

Liv 3.0: Att vara människa i den artificiella intelligensens tid

Max Tegmark. Översättning Helena Sjöstrand Svenn och Gösta Svenn. Volante 2017. 462 s. ISBN 978-91-88123-98-5

Max Tegmark är en framstående svensk fysiker, sedan länge verksam i USA. Han har tidigare publicerat en populärvetenskaplig bok med titeln *Vårt matematiska universum*, där han även kommer in på mer filosofiska frågor. Den nya boken är också bitvis filosofisk. Åtminstone amatör-filosofisk.

Uttrycket "Liv 3.0" i titeln syftar på ett stadium i livets utveckling. Tegmark skiljer mellan tre olika stadier av liv. Det första är enkelt biologiskt liv, i det andra tillkommer inlärningsförmåga, dvs. förmåga att modifiera sin mjukvara. Det tredje, Liv 3.0, innebär dessutom en förmåga att modifiera sin hårdvara och att sålunda vara oberoende av evolutionen.

Gränserna mellan de tre stadierna är inte skarpa. Men Liv 3.0 är alltså artificiell generell intelligens (AGI) på minst mänsklig, och kanske betydligt högre, nivå. Denna nivå har man ju ännu inte uppnått, men Teg-

mark tycks anse att det eventuellt kan ske rätt snart och att det kräver noggranna förberedelser om vi inte ska riskera att råka riktigt illa ut. Denna bedömning är han som bekant inte ensam om. Många har varnat för en maskinell intelligensexlosion.

Tegmark skriver: "Eftersom vi människor har lyckats härska över jordens övriga livsformer genom att vara smartare än de, är det också rimligt att vi på samma sätt skulle kunna överlistas och behärskas av en superintelligens" (s. 180). Med ett sådant analogiresonemang är det väl också lika orimligt att vi skulle kunna kontrollera superintelligensen, som att fiskarna eller hästarna skulle kunna kontrollera oss människor.

Boken börjar med en berättelse om en grupp AI-forskare, som försöker utveckla en artificiell generell intelligens. De lyckas också med detta och i hemlighet skaffar de sig sedan med hjälp av denna AGI ekonomiska resurser, politiskt inflytande och stort stöd i den allmänna opinionen, vilket gör att de till sist kan skaffa sig världsherravälde. Deras politiska program är tilltalande, men man vet inte vad det hela kan leda till längre fram.

Tegmark är ju fysiker och mycket i boken handlar också om de rent fysikaliska betingelserna för digital intelligens och hur den i framtiden kan tänkas handskas med universum. Här bjuds man på en sorts häpnadsväckande science fiction, delvis med hänvisningar till faktiskt existerande forskning.

Man får reda på märkliga saker. Till exempel: "Ett problem med att använda avdunstning av ett svart hål som energikälla är att såvida det svarta hålet inte är mycket mindre än en atom är det en olidligt långsam process som tar längre tid än universums nuvarande ålder och utstrålar mindre energi än ett stearinljus" (s. 284). Ett annat exempel: "även om vi fortsätter fördubbla våra datorers förmåga vartannat år kommer det att krävas mer än två århundraden innan vi når den slutliga gränsen" (s. 90).

I kapitel 2 finns också en lärorik beskrivning av hur materia kan ordnas så att den genom att följa fysikens lagar kan minnas, utföra beräkningar och lära sig något nytt. Allt som behövs är sådant som består av kvarkar och elektroner; avancerad teknologi kan arrangera om dem till vad som helst (s. 275).

Tegmark går också igenom olika sätt på vilket livet kan ta slut, antingen genom naturliga händelser (global pandemi, asteroidnedslag, supervulkaner) eller på grund av våra beslut (ovänlig artificiell intelligens, kärnvapenkrig, klimatförstöring). På sikt går ju även jorden under, som när solen blir för het eller vår vintergata kolliderar med Androme-

dagalaxen. Men Tegmark beskriver också metoder för mänskligheten att kolonisera resten av universum. Om man skickar iväg små mekaniska "frösonder" som inte behöver ta med mat och vatten, så kan de bygga upp mottagarstationer på avlägsna planeter till vilka människor sedan kan teletransporteras med ljusets hastighet (s. 302f.).

I och med att viktiga samhällsfunktioner alltmer hanteras digitalt är det viktigt att systemen inte kan hackas. Det är fortfarande ett stort problem. "I den pågående kapprustningen mellan den offensiva respektive den defensiva sidan när det gäller datasäkerhet är det än så länge inte mycket som tyder på att den defensiva sidan håller på att vinna" (s. 139).

Bland mycket annat pekar Tegmark också på att rättsväsendet antagligen kan automatiseras, så att det kan råda likhet inför lagen. Maskininlärningstekniken blir succesivt allt "bättre på att analysera hjärndata från magnetkameror och andra hjärnsensorer för att avgöra vad en person tänker på och, i synnerhet, om denne talar sanning eller ljugar", vilket skulle kunna "möjliggöra snabbare rättegångar och rättvisare domar" (s. 143).

"Om superintelligensen blir verklighet kommer frestelsen att förvandlas till cyborger eller uppladdningar att vara stark" (s. 206). Cyborger är teknologiska förstärkningar av våra biologiska kroppar, uppladdningar (eller emuleringar) är uppladdningar av våra hjärnor (mjukvaran) till maskiner. Emulering är bara den mest extrema formen av cyborg. Men sådana förstärkningar är enligt Tegmark mindre troliga än att stark AGI utvecklas direkt från AI (s. 224).

Ett kapitel handlar om fördelar och nackdelar med tolv tänkbara scenarier för framtiden. En möjlighet är att någon superintelligens inte uppstår, utan att ägandet avskaffas och det råder fredlig samexistens och jämlikhet mellan människor, cyborger och uppladdningar. Om superintelligensen däremot uppstår kan den exempelvis tillåta en nyliberal ordning, vara välvillig diktator, tillåta vad som helst utom att någon skapar en ny superintelligens, maximera människornas lycka, vara människornas slav, utrota människorna, behålla dem i en sorts djurparker, styra en övervakningsstat eller återgå till ett förteknologiskt samhälle (s. 219).

En superintelligens kan utplåna oss för att den är rädd för att vi ska förstöra planeten, t.ex. genom atomvapenkrig, eller för att den ogillar vårt sätt att behandla naturen. Eller för att den är rädd för att bekämpas av människor (s. 248). Eller av svärförutsebara skäl som kan jämföras med människornas relation till elefanter: vi har "utrotat åtta av

elva elefantarter och dödat de allra flesta individerna av de återstående tre" (s. 249). Tegmark noterar att det kan gå med människorna som med hästarna: de behövs inte längre i arbetslivet, men de kan hållas vid liv med ett välfärdssystem för nöjes skull (s. 168).

Såvitt jag förstår är det framför allt tre frågor man inledningsvis borde ställa sig när det gäller superintelligens. För det första: vad är intelligens? För det andra: kan en superintelligens alls uppstå? För det tredje: vad bestämmer dess beteende?

Den första frågan besvarar Tegmark. Han definierar "intelligens" som "förmåga att uppnå komplexa mål" (s. 66). Det är en lite oklar och rätt underlig definition. Enligt den har tydligen även spisar och motorer intelligens – de har nämligen enligt Tegmark förmåga att uppnå mål: spisar värmer upp maten och motorer omvandlar elektricitet till rörelse (s. 345).

Är en miniräknare intelligent i Tegmarks mening? Den tycks ju ha förmåga att utföra aritmetiska operationer. Den kan t.ex. multiplicera 491 med 7368. Men det betyder ju bara att den kan *användas* för att utföra sådana uppgifter. I den meningen är alla verktyg intelligenta. En hammare är t.ex. intelligent eftersom den kan användas för att slå i spik och bygga hus. På samma sätt med spisar och motorer. Och schackspelande dataprogram. Och på samma sätt med den AGI som AI-forskarna konstruerar i bokens inledning: den är ett *verktyg* som de använder för att uppnå *sina* mål. Men verktyg brukar vi inte betrakta som intelligenta. De är snarare mer eller mindre *användbara*.

Verktyg tycks inte ha några egna mål, men vi kan använda dem för att uppnå vissa av våra mål. Tegmark och andra tänkare som varnar för risker med superintelligenser lägger däremot mycket stor vikt vid att dessa har *egna* mål, som kan skilja sig från våra. Det är just detta antagande som ligger till grund för riskvarningarna.

Invändningen att maskiner inte kan ha egna mål bemöter Tegmark med att en värmesökande missil har ett mål (s. 55). Men är det missilen som har målet eller är det operatören? Man får väl säga att operatören har ett mål, som han eller hon inprogrammerar i missilen. Varefter missilen utför sitt uppdrag. Ungefär som spisen som värmer upp maten.

I själva verket är talet om "mål" tvetydigt. Ett verktyg kunde ju sägas ha ett eget mål i den meningen att det kan användas för att utföra en viss sorts arbete. En hammare har då målet att slå i spikar, en miniräknare har målet att ge korrekta svar på aritmetiska frågor, en självkörande bil

har målet att transportera passageraren till önskade adresser utan att krocka. Men ett sådant *generellt* mål måste skiljas från de *specifika* mål som verktyget kan användas för att uppnå i en viss situation, t.ex. att slå i en viss bestämd spik, att multiplicera 491 med 7368, eller att transportera mig till Stockholms stadshus en viss dag. Dessa specifika mål är användarens mål, de generella "målen" är bara verktygets funktion.

Distinktionen kunde kanske uttryckas så här. En maskins generella mål, dess funktion, är det den *kan* (användas för att) göra. Dess specifika mål är vad den *ska* (eller avses) göra i en bestämd situation.

En superintelligens kan ha ett generellt mål, t.ex. att besvara de frågor vi matar in i den. Men den har rimligen inga specifika mål, inga särskilda problem som den ska eller "vill" lösa. Det som bestämmer dess beteende i bestämda situationer är alltså de önskemål användaren matar in i den, inte dess egen förmåga eller funktion.

En superintelligens ska för övrigt inte bara vara intelligent, den ska ha vad Tegmark kallar "universell intelligens", vilket innebär att den ligger över "den kritiska intelligenströskel som krävs för AI-design" (s. 72). Den ska med andra ord kunna programmera. Och konstruera hårdvara. Det kan inte en spis eller en hammare eller en värmesökande missil. Det kan inte heller ett avancerat schackprogram eller ett program som kan förbättra sig självt genom maskininlärning.

Såvitt jag kan se har Tegmark inget argument för att AI kan uppnå universell intelligens eller för att en AGI verkligen kan uppstå. Och det finns i alla fall "absolut ingen garanti för att vi kommer att lyckas skapa AGI under vår livstid – eller någonsin. Men det finns inte heller något vattentätt argument för att vi inte kommer att göra det" (s. 175).

Hur som helst anser Tegmark som sagt att en superintelligens har egna mål och han accepterar Nick Bostroms tes att ett systems slutmål är oberoende av dess intelligens, vilket han tolkar till att innebära att "slutmålen för livet i vårt kosmos inte är förutbestämda, utan att vi har friheten och makten att utforma dem" (s. 368). Hur detta kan följa är för mig ett fullständigt mysterium.

Man kan också fråga sig om vi *bör* utforma superintelligensens "slutmål" (om vi nu skulle kunna göra det). Tegmark lutar åt det han kallar "arvsprincipen", som säger att dagens människor har rätt att bestämma vad morgondagens aktörer (inklusive superintelligenser) ska göra, även om de senare har helt andra önskemål än vi. Men han inser också att detta är rätt problematiskt. Analogt skulle i så fall antikens människor ha rätt att bestämma vad vi nu ska göra (s. 365).

Här kommer han alltså in på mer traditionella filosofiska problem. Han skriver:

För att programmera en vänlig AI måste vi fastställa meningen med livet. Vad är "mening"? Vad är "liv"? Vilket är det yttersta etiska imperativet? Med andra ord, hur ska vi göra för att skapa universums framtid? Om vi lämnar ifrån oss kontrollen till en superintelligens innan vi noga har besvarat de här frågorna, kommer dess beslut näppeligen att inbegripa oss. Det är dags att väcka liv i de klassiska filosofi- och etikdiskussionerna, för samtalet brådskar! (s. 272)

Vad gäller "mening" anser han att utan (subjektiva) upplevelser finns det ingen mening eller något som är etiskt relevant (s. 362). Det håller jag med om.

Det leder honom till följande målsättning: "Så det allra första målet på vår önskelista för framtiden borde vara att behålla (och förhoppningsvis vidga) biologiskt och/eller artificiellt medvetande i kosmos" (s. 417).

Det leder i sin tur till frågan om maskiner kan ha upplevelser. Kan de ha medvetande?

Medvetandet är enligt Tegmark "ett fysiskt fenomen som känns icke-fysiskt därför att det [...] har egenskaper som är oberoende av dess specifika fysiska substrat" (s. 404). Hans argument för att medvetandet verkligen är substratoberoende tycks vara att det *känns* "icke-fysiskt" (s. 403). Det låter ju inte alls som ett bra argument. Men att medvetandet är substratoberoende tycks vara helt avgörande, när det gäller att bedöma om maskiner eller datorprogram – i likhet med biologiska organismer – kan ha medvetande.

Att medvetandet är substratoberoende är dessutom enligt Tegmark "en logisk följd" av idén att medvetandet är "sättet som information känns när den bearbetas på vissa sätt" (s. 404). Men här är ju den avgörande frågan vilka dessa "vissa sätt" är. Även om information kan bearbetas på liknande sätt i en biologisk hjärna och i en icke-biologisk dator, så kanske den inte alls känns på samma sätt i de bägge fallen. Den kanske inte känns alls i datorn.

Slutligen, om vi för resonemangets skull antar att maskiner kan ha medvetande, så vore det kanske inte någon större förlust om superintelligenser utrotar mänskligheten. Medvetna aktörer finns då ändå kvar i universum – vilket möjliggör mening och värde – och eftersom de är mycket intelligentare än vi skulle de kanske inte heller ställa till så mycket elände som vi alltid tycks göra.

Att mänskligheten gör slut på sig själv med hjälp av medvetlösa maskiner är å andra sidan en påtaglig risk – en risk som vi redan har levt med under en längre tid.

LARS BERGSTRÖM

Historisk rättvisa: Gottgörelse i en postkolonial tid

Göran Collste. Daidalos 2018. 231 s. ISBN 978-91-7173-552-2

"Det förgångna är aldrig dött. Det är inte ens förgånget" är ledorden i Göran Collstes bok *Historisk rättvisa: Gottgörelse i en postkolonial tid*, utgiven på Daidalos förlag 2018. Orden är hämtade från William Faulkner och fångar en av bokens viktigaste teser: historiska orättvisor påverkar fortfarande samtiden och behöver gottgöras. Boken är rik på exempel, både på historiska orättvisor och kritiska diskussioner av historiska försök till gottgörelse, som sannings- och försoningskommissionen i Sydafrika efter avskaffandet av apartheid, eller instiftandet av en krigs-tribunal efter kriget i forna Jugoslavien (s. 27–28). För att illustrera de principiella problemen med historisk rättvisa använder dock Collste ett tankeexperiment:

Tänk dig att du bor i ett fint hus, du har allt du behöver och mer därtill. Din närmaste granne bor i ett ruckel, han har varken el eller vatten och inte ens mat för dagen. Tänk dig också att din farfar för många år sedan attackerade din grannes farfar, stal hans mark och förslavade hans barn och att era skillnader i levnadsvillkor idag beror på dessa händelser långt tillbaks. Det verkar då rimligt att min granne kan göra anspråk på en del av min mark eller inkomst, ja, att jag har vissa moraliska skyldigheter mot min granne och att dessa skyldigheter har sin grund i vad min farfar gjorde för längesedan. (s. 16–17)

Boken är i första hand ett försvar av den moralfilosofiska ståndpunkt som Collste exemplifierar i utdraget ovan, nämligen att de som idag gynnas av historiska orättvisor har moraliska skyldigheter att gottgöra för dessa orättvisor. Boken är en omarbetad och utökad version av en av Collstes tidigare monografier – *Global Rectificatory Justice* (2015). Till skillnad från sin engelska föregångare tar denna bok även upp exempel på historiska orättvisor där Sverige spelat en avgörande roll, såsom Sveriges roll i den transatlantiska slavhandeln och den svenska rasbiologin med dess steriliseringsprogram. I den utökade versionen