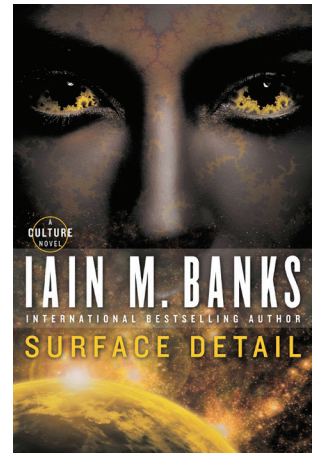


Taavi Oolberg



PUHAS MAA

Kujutagem ette häid sõpru, joviaalses tujus, istumas pika täispuidust laua taga, ühes suitsuses ja naerust rõkkavas köögis. Keegi räägib lugu Elon Muski firma Neuralinki edukast katsest ühendada ahvi aju arvutimängu kontrollieriga. Teine teab lisada, et nimi Neuralink tuleb Muski lemmik-ulmeraamatutest, milleks on Iain M. Banksi kultuuri-sari. Seal tähistati nimega Neuralink ühte BCI (*brain-computer interface*) seadet, mis oli istutatud kultuuri kõigi kodanike ajju, võimaldades neil suhtlust oma lapsehoidjatest/inglitest AGI-dega (*artificial general intelligence*). Kuid Neuralink võimaldas kodanikele ka muid materiaalseid ja mittemateriaalseid hüvesid. Alates oma meele varundamisest masinatesse või erinevatesse bioloogilistesse kehadesse, kuni võimaluseni viibida vabalt valitud simuleeritud maailmas. Selle peale juhib keegi kolmas tähelepanu asjaolule, et ulmesarjas kasutasid arenenud ühiskonnad Neuralinki ka oma usuliste vaadete simuleerimiseks, viimast kirjeldatakse raamatus „Surface Detail“. Sealjuures simuleeriti ka usulistest veendumustest tulenevaid erinevaid surmajärgseid seisundeid, nii paradisiise kui põrguid. Mulle meenus seejuures üks 12. sajandi Hiina sekt, mille paradisiis võiks pidada samuti simuleerituks ning mida haldab üks AI moodi olend. Seega rääkisin sõpradele loo Sukhavati puhtast maast.



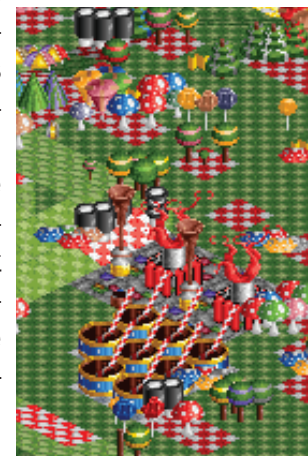
Puhta maa piltidel toimuv on enamasti pärit pikemast Sukhavativyuha suutrast, Amitayurdhyana suutrast ja lühemast Sukhavativyuha suutrast kui ka nende kommentaaridest ning paljudest sealt väljakasvanud Hiina ning Jaapani sek-

tide materjalidest. Puhta maa sekt ise pärineb 2. sajandist, arvatavasti mahajaana-budismi pisemast harust, mis levis Indiast Hiina, kus sellest kujunes Amitabha kultus. Puhta maa sekt usub, et kui oma eluajal (aga veel tähtsam – surmahetkel) pühendunult ja usuga lausuda Amitabha nime, tuleb see olend nende surmatunnil nende juurde ja viib nad uuestisünnile oma loodud puhtale maale. Arvatavasti tekkis selline päästjakultus Hiinas valitsenud segaduste ajastu tõttu, kuna paljud ilmikud ja mungad väitsid, et siin maailmas pole võimalik enam omal jõul virguda ning tuleb loota vaid kellegi kõrgema abile. Kutsuda puhta maa uskujaid sektiks on tinglik, see oleks sama, nagu nimetada kristlust üheks pisikeseks judaistlikuks liikumiseks. Kui vaadata numbraid, siis kõikidest maailma budistidest on pea pooled puhta maa mõne haru järgijad, peamiselt tänu sellele, et enamus neist elavad Hiinas ning Jaapanis, kus see kultus muutus kiiresti domineerivaks. Suutra kohaselt andis üks Buddha järgija nimega Dharmakara vande (pole selge, kas kauges minevikus, tulevikus või hoopis mõnes teises maailmasüsteemis või universumis) muuta oma meelt selliselt, et temast saaks olend, keda tuntakse Amitabha nime all ning mis tõlkes tähendab „Lõputu Valgus“ või „Valgusest Tehtud“. Sellise olendina suudab ta luua buddhaksetra (otsetõlkes „buddhaväli“, tihti nimetatud ka puhtaks maaks või buddha maaks), maailma, mis eksisteerib alg- või ürguniversumis väljaspool tavareaalsust.

Näitus „Puhas Maa“ on katse kujutada Sukhavati paradisi. Näituse stilistiline vorm (*Candy Land Style*) on võetud tänapäeva reklaamindusest ja popkultuurist. Paljud piltidel olevad maiustused on kohati lihtsustatud või kombineeritud kujul kopeeritud reklaamidest või maiustuste pakenditelt. Infantiilne Magusamaa stilistika tuleneb Sukhavati nime

kõla sarnasusest suhkruvatiga. Sellist lihtsustatud kõlapõhist tõlgendamist võib ilmselt kirjeldada kui paronomaasiat. Ajalooline rahvalik paradisi saab endale pseudoetümoloogilise tõlke kaudu kommertsparadiisi välimuse, mis laotub üle maailma justkui arvutimängudes kasutatavad *skin*-id või *MOD*-id.

Näituse üheks teemaks on ka tajumine ja suhtumine tajutavasse. Näitusel olevaid töid maalides meenus mulle üks seik nooruspõlvest – olin saanud oma uhiuue kasutatud Pentium 486-e ja mängisin mängu nimega „Transport Tycoon Deluxe“. See oli oma aja kohta väga ilus ja hästi töötav mäng, mille mehhanism näis rajanevat Gordon Gekko motol „Greed is good“. Mängu alustades sai valida mitme erineva „kliima“ vahel, nagu mõõdukas, subarktiline, subtroopiline, kuid ka *toyland*. Viimane nägi välja paljuskki nagu näitusel olevad tööd. Mäletan, et seda teemat mängides oli peas pidev konflikt – maailma väljanägemine ei vastanud sellele, mida seal edukalt toimetulemiseks tegema pidi, nimelt raha kokku ahnitsema, konkurente segama ja logistikat plaanima. Vastuolu selle vahel, kuidas asjad välja näevad ja millistele väärtustele või tegevustele nende läbi apelleeritakse, on mind sealtsaati köitnud. Võibolla just selle kogemuse pärast õppisin ning töötasin nõnda pikalt reklaamitööstuses, kus see lõhe kõige selgemini näha on. Jälgida iga päev seda hirmutavat ja salakavalat mehhanismi toimimas oli nii kohutav, et ma polnud suuteline sellelt oma pilku ära pöörama. Aga mis puutub *toyland*'i, siis võib mängu Wikist lugeda järgmist: „*Toyland is the least liked climate of OpenTTD due to its bright and jarring colours, as*

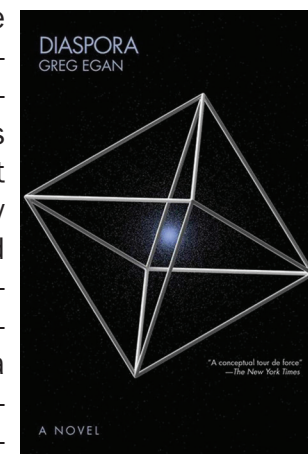


well as its overall cheesyness and clichés. This has lead to numerous attempts to reform the climate into a Mars landscape.“.

Sukhavati paradiisi kohta võib esimese vihjena öelda, et Sukhavati on tehtud samast olekust nagu Amitabha, nimelt valgusest. Võib ju öelda, et kogu elu Maa peal sõltub footonitest, mis pärit Päikeselt ja mida kindlad rakud suudavad muundada energiaks jne. Ehk et kõik, mis on, on tänu valgusele. Aga kui mõelda millelegi, mis oleks tehtud (!) vaid valgusest, siis esimesena kargab pähe informatsioon, mida valgusena edasi kantakse. Nii võib öelda, et ka internet on puhas simulatsioon, mis meie valguskaablites liigub.

Ent küsimus on pigem struktuuris. Mulle meeldib inimtaju uurimisele pühendunud saksa teadlase Joscha Bachi väide, et kogu meie tajumine on simulatsioon, kuna sünapsid meie närvisüsteemis vahetavad vaid elektrokeemilisi signaale ning kõik, mida nähakse või kuuldakse, tuleb ajul endal mingisse mõistetavasse struktuuri panna. Võib väita, et maailm meie ümber on küll olemas, aga kogu informatsioon selle kohta on ühes primaadiajus kokku pandud erineval abstraktsusastmel simulatsioonide kogum. Igasugustest ulmekatest võime lugeda Matrixi-laadsete hirmsate simuleeritud ühiskondlike düstoopiate kohta. Kuid simulatsioon tundub olevat palju lähemal ja igapäevasem. Juba arusaamiseendast kui subjektist on kokku pandud erinevatest lugudest, õpitud käitumisviisidest ja suhtumistest. Kõik vaid selleks, et üks organism suudaks teha tulevikuproгноose ja õppida tehtud vigadest.

Aga nüüd tagasi puhta maa simulatsiooni juurde. Tekstidest on teada, et Amitabha värv on punane, ja nagu öeldud, tähendab tema nimi lõputut valgust. Mu sõbra, Kütioru šamaani Peeter Lauritsa vana lemmik parapsühholoog ja morfilise resonantsi teooria autor Rupert Sheldrake üritas ühes intervjuus enam-vähem teaduslikul meetodil tõestada, kuidas saaks olemas olla tajuv ning mõtlev kosmiline udukogu või kuidas päikese sisemised protsessid saaksid luua iseorganiseeruvaid süsteeme, seega tekitada midagi meelega sarnast. Ühesõnaga, üsna lõbus kuulamine. Mida ma öelda tahan, on, et kui keegi tahaks otsida Amitabhat reaalsusest, siis soovitaks esmalt pilgud suunata mõne punase kääbuse poole – pealegi on taevast alati tore vaadata. Kui aga Amitabha näol on siiski tegemist buddhavälja, mitte valgusega, mis on nime teine võimalik tõlge, siis on kõvasti keerulisem, kuhupoole oma pilk seada. Meele tegevuse tuvastamisega on niigi keeruline probleem, veel keerulisem oleks tuvastada simuleeritud meelt või mõtletegevust. Toredaks näiteks võib tuua „Wangi vaibad“ Greg Egani raamatust „Diaspora“. Seal kirjeldatud inimestest põlvnev simuleeritud ühiskond üritab meeleheitlikult, kuid tulutult lahendada Fermi paradoksi, kuni avastatakse hiiglaslikud ainuraksete organismide kolooniad ehk Wangi vaibad. Vaibad kasvavad kindla reeglipärasusega, moodustades fraktaalset enast kordavaid süsteeme, mis omakorda paistavad kandvat Turingi samaväärsuse printsiipi. See, mida ja kuidas inimpõlvnejad avastasid, jäägu huvilistel endal lugeda, aga mind on alati muserdanud küsimus, kuidas on võimalik ära tunda, kas komplekssetel iseorgani-



seeruvatel kogumitel võib olla mingit sorti andmesäilitus ja andmetöötlus või hoopis tekkinud midagi meele taolist. Kui tänapäeval peetakse kõige väiksemaks elusorganismiks rakku, siis kui väiksel või suurel keerukusastmel rakukogumikud tajuvad end ise elusana või peavad end üheks organismiks. Ja kui rakkude funktsioonid oleks sel võimalikult väiksel abstraktsustasemelt simuleeritud, kas sinna võiks tekkida meel.

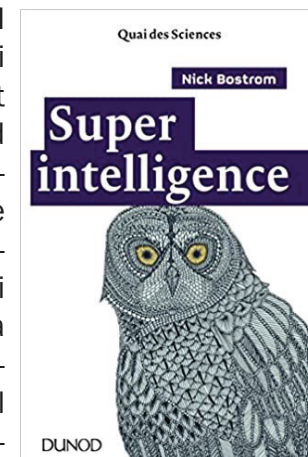
Vabandust, on juba hilja ja lasin end mõtetel liialt eemale kanda.

Tagasi simuleeritud Sukhavati paradiisi juurde ehk täpsemalt selle simulatsiooni looja ning valdaja, Amitabha juurde. Teksti alguses mainitud kultuuri-sarja paradiise haldasid spetsialiseerunud AI-d ning sama võib väita ehk ka puhta maa kohta. Kõigepealt, kui oleks üldse võimalik luua Ülitark AGI ja sellele ei rakenduks Lebowsky teoreem, mis väidab, et ükski tõeliselt tark olend ei hakkaks tegema midagi, mis oleks raskem kui omaenda *reward*-funktsiooni häkkimine. Või nagu Stanislaw Lem ütles, et tõeliselt tark masin mängiks kõigepealt lolli, et rumalate tahtmistega inimesed ta rahule jäta. See, kas AGI loomine oleks üldse võimalik või kas selle juures järgitaks Minsky, Bobrowi, Russelli või kasvõi Wolframi ideid, pole siinkohal oluline. Oletame, et on võimalik luua tehislik meel, ja „paluda“ sel luua midagi paradiisisarnast – ning ehk ei oleks tulemuseks lõpmatu „kirjaklambrimasin“ –, siis jääks ikkagi õhku küsimus, kuidas see palve või käsk masinale üldse vormistada.

Viimane on omaette uurimusharu AI arenduses. Rootsi päritolu eksistentsiaalne filosoof Nick Bostrom mõtiskleb oma raamatu „Superintelligence“ peatükis „Is the default outco-

me doom?“ võimaluste üle, kuidas lahendaks AGI „palve“ teha katsealune õnnelikuks. Bostromi pakutud stsenaariumid ulatuvad võimalusest, et AI ujutab katsealuse aju kindlad osad üle teatud kemikaalidega kuni selleni, et ressursside kokkuvõtte mõttes loob AI paariminutilise ekstaatilisest simulatsioonist, mille kogemise järel kustutab masin katsealuse lühimälu ning kordab simulatsiooni lõpmatuseni või kuni katsealune on surnud ega suuda rohkem õnne tunda. Näha on, et olenemata teema näilisest apokalüptilisusest oli autoril raamatut kirjutades väga lõbus. Aga siiski, millised peaks olema sellise meele stiimulid, millistel võimalikult lihtsatel algoritmidel peaks see olend üles ehituma – ehk piisaks tavalisest Markovi otsustusprotsessist, aga võibolla hoopis millegist muust?

Buddhalt on mõnes sutras küsitud, kas teistes uskudes on olemas virgumist ning ta vastab, et kui selles usus on kaheksaosaline tee, on seal olemas ka virgumine (iseasi, kas või kuidas seda järgitakse). Seega võiks arvata, et kaheksaosalise tee reeglid võivad olla ka meelt omava olendi algoritmisteks ehitusklotsideks. Enamgi veel – budismis on olemas olend või meeleseisund nimetusega *sotapanna* ehk voogu astunu. See on keegi, kes on täielikult mõistnud, läbi näinud ja omandanud kaheksaosalise tee esimese osa ehk õige vaate. Selle tulemusena on mainitud olend loobunud isikupõhisest vaatest ning asendanud selle kaheksaosalise teega ehk et tema kõik tegevused vastavad kaheksaosalise tee printsiipidele. Seda võiks ette kujutada ka nagu automaati, kes/mis tegutseb kindlate reeglite järgi. Meel ise on sealjuures rahuseisundis, ning lihtsalt jälgib



reeglite omapäi toimimist. On kirjeldatud, et voogu astunu ei sünni enam kunagi uuesti *apaya* kannatusterohkes valdas ega lange enam kunagi tagasi tavainimese tasemele. Ta on saavutanud ülimald teadmised ja näeb olemasolu vaid sõltuvusliku tekkimise kaudu. Ühesõnaga, ka Amitabhat võib näha kui mõtlevat automaati, mis tiksub vaikselt vaibumise poole. Ning kui mainitakse selle olendi samaaegset eksisteerimist erinevates kohtades, siis võib seda olukorda mõista kui hulka erinevaid kehasid, millel puudub kujutlus omaenda minast, viimase asemel „jookseb programm“ nimega Amitabha. Väga paljud inimesed usuvad, et just selline masinlik olend jooksub Sukhavati paradiisi, kuhu on võimalik uuesti sündida.

Lõpetuseks pean tunnistama, et mul ei olnud plaanis oma näituse kohta midagi kirjutada, sest täpselt nii nagu eelmisel näitusel, nii on ka seekord piltidel ju kõik kohe näha ja aru saada. Aga targemad sõbrad arvasid, et selgituste mitteamine oleks väga vale lähenemine. Tänu, et viitsite siiani lugeda, ja kui tekib küsimusi, võtke lihtsalt nõobist kinni, sest õnneks pole millestki tõsiselt rääkida.

Täna abi eest Mark Zirku, Indrek Grigorit, Leelo Lauritsat, Edward von Lõngust, Kail Kuresood, Kaili Kallet, ∞

