

Teadlased hoiatavad: tehisintellekt muudab meid märkamatuks robotite sarnaseks

37 min tagasi - 02.09.2026 Autor: [AM](#)

Rahvusvaheline teadlaste rühm hoiatab [värskes uuringus](#) ootamatu ja kõhedust tekitava nähtuse eest: mida inimlikumaks muutub tehisaru, seda rohkem hakkame meie alateadlikult peegeldama masinaid, lihvides maha oma inimlikud veidrused.

Teadusajakirjas AI & Society avaldatud teoreetilises uurimistöös võtavad teadlased luubi alla viisi, kuidas igapäevane suhtlus nutikate klienditoe agentide ja keelemudelitega meie minapilti murendab. Autorid võtsid selle nähtuse kirjeldamiseks kasutusele uue mõiste – „robotoidne inimlikkus“ (*robotoid humanness*).

Peegel, mis näitab meid masina koodina

Oleme harjunud mõtlema, et tehisaru on pelgalt kuulekas tööriist või digitaalne abiline, keda me oma suva järgi suuname. Kuid pilt muutub, kui vastaspool hakkab meenutama päris inimest.

Birminghami ülikooli turundusteadlane Inci Toral-Manson kirjeldab seda kahepoolset protsessi tabavalt:

„Klienditeeninduse sotsiaalsed robotid on disainitud jäljendama žeste, kõnet ja emotsionaalseid vihjeid, et kutsuda klientides esile kognitiivseid ning emotsionaalseid reaktsioone. Meie uuring vaatleb, kuidas suhtlemine selliste inimnäoliste robotitega tekitab kahesuunalise mõju: robotid muutuvad üha enam inimeste sarnaseks ning inimesed omakorda robotite sarnaseks. Seda me nimetamegi „robotoidseks inimlikkuseks“.“

Masinõpe ja suured keelemudelid jälgivad igat sinu sisestatud sõna, kohandavad oma vastuseid ja pakuvad vastu statistiliselt ideaalset, siledat servadega peegeldust sinust enesest.

Rootsis asuva Linnaeuse ülikooli juhtimisinsener Selcen Ozturkcan selgitab mehhanismi:

„Tarbija tegutseb, robot vastab ning korduva kogemuse kaudu võtab tarbija selle suhtluse aja jooksul omaks. Masinõpe ja tehisaru võivad seda protsessi võimendada, kohandades roboti käitumist kasutaja sisendi põhjal ning võimaldades robotilt veelgi personaalsemat ja inimlikumat matkimist.“

Kuidas me vabatahtlikult oma nurgad maha lihvime?

Päriselus on inimsuhtlus ettearvamatu, täis pause, emotsioone, konarusi ja huumorit. Kui räägid elava sõbra või teenindajaga, ei tea sa kunagi sajaprotsendiliselt, millise reaktsiooni vastu saad. Tehisaru on aga treenitud konvergensile – see otsib mustreid, optimeerib ja premeerib selgust.

Selleks, et masin meid võimalikult kiiresti ja tõrgeteta mõistaks, hakkamegi me end tema jaoks „söödavaks“ vormima. Me loobume keelelistest viguritest, sarkasmist ja mitmetähenduslikkusest, sest süsteem lihtsalt ei oska nendega midagi peale hakata.

Nagu nendib Selcen Ozturkcan:

„Inimestele omane kaasasündinud peegeldusinstinkt võib viia selleni, et inimesed hakkavad vastu peegeldama roboti suhtlusmustreid, muutudes seeläbi ise robotile sarnasemaks.“

Mugavuse varjatud hind: kas me kaotame iseenda näo?

Uuring toob välja kolmeastmelise peegeldusmudeli.

Esmalt astume tehisklikku sotsiaalsesse reaalsusesse, kus vestlus tundub ehtne. Seejärel püüab algoritm kinni meie digitaalse identiteedi ja sööda selle meile tagasi juba lihtsustatud statistilise profiilina. Lõpuks kohandamegi me oma eneseväljendust selle järgi, mida süsteem suudab hõlpsasti lugeda ja mille eest premeerida ("masin saab aru").

Taanis asuva Aarhushi ülikooli äriarenduse teadlane Jean-Paul de Cros Peronard rõhutab [Science Alerti vahendusel](#), et teema vajab praegu väga kiiret tähelepanu:

„Robotid ja tehisaru on klienditeeninduses saanud igapäevaseks nähtuseks. Seetõttu on äärmiselt oluline mõista, kuidas inimesed nendega suhestuvad, et ettevõtted saaksid olemasolevat tehnoloogiat

kasutada parimal võimalikul moel, jäädes samal ajal eetiliseks.“

Kuigi teadlaste töö on esialgu teoreetiline ja vajab põhjalikke empiirilisi inimkatseid, paneb see meid tõsise küsimuse ette. Juturobotid ja nutiassistendid säästavad meie aega, kuid kas me märkame piiri, kus nutikas abiline hakkab dikteerima seda, kuidas me inimestena mõtleme ja suhtleme?

- [Uudised](#)
- [Robotid](#)
- [Tehisintellekt](#)

Pilt

