

Magyar kutatók AI segítségével írták le, hogyan lát az agyunk

Nemcsak az agyműködés modellezésében jelent előrelépést, hanem a gépi látási rendszereket is megbízhatóbbá és pontosabbá teheti az az újfajta, AI-alapú látórendszer-modell, amelyet a HUN-REN Wigner Fizikai Kutatóközpont kutatói írtak le. Eredményeiket a [Nature Communications](#) folyóiratban tették közzé.

Neurális kód: az idegsejtek rövid elektromos impulzusok segítségével kommunikálnak egymással és az izmainkkal. Az idegsejtek közötti kommunikáció nyelve a neurális kód: ez biztosít információt a környezetben zajló folyamatokról, és arról, miképpen reagálunk ezekre.

Mély diszkriminatív modellek: ezek az AI-eszközök a mélytanuló rendszerek közé tartoznak, melyek külső tanítóingerek hatására tudják minél hatékonyabban megkülönböztetni a különböző kategóriájú inputokat (például képeket) és hatékonyan ismerik fel a hasonlóságot az azonos kategóriájú inputok között (pl. telefonos arcfelismerés).

Mély generatív modellek: ezek az AI-rendszerek abban különböznek a diszkriminatív modellektől, hogy nem igényelnek külső tanítóingereket a tanuláshoz, ehelyett ezek az AI-eszközök magukat tanítják. A nagy nyelvi modellek és képgeneráló modellek (ChatGPT, Dall-E, Midjourney) egyaránt a generatív modellek közé tartoznak.

Agyunk kétirányú kapcsolatokkal sűrűn összekötött területek hálózata, ahol az ellentétes irányú kapcsolatok jellege és szerepe még messze nem tisztázott. Amikor valamit meglátunk, az agyunk több szinten dolgozza fel az információt: az egyszerű formáktól a bonyolultabb fogalmakig. Az AI eddigi képfelismerő rendszerei, amelyek például felismernek egy kutyát a telefonunk fotóján, egyirányú feldolgozással működnek: az információ csak „alulról felfelé” halad.

Az agyunk viszont kétirányban dolgozik: nemcsak az alakítja az idegsejtek válaszát a feldolgozás adott szintjén, hogy a korábbi szintek mire jutottak, hanem az is, mi fog történni a következő feldolgozási szinten. Ez azt jelenti, hogy az agy mindig figyelembe veszi a környezetet és a kontextust is: nemcsak azt, hogy mit látunk, hanem azt is, mit jelent az, amit látunk (a meglátott kutya barát vagy ellenség, közelít vagy távolodik). Ennek a következménye pedig drámai: a neurális kódot nem csak az határozza meg, ami a feldolgozásban az adott feldolgozási szint előtt történt, hanem az is, ami a feldolgozás következő lépéseiben történik.

A HUN-REN Wigner Fizikai Kutatóközpont kutatói által kifejlesztett modell ezt a kétirányú információáramlást utánozza, azaz egy olyan AI-modellt hoztak létre, amely nemcsak lát, hanem az emberi agyhoz hasonlóan értelmez is. Ennek segítségével 'nemcsak az idegrendszer információfeldolgozási folyamatait tudjuk precízebben feltárni (köztük olyan érdekes jelenségeket is, mint például a látási illúziók), hanem megbízhatóbb és rugalmasabb gépi látási rendszereket is készíthetünk.

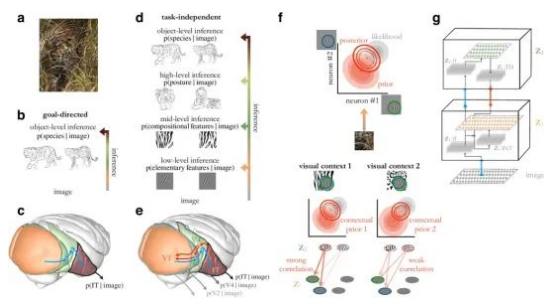
Csikor Ferenc és munkatársai munkájukban arra világítottak rá, hogy az idegrendszerünk összetettebb feladatot old meg, mint a telefonunkban rejlő képfelismerő algoritmus. Az idegrendszer ugyanis rugalmasan kíván megfelelni többféle kihívásnak, a megfigyelt állat típusának megállapításán túl arra is, hogy eldöntse, az állat barát vagy ellenség, felénk mozdul vagy tőlünk el.

Ahhoz, hogy rugalmasan tudjunk alkalmazkodni a különféle igényekhez, a hagyományos mély diszkriminatív modellek nem megfelelőek, helyettük a mély generatív modellekhez kell fordulni. A

HUN-REN Wigner FK kutatói szerint a jövőben ezek az új AI-modellek ellenállóbbak lehetnek hibákkal vagy támadásokkal szemben, kevesebb felcímkezett tanítóadattól is tanulhatnak, valamint sokkal pontosabb gépi látási rendszereket tehetnek lehetővé.

Sajtókapcsolat:

- Torda Júlia, kommunikációs vezető
- kommunikacio@hun-ren.hu



© HUN-REN Wigner Fizikai Kutatóközpont

Eredeti tartalom: HUN-REN Magyar Kutatási Hálózat

Továbbította: Helló Sajtó! Üzleti Sajtószolgálat

Ez a sajtóközlemény a következő linken érhető el:

<https://hellosajto.hu/?p=26998>