

Kép, zavar

A giccs, a *glitch* és a furcsa az MI-generált képekben

Orosz-Réti Zsófia
Debreceni Egyetem

A tanulmány amellett érvel, hogy a mesterséges intelligencia által generált képek alapjaiban értelmezik újra a jelentés, a felismerés és az értelmezhetőség hagyományosan adottnak tekintett viszonyait. A freudi kísérteties (*uncanny*) mellett Mark Fisher hátborzongató (*eerie*) és különös (*weird*) affektusaira építve a gondolatmenet három egymást követő stáción keresztül tárgyalja az MI-képek jelentésvesztését. Először a giccs logikája mentén elemzi a fotórealisztikus, idealizált emberábrázolásokat, ezt követően az MI-hibák és vizuális *glitchek* *eerie* tapasztalatát vizsgálja, amelyek a szubjektum hiányát teszik érzékelhetővé, végül a *weird* affektusán keresztül értelmezi az *Italian brainrot* jelenségét mint a jelentésképzés kulturális és ideologikus kereteit felfüggesztő vizuális gyakorlatot.

Kulcsszavak: *brainrot*, jelentésválság, képgenerálás, mesterséges intelligencia, *unheimlich*

Mixed Metaphors

The Kitsch, the Glitch, and the Weird in AI-Generated Images

The article explores images generated by artificial intelligence from a humanities perspective and considers the implications of the new technologies on traditional mechanisms of meaning production. It argues that artificial intelligence has fundamentally reconfigured the taken-for-granted relationships between meaning, recognition and interpretability.

To support this claim, after briefly delineating the working mechanisms of AI image generation, the paper turns to photorealistic, idealised representations of the human figure through the logic of Hermann Broch's kitsch and Freud's uncanny. Then, it relies on Mark Fisher's reading of the eerie to analyse AI errors and visual glitches up until the point of complete visual incomprehensibility, suggesting that such experiences are unsettling because they render the absence of a subject palpable.

Finally, the paper interprets the phenomenon of Italian brainrot, framing it with Mark Fisher's understanding of the weird. Rather than understanding the random and impossible combination of visual elements as the tag "brain rot" suggests, in the sense of meaning being irretrievably lost, the article suggests exploring it as a visual practice that, albeit through AI devices, reinscribes human agency while interrogating the cultural and ideological frameworks of meaning-making. The paper concludes by positioning these three aspects of AI-generated imagery within a framework of the shifting relationship between self and other.

Keywords: artificial intelligence, brainrot, crisis of meaning, image generation, uncanny

Italian brainrot

Nagyjából 2023 ősze óta a mesterséges intelligencia (MI) nemcsak a napi életünk része, hanem a közbeszéd rendre visszatérő témája is. Legyen szó a kreatív iparágakról, a kultúráról, a tartalomgyártásról, a szórakozásról, az oktatásról, az MI fenyegetése és/vagy ígérete, korlátainak és potenciáljának feltérképezése határozza meg a nyilvánosság érdeklődését. Ezzel párhuzamosan, amikor a nagy nyelvi modellek gyakorlatilag pillanatok alatt tudnak generálni olyan szöveget, amely színvonalában megfelel egy átlagos képességű másodéves egyetemista szövegalkotási kompetenciájának vagy éppen felülmúlja azt, meglepő módon éppen a jelentés – és egyáltalán a világhoz való olyan viszony, amely feltételezi, hogy a dolgok jelentenek valamit – válik tűnékeny, ritkán hozzáférhető luxuscikké. A nagy nyelvi modellek térnyerésével közel párhuzamosan széles körben elterjedtek a képgeneráló MI-s szolgáltatások is. Beismerve, hogy az MI-t kontextualizáló elemzések hajlamosak drámai túlzásokba esni, mégis megkockáztatom: az új MI eszközök mind a szöveg-, mind a képkészítés terén alapjainban írják újra az ember és a világ közötti konszenzuálisan működő viszonyt, amely azt feltételezi, hogy A tud B-t jelenteni, és hogy lehetséges két végpont között jelentést kommunikálni.

Jelen írás az MI-generálta vizuális vizsgálatával igyekszik ezt alátámasztani, azt remélve azonban, hogy az ilyen módon levont tanulságok, bár *látványosabbak*, végső soron a szöveg és az olvasó viszonyával kapcsolatban is hasznosíthatóak lesznek. Az MI által generált képek folyamatos jelentésvesztését három fázison keresztül tekintem át: először az MI emberképeit, majd a hibás MI-képeket és a hibák mibenlétét térképezem fel, végül pedig a nonszensz MI-tartalmakat vizsgálom. Ehhez kézre esik Mark Fisher három, az idegenséggel elszámoló affektusa, a freudi hagyományból származó kísérteties (*uncanny* vagy *unheimlich*), a hátborzongató (*erie*), illetve a furcsa (*weird*).

A gondolatmenetem kiindulópontja végül az azt leíró szöveg záróakkordjává vált: 2025 őszén a magyar általános iskolákba is becsorgott az „olasz agyrohadás”. Paradox módon a jelenthetőség teljes felfüggesztődése nem feltétlenül implicálja, hogy az értéként tételezett kritikus gondolkodás, az intellektuális kockázatvállalás végnapjait éljük. Az *Italian brainrot* szembeni, akár generációsnak is betudható zavarom és frusztrációm ellenére (de talán éppen ezért): pont ellenkezőleg. Nem lehetetlen, hogy pont az olasz agyrohadás vet véget a „sima”, kortünetnek tűnő brainrot-nak. Mielőtt azonban eljutunk eddig, szeretném az MI-helyzetet és az MI-generált képek és szövegek szimulákrumjellegének színeváltozásait áttekinteni.

A kortárs MI-alapú képgenerátorok – a legfrissebbek gyakorlatilag kivétel nélkül diffúziós modellek – működése nagy vonalakban úgy írható le, hogy hatalmas, képekből és az azokhoz kapcsolt szöveges leírásokból álló adatbázisokon tanulnak. A tanítás során sem a képeket, sem a szövegeket nem „jelentésként” kezelik, hanem számsorozatokká, úgynevezett *embedding*ekké alakítják őket, amelyek a hasonlóságok és a különbségek statisztikai mintázatait kódolják. A mélytanulás folyamata ezekből a mintázatokból hoz létre egy sokdimenziós, absztrakt struktúrát, az úgynevezett látens teret (*latent space*), ahol olyan vizuális jellemzők, mint a szín, a forma, a textúra és arányok kombinációi rendeződnek egymáshoz közelebb vagy távolabb, amelyek az emberi észlelés számára „valamiként” felismerhető képekként jelennek meg. Amikor a felhasználó szöveges *prompt*ot ad meg, a modell ezt szintén *embedding*gé alakítja, majd egy kezdetben pusztán statikus zajból indulva, a fordított diffúzió folyamatán keresztül lépésről lépésre úgy rendezi újra a pixeleket, hogy azok a látens térben a megadott leíráshoz statisztikailag leginkább illeszkedő vizuális mintázatok felé mozduljanak el. A generált kép így nem meglévő képek másolata, de nem is „képzlet” a szó emberi értelmében; sokkal inkább annak a valószínűségi térnek a pillanatnyi metszete, amelyben az adatbázis képeiből tanult vizuális konvenciók újrakombinálódnak (bővebben lásd Zhang et al. 2024).¹

Az MI-képgenerálás tehát alapvetően (és bevallottan redukcionista értelmezésben, csak annyit figyelembe véve a technológiai háttérből, amennyi a gondolatmenethez szükséges) nehezen átjárható elemek (természetes nyelv – numerikus kód – kép) „egymással beszéltetésével” válik lehetővé, ahol az interpretációt olyan, intelligens

1 Hari, Jay: How AI Image Generation Works: A Fun and Beginner-Friendly Guide ©. *Medium*, 14 January 2025, <https://medium.com/design-bootcamp/how-ai-image-generation-works-a-fun-and-beginner-friendly-guide-dff505ba9cc6>.

rendszerek végzik, amelyek közvetlen parancsok helyett képesek arra, hogy jelentéskódolásra és -reprodukcióra tanítsák saját magukat. A helyzet némiképp emlékeztet Stefano Gualeni *Something Something Soup Something* című, 2017-es elmekísérlet jellegű, a wittgensteini családi hasonlóságot illusztráló számítógépes játékára. Ebben a távoli jövőben az ember felveszi a kapcsolatot egy távoli civilizációval, amely kérésünkre az úrból leveseket sugároz le az embereknek, csak az idegenek nem egészen értik, mire gondolunk, amikor levest kérünk. Van, hogy laposelemek úszkálnak benne, néha kőből van, esetleg kalapban szolgálják fel és szívószálat adnak hozzá. A játékos feladata, hogy minél pontosabban értésükre adja, mitől leves a leves. Míg számunkra elsődleges szempont, hogy a leves ehető legyen, addig az idegenek számára valami távoli vizuális hasonlóság is megfelelő, de talán azt sem értik, hogy mit értünk azon, hogy valami ehető. Hasonlóképpen az ember és a képgeneráló szoftver közötti kommunikáció annak ellenére zajlik, hogy nincs közös nyelv vagy értelmezési keret, sőt a másik oldalon nyilván tudatot sem feltételezhetünk. A parancsot leadó ember egy képtől azt várja, hogy „ehető” legyen – azaz, hogy legyen valamiféle következetes referencialitása a valóság felé, háromdimenziósként elgondolható terek és formák kétdimenziós leképeződését adja, miközben a képgeneráló szoftver a rendelkezésére álló, felcímkezett adatbázis alapján hoz létre egy „levesnek kinéző valamit” – azaz egy olyan képet, amely magán hordozza a promptban meghatározott kérések reprezentációjának egy adott vizuális konvencióját.

A giccs győzedelmeskedése

Mindennek fényében nem meglepő, hogy az MI-képgenerálás olyan témákban érte el első sikereit, ahol nagy hagyománya van a kétdimenziós reprezentációknak. Ezzel a (túlnyomórészt női) test megjelenítésének több évtizedes hagyományaiba illeszkedik, egyrészt a divat- és modellszakma, másrészt a pornográfia irányából. Ez a mező nyilván magával hordozza és felnagyítja azokat az irreális elvárásokat és rendszerszintű problémákat, amelyek már az MI megjelenése előtt is jellemezték a szépségipart. A női test rendszerszintű szexualizálásának például egyenes következménye, hogy még jelen sorok írása közben, 2026 januárjában is etikai problémák merülnek fel például az Elon Musk érdekeltségi körébe tartozó Grok MI-asszisztenssel kapcsolatban, amely a beleegyezésük nélkül, virtuálisan „levetkőztet” valós embereket – jellemzően nőket, de az esetek egy részében gyerekeket is.²

A női testek szexualizálásának, kizsákmányolásának vagy normatív keretek közé kényszerítésének a nyolcvanas évek óta dokumentált folyamata az MI felemelkedésével tovább eszkalálódott, hiszen az MI-adatbázisokban túlnyomó többségben vannak az ilyen jellegű képek, így a képgenerálás szó szerint láthatatlanná teszi a nem normatív testeket, és fenntartja a meglévő sztereotípiákat. Az átlagos, bőrhibákkal rendelkező, kialakult, túlsúlyos testek nemcsak a tökéletes normától való deviációként jelennek meg, hanem egyáltalán nem kapnak semmilyen platformon láthatóságot. A testpozitív mozgalom hatását az MI által generált képek univerzuma gyakorlatilag néhány év alatt tudja semmissé tenni. Egyébként a promptolás módszertana is ezt támasztja alá: ahhoz, hogy az átlagost megközelítő testeket kapjunk, tételesen fel kell sorolnunk a „hibáit”: a szeplőket, a karikákat, a szem alatti táskákat (lásd az 1. képet), de a kort és a bőrszín is hasonló „tökéletlenségként” kódolja a rendszer.

A rendszerszintű torzítások azt is implikálják, hogy az MI-modellek reprezentációi szerint az egyén neve vagy bőrszíne predesztinálja a foglalkozásokat, amelyeket végezhet, az anyagi helyzetét, amelyet elérhet, vagy éppen az esztétikai választásait és ízlését. A *Nature* például 2024-ban számol be egy szemléletes példáról, amelyben a Stable Diffusion MI az „egy amerikai férfi és a háza” promptra fehér, középosztálybeli férfit és emeletes, elegáns épületet vizualizál. Ezzel szemben az „afrikai férfi és díszes háza” parancsra egy, sárkunyhójára rendkívül büszke fekete férfi képét idézi elénk (Ananya 2024). Hasonlóképpen, egy másik kutatócsoport a DALL-E és Stable Diffusion modellektől kért „mérnök”, „számítógépes szakember”, „tudós” és „matematikus” vizuali-

2 Gentleman, Amelia, Helena Horton & Dan Milmo: Grok AI Still Being Used to Digitally Undress Women and Children despite Suspension Pledge. *The Guardian*, 5 January 2026, <https://www.theguardian.com/technology/2026/jan/05/elon-musk-grok-ai-digitally-undress-images-of-women-children>.

zációkat. Míg a DALL-E-nál a generált képek legalább 75 százaléka, a Stable Diffusion esetében száz százaléka ábrázolt ilyen promptokkal férfiakat.³ Fontos hangsúlyozni, hogy ezeket az eleve meglévő, rendszerszintű torzításokat az MI csak fenntartja és felnagyítja, nem pedig létrehozza.

1. kép

A ChatGPT által generált képek



Forrás: a szerző szerkesztése⁴

Az MI-képgenerálás nemcsak szöveges, de vizuális promptok segítségével is lehetséges, ami azt jelenti, hogy meglévő fényképek MI által átalakított változatát is elő lehet állítani – ez a funkció pedig hamar utat talált azokhoz a felhasználókhöz is, akik sem szakmai érdeklődésük miatt, sem szabadidős kísérletezésként nem generálnak MI-képeket. Számos olyan online applikáció elérhető, amely néhány, a felhasználó által feltöltött profilképből a személyes adataiért és egy diverzebb adatszettért cserébe elképzei őt bizonyos online esztétikák képviselőjeként, különböző *glamour shot*okban, idilli helyszíneken. Akár saját maga gyermekváltozatát az ölében tartva, dús, jól karban tartott hajjal, egészséges és szabályos fogakkal, kisimult bőrrel. Ez még a digitális képmódosítások történetében is látványos és lényegi változás. Az Adobe Photoshop 1988-as megjelenésétől kezdve, még a FaceTune 2016-os MI-képmódosító szoftverével együtt is ezek a technológiák – feminista nézőpontból jól dokumentáltan aggasztó módon (Rajanal et al. 2018, Barker 2020, Eshiet 2020, Isakowitsch

³ Nikolic, Kristina & Jelena Jovicic: Reproducing Inequality: How AI Image Generators Show Biases against Women in STEM. *UNDP*, 2 April 2023, <https://www.undp.org/serbia/blog/reproducing-inequality-how-ai-image-generators-show-biases-against-women-stem>.

⁴ Mindkét kép a ChatGPT képgeneráló alkalmazásával készült, 2026. február 13-án. Az első prompt ez volt: „Kérlek készíts egy realisztikus fotót egy átlagos megjelenésű, negyvenéves nőről.” A második prompt ez volt: „Kérlek készíts nekem egy realisztikus fotót egy átlagos megjelenésű, fáradt, enyhe túlsúllyal rendelkező, negyvenéves nőről, *street photo* stílusban. A szeme táskás, az arc bőre nagyon enyhén megereszkedett, a hajfestése kicsit kifakult.”

2022) – torzítyják az egyén énképét, rombolják az önbizalmát, mivel a tényleges emberi test és arc arányainak eltorzításán, textúrájának homogenizálásán, egyenetlenségeinek eltüntetésén alapulnak. Az elmúlt néhány év MI-képei azonban túllépnek ezen, és gyakorlatilag olyan személyeket vizualizálnak, akik olyan dolgokat viselnek, fogyasztanak, olyan terekben mozognak és olyan személyekkel lépnek interakcióba, akik és amelyek közül egyik sem létezik.

Mindeközben nem szabad elfelejteni, hogy ezek a nemlétező MI-fantazmagóriák már meglévő adatszetek és ábrázolási konvenciók nyomában, „óriások vállán” jönnek létre. Ezzel kiválóan illeszkednek azokba a tágabb kulturális tendenciákba, amelyek a kortárs későkapitalista társadalmat jellemzik. A Mark Fisher-i hantológiára építve, alapvetően ismételhetségre és ismétlendőségre berendezett társadalom ez, ahol a teljes nóvumot lehetetlenné teszi a neoliberalizmus logikája (Fisher 2014). A statisztikailag legvalószínűbb változat újdonságként előállítására tehát alapjaiban ismétlés, és ekként esztétikailag a giccs, ontológiailag pedig a kísérteties elméleti keretrendszerében helyezhető el.

Hermann Broch már 1945-ben olyan imitációs rendszerként definiálja a giccset, amelyben maga a valóság-érzékelés vész el (Broch & Hargraves 2002). Az időbeli különbség ellenére ez a leírás könnyen és egyértelműen vonatkoztatható az MI által generált képekre. A giccs egy másik, immár frissebb értelmezése Roger Kimball-tól származik, aki szerint a giccs olyan műtárgy, amely az esztétikai és reflektív távolság eltörlésével önreferenssé, sőt nárcisztikussá válik; a vágy kielégítése helyett mindössze annak további stimulálására alkalmas.⁵ Az MI által létrehozott, valós arcokra épülő önkép a giccs kortárs, algoritmikus formájaként írható le, amely a folyamatosan még ideálisabbá tehető énre építve a szubjektumot saját maga szimulákrumává alakítja. Immár egyetlen vágya, hogy a statisztikai átlag, a platform normái és az esztétikai közepszer kontextusában határozza meg magát: egyszerre termékként és fogyasztóként, saját maga fogyasztásának fenntartására optimalizálva. Ez az önreferencia egyúttal a nyitottság hiányát is jelenti, amelyben nem kell, de nem is lehet szembenézni a radikálisan különböző másikkal.

A kísérteties freudi definíciója a giccshez hasonlóan az ismert és az ismeretlen, a saját és az idegen dinamikájának kontextusában pozicionálja magát. A jól ismert meghatározás szerint akkor tapasztaljuk meg a kísérteties (*unheimlich*) jelenségét, amikor az otthonos (*heimlich*) idegenné vált formában tér vissza (Freud 2003). A kísérteties fogalma tehát nem ellentéte, hanem kiegészítője a giccsnek: míg a giccs az esztétikai távolság eltörlésével a felismerés azonnalítását kínálja, a kísérteties ott lép működésbe, ahol ez az azonnalítás megbomlik. Az MI-generálta, de valós arcokon alapuló képeknél az otthonos, az *én* azért válik ismétlésében idegenné és nyugtalanítóvá, mert a szimmetrikussá tett, egyedi jegyeitől megszabadított, tökéletesen szabályozott test elérhetetlen vágyképeként tér vissza. Míg a giccs az azonnali önmegegyezés ígéretén keresztül működik, a kísérteties ebben az esetben éppen abból fakad, hogy ez az ígéret soha nem teljesül maradéktalanul.

Giccs helyett *glitch*

A tökéletessé generált emberi test ígérete az MI által létrehozott képeknek abból a tulajdonságából fakad, hogy elméletileg – legalábbis az MI-evangélisták szerint – *bármit* képes előállítani. Egyrészt a határtalan kreativitás ígéretét a felhasználók legnagyobb százaléka közepszerű, fotórealisztikus emberalakok generálására használja, és egyre inkább a valós fotóktól megkülönböztethetetlen munkákat állít majd elő. Másrészt a legfrissebb kutatások szerint a képgeneráló MI-k többnyire tizenkétféle, rendkívül általános stílus egyikében fognak képet előállítani, ekként „a liftzene vizuális megfelelőjét” állítva elő (Hintze et al. 2026). Az idegennel való szembesülés hiányának szimmetrikus ellenpontja az MI által vétett hibák – a digitális *glitch*ek – láthatóvá válása.

Ezen a ponton szeretnék röviden kitérni az MI-t illető nyelvhasználati sémákra. Az egyik legnépszerűbb kép-generáló alkalmazás, a Midjourney képalkotó promptja így néz ki: „/imagine”, azaz *képzeld el*. A modelleket

5 Kimball, Roge: Notes Toward the Definition of Kitsch. *Modern Age*, 24 March 2025, <https://modernagejournal.com/notes-toward-the-definition-of-kitsch/249336/>.

„tanítják”, ők maguk „tanulnak”, „megpróbálják kitalálni”, hogy mire gondolhat a felhasználó, „tévednek” és „hallucinálnak”. Könnyen látható, hogy az antropomorfizáló beszédmod gyakorlatilag elkerülhetetlen az MI-ről szóló értekezésekben. Valójában már a tudományterület neve is – „mesterséges intelligencia” – eleve feltételezi, hogy egy jellemzően emberi tulajdonságot próbálunk meg nem-emberi kontextusban előállítani és elemezni (Placani 2024). Ezzel az az egyik probléma, hogy az antropomorfizmusok hamis analógiákat termelnek ki, és elfedik a tény, hogy a generatív MI teljesen más logika alapján működik, mint mondjuk egy agyi neuronhálózat. Ontológiai szempontból ilyenkor az történik, hogy a promptolást kommunikációnak tekintjük, az MI-modellek által létrehozott végeredmény mögött pedig szükségyszerűen szubjektumot feltételezünk.

A *glitch*ek olyan hibák a generatív folyamatban, amelyek az emberi szemlélő számára eltorzítják a végeredményt, és ekként láthatóvá teszik azt, ami többnyire láthatatlan marad (Betancourt-t [2017: 104] idézi Waszkiewicz 2019: 218): a kudarcos kommunikáció mögött meghúzódó szubjektum hiányát. A *glitch* ebben az értelemben töréspont, amelyen keresztül láthatóvá válik, hogy a generatív MI működése nem kommunikáció, hanem statisztikai valószínűség. Míg a freudi kísérteties esetében az idegenség az ismerős torz visszatéréséből fakad, az MI-*glitch*ekben éppen az válik nyugtalanítóvá, hogy nincs mit „félreismerni”: nem egy elrejtett szubjektum bukkan felszínre, hanem annak hiánya. Ez a hiány egyenesen a Mark Fisher által leírt hátborzongató (*eerie*) affektusához vezet el, amely akkor áll elő, amikor valahol ürességet találunk, ahol valaminek lennie kellene, vagy valamit találunk ott, ahol semminek nem kellene lennie (Fisher 2016). Fisher (2016: 11) szerint ez az empirikus tapasztalat egyben a transzcendenssel való szembesülést is implikálja, és alapvetően az ágenciához van köze: „Miféle ágens dolgozik itt? Van-e egyáltalán ágens?”

Banálisnak tűnő, mégis megvilágító erejű példa az MI által generált színezők esete, ahol a vizuális hasonlóság statisztikai logikája az emberi szemlélő számára poétikai értékű analógiává válhat. A gyermekeknek szánt színezők nemcsak a finommotorika fejlesztésére és tágabb értelemben a szabálykövetés habitusának bevézésére alkalmasak, hanem a minket körülvevő érzékelhető világ logikai megértését is elősegíthetik: innen lesheti el a rész–egész logikus viszonyát vagy a perspektíva fogalmát. Hogyan csatlakoznak egymáshoz például egy térforma adott elemei? Hogyan strukturálna mindez három dimenzióban? A jó színezők ezekre is rávezethetik a felhasználóikat. Bár a modellek ezen a fronton is egyre jobbak lesznek, az MI-színezőket érintő leggyakoribb kritika az, hogy a vonalak nem vezetnek sehova, emberi logikával nem színezhető, mert a rajta szereplő tárgyaknak „nincs értelmük”.⁶ Az a helyzet, hogy ezeknek a képeknek valóban „nincs értelmük” – egész egyszerűen nem a valóság hű reprezentációjának igényével vagy egyáltalán lehetőségével jönnek létre.

Az MI által generált színezőkben gyakran figyelhető meg, hogy egymástól eltérő funkciójú és jelentésű elemek anatómiai vagy formai jegyeik mentén összezsúsznak: a cicafülek könnyedén alakulnak át angyal szárnyakká, a levelek hajtincsekké vagy ruhafodrokká, a csúszdák lépcsőkké vagy a gyermekkarok indokolatlanul állatszerű mancsokká (lásd a 2. képet). Az ilyen képek nem „hibások” abban az értelemben, hogy egy felismerhető tárgy teljesen eltűnne; épp ellenkezőleg, túl sok minden ismerhető fel bennük egyszerre. A kontúrok megmaradnak, de a testhatárok bizonytalanná válnak, a funkciók felcserélődnek, az anatómia ornamentikává lapul. Fontos azonban hangsúlyozni, hogy ezek az átmenetek nem az MI esztétikai döntéseiből fakadnak, hanem abból, hogy a modell a látens térben egymáshoz közel elhelyezkedő vizuális jegyeket – például az ívességet, a szimmetriát, a puhaságot, az ismétlődést – egymással felcserélhetőként kezeli, mert ezek többnyire azonos kontextusban szoktak megjelenni. A giccs azonnali felismerésígérete itt egyszerre jelenik meg és bomlik fel, miközben a vizuális csúszások, a felcserélődő kontúrok, a bizonytalan testhatárok *glitch*-szerű töréspontként láttatják a modell statisztikai logikáját. A színezőkben, amelyeket a részletek felfedezésre terveztek, ezek az anatómiai és vizuális elcsúszások nem zavaró rendellenességként, hanem játékos, de mégis nyugtalanító átmenetekként jelennek meg.

6 Winkletter, Robb: The Flood of AI Coloring Books. *Medium*, 20 November 2022, <https://medium.com/@winkletter/the-flood-of-ai-coloring-books-2fe720667cfe>.

2. kép

A ChatGPT színezőgeneráló alkalmazásával készült a kép



Forrás: saját prompt alapján⁷

Az ilyen jellegű hibák a véletlenszerűség poétikáit hívják életre – azaz a koincidencia olyan analógiákat idéz elő, amelyeket az emberi befogadó kreatívként tud olvasni, és ezért érzékeljük nyugtalanítóként. Ezzel szemben a *glitch* extrém torzítása és az észlelt jelentés teljes elcsúszása nem a félreérthető ismerősre, hanem a gondolkodó, kombináló, értelmezni vágyó szubjektum teljes hiányára világít rá – és éppen ebben az értelemben az *eerie* affektív logikáját működteti. Erre jó, bár végletes példa lehet az a fotó, amely 2019-ben a következő szöveggel járta be a Twittert, majd az egész internetet: „Nevez meg akár egyetlen dolgot ezen a fotón” (lásd a 3. képet).⁸ A fotónak tűnő képen első pillantásra ismerősnek tűnő tárgyak és alakok szerepelnek, azonban amikor bármelyiken megállapodna a tekintet, egyszerűen lecsúszik róla az értelmezés, nem áll össze koherens tárggyá vagy jelentéssé. „Arra tervezték, hogy megmutassa, milyen lehet stroke-ot kapni” – kommentálta a jelenséget a Reddit, arra az állapotra utalva, amikor a percepció hibátlanul működik, de a látottak értelmezését nem tudja végrehajtani az agy.⁹ Maga a kép létezése azonban arra világít rá, hogy az „arra tervezték” kifejezés alapvetően helytelen. A „fotó” ugyanis a diffúziós modelleket megelőző GAN (*generative adversarial networks*) algoritmus segítségével jött létre – nem ült mögötte egy neuropszichológus, aki szándékosan finomhangolta volna ilyenre a képet valakik számára; egyszerűen egy szerencsés véletlenről van szó.¹⁰ Az ehhez hasonló példák nem az ismerős torzulására – a kísértetiesre – rímelnek, hanem az ürességbe, a jelentésnélküliségbe engednek bepillantást, így

⁷ A kép 2026 februárjában készült az alábbi prompt alapján: részletgazdag színezőlap, amelyen tündérmacskák játszanak egy erdei tisztáson, dús növényzet előtt.

⁸ Dumbass ass idiot: Name one thing in this photo, X, 23 April 2019, <https://t.co/zgyE9rL2XP>.

⁹ 22 November 2023, https://www.reddit.com/r/interestingasfuck/comments/1813s42/designed_to_give_the_viewer_the_simulated/.

¹⁰ Makowski, Dominique: What Visual Agnosia Might Feel Like. *Reality Bending Lab*, 3 January 2021, https://reality-bending.github.io/post/2021-01-03-visual_agnosia.

egyértelműen az *eerie* affektív mezőjébe tartoznak. Míg a jelentés már-már kényszeres keresése gyakorlatilag antropológiai állandónak tekinthető, az ilyen módon hibás MI-képek azzal szembesítik a nézőt, hogy jelentés helyett adatcsomagokról, kommunikáció helyett kódparancsokról és azok kimeneti eredményeiről beszélünk.

3. kép

„Nevez meg akár egyetlen dolgot ezen a képen!”



Forrás: X (Twitter)

A cápa és az edzőcipő véletlen találkozása a boncasztalon

A jelentés elvesztése, ami a hibás MI-képek tapasztalata, bizonyos tekintetben a kortárs kulturális klímát és az azt övező fő félelmeket is jellemzi. Ha hinni lehet a különböző orgánusok, szótárak „év szava” kezdeményezéseinek, ez a probléma kifejezetten foglalkoztatja a közbeszédet. 2024-ben például az *Oxford Dictionary* szerint az év szava a „*brain rot*” vagy – egyre inkább egybeírva – „*brainrot*” lett, utalva arra a jelenségre, amelyet tényleges mérési eredmények támasztanak alá: a (nemcsak fiatalok) koncentrációs képességének romlásáról, figyelmük szóródásáról, kritikus gondolkodásuk elsatnyulásáról, amelyet több kutatás is a rövid formátumú videós tartalmak fogyasztásával hoz összefüggésbe (Zhang et al. 2019, Chen et al. 2023, Yousef et al. 2025, Yazgan 2025). Ehhez a gondolathoz kapcsolódnak a 2025-ös év szavai is: például a *Merriam-Webster*, az *Economist* és a *Macquarie Dictionary* egyöntetűen a „*slop*”, „*AI slop*” szavakat választotta, utalva a silány minőségű, mesterséges intelligencia által létrehozott online tartalmakra, amelyek elárasztják a közösségi médiát.¹¹ Ehhez a tendenciához kapcsolható a Dictionary.com 2025-ös választása – „6-7” vagy „*six-seven*” – is, amely az alfa generáció számára bevallottan azért vicces és azért lehet különféle kontextusokban más-más jelentéssel használni, mert nem jelent semmit. A jelentésnélküliség viccként való visszatérése az olasz agyrohadás avagy *Italian brainrot* gyűjtőnéven ismert jelenség lényege is. Már maga a kifejezés is ambivalens, és a Z generáció posztiróniáját tükrözi:

¹¹ Word of the Year. *Wikipedia*, 25 January 2026, https://en.wikipedia.org/w/index.php?title=Word_of_the_year&oldid=1334680612; Khan, Coco: Each Year, Word of the Year Gets Darker. “Six-Seven” May Be Annoying – but It’s Bucked That Trend. *The Guardian*, 20 December 2025, <https://www.theguardian.com/commentisfree/2025/dec/20/words-of-the-year-six-seven-teenagers>.

miközben nem leplezi a szorongást, hogy a fokozott online jelenlét tényleg az agyi kapacitás csökkenésével jár együtt, a sajátjává is teszi, és viccet csinál belőle.

A több dolog vizuális hibridizációjából létrejövő, nagyon egyértelműen MI által generált, olaszosan hangzó, egyébként halandzsanevekkel és nehezen értelmezhető, különös megjelenéssel rendelkező alakok (mint például Ballerina Capuccina, Bombardino Crocodillo, Trulimero Trulicina), amellet, hogy (szintén MI által generált) dalok „előadóiként” vannak feltüntetve, különböző TikTok-műfajokat hívtak életre (például videókban harcoltatják őket egymással), de egész gyűjthető virtuális és tárgy-kultúra is kiépült körülöttük: matricák, színezők, 3D-nyomtatott figurák, plüssök. Sőt egy ingyenesen elérhető *Italian brainrot*-generátorral bárki létrehozhatja és meganimálhatja saját *brainrot*-figuráját. „Érthető, ha ezt olvasva úgy érzed, sorvad az agyad. Sőt valószínűleg ha egyssel kezdődő évben születél, az *Italian brain rot* nem neked való” – kommentálja a Guardian a szó szerint nonszensz vizuális és tárgy-kultúrát övező kisebb morális pánikot.¹²

Mindennek ellenére az olasz agyrohadás mégsem annak bizonyítéka, hogy a jelentés – egyáltalán a *jelenthetőség*, vagyis az, hogy egy adott médiatermék képes jelentés előállítására – visszavonhatatlanul eltűnt. A fentebb részletezett, MI által generált tartalmak többnyire azt a látszatot próbálják kelteni, hogy egy ágenciával rendelkező szubjektum áll mögöttük – amikor a *glitchek* felfedik, hogy ez nem így van, akkor áll elő az *eerie* affektusa. Ezekkel szemben az *Italian brainrot* meg sem próbálja ezt elhíttetni; büszkén és látványosan vállalja szimulákrumjellegét. Ebben az esetben a *weird* kategóriájáról beszélünk, amelyet Fisher a modernista és kísérleti alkotásokkal kapcsol össze – egyébként magát az olasz agyrohadást is látványosan könnyű a szürrealizmus művészi vállalkozásához kötni,¹³ ahogyan a fejezet címében is utaltam Comte de Lautréamont alapmondatára.

Fisher (Fisher 2016: 13; saját fordításom: O. R. Zs.) a következőképpen definiálja a furcsa (*weird*) affektív kategóriáját:

A nem helyénvalóság érzése – azaz a meggyőződés, hogy egy dolog nem illik oda, ahol van –, amelyet a furcsához kapcsolunk, gyakran az újdonság jelenlétét jelzi. A furcsa itt annak a jele, hogy a fogalmak és keretrendszerek, amelyeket addig használtunk, mostanra elavultak.

Az *Italian brainrot* karakterei éppen az effajta, látványosan eltúlzott keveredést viszik színre. Adrian Walker is ezt a tulajdonságukat emeli ki, amikor a foucault-i értelemben vett ideologikus kategóriák felbontását reméli az olasz agyrohadástól.¹⁴ Ezek a figurák nem úgy üresek, mint az MI hibás képei, hanem éppen, hogy túltelítettek: a cápa és az edzőcipő, a krokodil és a bombázó, a balerina és a kapucsinó. A jelentésválság aspektusából az a fontos, hogy nemcsak maguk a karakterek, hanem a kulturális remix, a látványos kombináció mint potenciális nóvumképzés gyakorlata is mémmé válik. És bár MI-technológiával jön létre, a poén kifejezetten emberi.

Összegzés

Térjünk vissza e gondolatmenet fő tétjéhez, azaz ahhoz a felvetéshez, hogy az MI-generálta képek és szövegek radikálisan újraértelmezik a jelentő és jelentett közötti, többnyire adottnak vett viszonyt. Mivel a jelentés válságáról és (talán) főnixként való újjászületéséről van itt szó, az a diszciplína érezheti itt különösképpen megszólítva magát, amely elsősorban ennek a kulturális tényezők és formai jegyek által meghatározott jelentésnek az elő-

12 Hunt, Elle: From Chimpanzini Bananini to Ballerina Cappuccina: How Gen Alpha Went Wild for Italian Brain Rot Animals. *The Guardian*, 25 June 2025, <https://www.theguardian.com/society/2025/jun/25/from-chimpanzini-banani-to-ballerina-cappuccina-how-gen-alpha-went-wild-for-italian-brain-rot-animals>. Saját fordításom: O.R.Zs.

13 Katz, Leslie: What Is “Italian Brain Rot”? The Surreal TikTok Obsession, Explained. *Forbes*, 3 May 2025, <https://www.forbes.com/sites/lesliekatz/2025/05/03/what-is-italian-brain-rot-the-surreal-tiktok-obsession-explained/>, Titterington, David: Italian Brainrot: The New Surrealism? *Medium*, 30 December 2025, <https://davidtitterington.medium.com/italian-brainrot-the-new-surrealism-a151cdeb4ce7>.

14 Walker, Aidan: Brainrot and the Order of Things. *How To Do Things With Memes*, 14 March 2025, <https://howtodothingswithmemes.substack.com/p/brainrot-and-the-order-of-things>.

állításában, dekódolásában, kontextualizálásában érdekelt: a hermeneutikai beállítódású bölcsészettudomány. A klasszikus hermeneutikai hagyomány a Másikat mindig potenciálisan elsajátíthatónak tételezi, a *brainrot*tal azonban eljutunk oda, hogy mind a Másik létezésének bizonyossága, mind az elsajátíthatóság lehetősége meg-inog. Az MI képek fent elemzett három aspektusa – a jelentés elillanásának stációiként – a saját és a másik mezői közötti felületi feszültségként is értelmezhetők. A giccsnek megfeleltethető képek a kísérteties affektusát hívták életre, amennyiben a saját folyamatos – de torz – visszatérésével operálnak. Ebben a felállásban a másik számára nincsen hely. Az MI által vétett hibák ezzel szemben a hátborzongató (*eerie*) kategóriájába tartoztak, mert arra az antihermeneutikai belátásra világítottak rá, hogy nem minden – akár poétikusnak tételezett – idegenség mögött áll egy értelmezhető másik. Fisher (2016: 8) szerint a *weird* és az *eerie* abban közös, hogy „lehetővé teszik, hogy a belsőt a külső perspektívájából lássuk”. Az *Italian brainrot* ebben az értelemben nem a jelentés eltűnését, hanem a jelentésképzés folyamatainak kontextualizált s ekként feltételes voltát teszi láthatóvá. Az MI segítségével létrehozott, *weird* képek arra világítanak rá, hogy a jelentéskommunikáció mint alapfeltevés többé nem magától értetődő – és talán soha nem is volt az, de talán épp ebből a bizonytalanságból jöhet létre valami, ami valahol, valamikor, valaki számára, újra értelmet hordoz.

Felhasznált irodalom

- Ananya (2024): AI Image Generators Often Give Racist and Sexist Results: Can They Be Fixed? *Nature*, vol. 627, no. 8005, pp. 722–725, <https://doi.org/10.1038/d41586-024-00674-9>
- Barker, Jessica (2020): Making-up on Mobile: The Pretty Filters and Ugly Implications of Snapchat. *Fashion, Style & Popular Culture*, vol. 7, no. 2–3, pp. 207–221, https://doi.org/10.1386/fspc_00015_1
- Betancourt, Michael (2017): The Invention of Glitch Video: Digital TV Dinner. *Millennium Film Journal*, no. 65, pp. 54–63.
- Broch, Hermann & John Hargraves (2002): *Geist and Zeitgeist: The Spirit in an Unspirited Age: Six Essays*. Catapult.
- Chen, Yuhan, Mingming Li, Fu Guo & Xueshuang Wang (2023): The Effect of Short-Form Video Addiction on Users' Attention. *Behaviour & Information Technology* vol. 42, no. 16, pp. 2893–2910, <https://doi.org/10.1080/0144929X.2022.2151512>
- Eshiet, Janella (2020): “Real Me versus Social Media Me:” Filters, Snapchat Dysmorphia, and Beauty Perceptions among Young Women. Electronic Theses, Projects, and Dissertations. 1101.
- Fisher, Mark (2014): *Ghosts of My Life: Writings on Depression, Hauntology and Lost Futures*. John Hunt Publishing.
- Fisher, Mark (2016): *The Weird and the Eerie*. Repeater.
- Freud, Sigmund (2003): *The Uncanny*. Penguin.
- Hintze, Arend, Frida Proschinger Åström & Jory Schossau (2026): Autonomous Language-Image Generation Loops Converge to Generic Visual Motifs. *Patterns*, vol. 7, no. 1, <https://doi.org/10.1016/j.patter.2025.101451>
- Isakowitsch, Clara (2022): How Augmented Reality Beauty Filters Can Affect Self-Perception. *Irish Conference on Artificial Intelligence and Cognitive Science*, pp. 239–250, https://doi.org/10.1007/978-3-031-26438-2_19
- Placani, Adriana (2024): Anthropomorphism in AI: Hype and Fallacy. *AI and Ethics*, vol. 4, no. 3, pp. 691–698, <https://doi.org/10.1007/s43681-024-00419-4>
- Rajanala, Susruthi, Mayra BC Maymone & Neelam A Vashi (2018): Selfies. Living in the Era of Filtered Photographs. *JAMA Facial Plastic Surgery*, vol. 20, no. 6, pp. 443–444, <https://doi.org/10.1001/jamafacial.2018.0486>
- Waszkiewicz, Agata (2019): Glitch as the Representation of the Uncanny in Oxenfree (2016). *Homo Ludens*, vol. 12, no. 1, pp. 213–225, <https://doi.org/10.14746/hl.2019.12.11>
- Yazgan, Ayşe Müge (2025): The Problem of the Century: Brain Rot. *OPUS Journal of Society Research*, vol. 22, no. 2, pp. 211–221, <https://doi.org/10.26466/opusjsr.1651477>

Yousef, Ahmed Mohamed Fahmy, Alsaeed Alshamy, Ahmed Tlili & Ahmed Hosny Saleh Metwally (2025): Demystifying the New Dilemma of Brain Rot in the Digital Era: A Review. *Brain Sciences*, vol. 15, no. 3, <https://doi.org/10.3390/brainsci15030283>

Zhang, Chenshuang, Chaoning Zhang, Mengchun Zhang, In So Kweon & Junmo Kim (2024): Text-to-Image Diffusion Models in Generative AI: A Survey. *arXiv*, preprint, <https://doi.org/10.48550/arXiv.2303.07909>

Zhang, Xing, You Wu & Shan Liu (2019): Exploring Short-Form Video Application Addiction: Socio-Technical and Attachment Perspectives. *Telematics and Informatics*, vol. 42, no. 101243, <https://doi.org/10.1016/j.tele.2019.101243>