

A kulturális evolúció hozza létre a nyelv statisztikai szerkezetét

BEVEZETÉS

A nyelv két sajátossága biztosítja azt az alapot, amelyre a nyelvi struktúra egésze épül, meglepő azonban, hogy eredetük máig kevésbé ismert. Egyrészt a nyelv kisebb, szekvenciálisan (újra)kombinálható részekre tagolódik: a hangokat szavakká, a szavakat mondatokká lehet szervezni. Másrészt a szavak gyakorisági eloszlása rendkívül aránytalan. Míg a világ nyelvei több szempontból különböznek, abban szisztematikus hasonlóságot mutatnak, hogy meghatározott gyakorisági eloszlást követő részekből állnak.³ A tanulmányban – a kulturális evolúció kísérleti modelljére támaszkodva – bemutatjuk, hogy a nyelv ezen sajátosságok által tanulhatóbbá válik, továbbá hogy ezek generációról generációra, a kulturális átadás során alakulnak ki a nyelvhasználók ismétléses (iteratív) tanulása révén.

A nyelvet elsajátító csecsemők számára alapvető kihívás a megfelelő alegységek felismerése. A beszélt nyelvben ugyanis, számos írott nyelvvel ellentétben, a szóhatárok nincsenek egyértelműen jelezve. Ugyan többféle támpont segíthet a szóhatárok azonosításában (fonotaktika, a hangsúly jelölése, statisztikai információ stb.),⁴ egyik sem teljesen megbízható, és mindegyiket maguktól kell a csecsemőknek felfedezniük. A megfelelő egységek felismerése ugyanakkor kulcsfontosságú lépés ahhoz, hogy később eredményesen tudják azokat kombinálni. De hogyan fedezik fel a csecsemők a részeket anélkül, hogy tudnák, mit is kell keresniük?

Egy lehetséges válasz, hogy a gyermekek képesek felismerni azokat a statisztikai szabályszerűségeket, amelyek segítik a beszélt nyelv szóhatárainak kijelölését. Egy ilyen típusú, már széleskörűen tanulmányozott szabályszerű-

¹ A Jeruzsálemi Héber Egyetem Pszichológiai Tanszékének kutatója, Jeruzsálem, Izrael.

² Az Edinburghi Egyetem Filozófiai, Pszichológiai és Nyelvtudományi Intézetének kutatója, Edinburgh, Egyesült Királyság, e-mail: inbal.arnon@mail.huji.ac.il.

³ George ZIPF, *Human Behavior and the Principle of Least Effort* (Cambridge, MA: Addison-Wesley, 1949).

⁴ Peter W. JUSZYK, „How infants begin to extract words from speech”, *Trends in Cognitive Sciences* 3, 9. sz. (1999): 323–328, [https://doi.org/10.1016/S1364-6613\(99\)01363-7](https://doi.org/10.1016/S1364-6613(99)01363-7).

ség a szótagok egymásra következtési valószínűsége.⁵ Eszerint a gyermekek alapszintű statisztikai mintázatokra támaszkodhatnak abban, hogy megállapítsák, melyik szótag követ nagy valószínűséggel egy másikat egy szón belül, és hol vannak a szóhatárok. Ha a nyelv kombinálható részek összessége, akkor a részeken belüli következtési valószínűség nagyobb lesz, mint az egységhatárok között. Ez a legtöbbször így is van: a szavakat alkotó szótagok következtési valószínűsége nagyobb, mint azoké a szótagoké, amelyek túlnyúlnak a szóhatárokon.⁶ Sok különböző hang jelenhet meg egy szó után, mert sokféle szó követheti azt, míg egy szón belül kevesebb féle hang szerepelhet a kezdőhangot követően. Vegyük például a „aranyos kisbaba” (*pretty baby*) szekvenciát. Különböző szavak jelenhetnek meg az „aranyos” szó után (például autó, fiú, kalap, macska stb.), de csak kevés hang szerepelhet az „ara-” [*pre*] után egy értelmes szekvenciában (arabeszk, arabul stb.). Emiatt nagyobb a szó belsejében található szótag (példánkban a „-nyos” [*tty*]) következtési valószínűsége, mint a szóhatáron túlnyúló szótagoké (a „kis-” [*ba*] a „-nyos” [*ty*] után). Bőséges empirikus bizonyíték támasztja alá, hogy a csecsemők, a gyermekek és a felnőttek is érzékelik ezeket a valószínűségeket, amit egy új, tagolatlan beszédflow részre bontásakor is hasznosítani tudnak.⁷ Továbbá mindez nem csak a nyelvre korlátozódik. A valószínűségek változatai más területeken is segítik a határkijelölést: a nagy valószínűség az elemek csoportosítását, míg az alacsony az elemek elkülönítését vonja maga után. A csecsemők például a következtési valószínűséget vizuális alakzatok esetében is használni tudják, hogy ismétlődő elemeket fedezzenek fel alakzatok folytonos sorozatában.⁸

A statisztikai tanulásról szóló szakirodalom megoldást kínál arra, hogy a nyelvtanulók hogyan bonthatnak részekre egész megnyilatkozásokat, feltéve, hogy ezek a részek már eleve léteznek. Ha tehát a nyelv kombinálható részekből áll, a statisztikai tanulás kiszűrheti azokat. Marad azonban egy megválaszolandó kérdés. Hogyan válik a nyelv olyan részekből felépülő rendszerre, amelyek kombinációja egy ilyen típusú tanulási stratégiát hatékonyra tesz?

⁵ Jenny R. SAFFRAN, Richard N. ASLIN és Elissa L. NEWPORT, „Statistical learning by 8-month-old infants”, *Science* 274 (1996): 1926–1928, <https://doi.org/10.1126/science.274.5294.1926>.

⁶ Adele SAKSIDA, Alessandro LANGUS és Marina NESPOR, „Co-occurrence statistics as a language-dependent cue for speech segmentation”, *Developmental Science* 20, 3. sz. (2017): e12390, <https://doi.org/10.1111/desc.12390>.

⁷ Jenny R. SAFFRAN és Natasha Z. KIRKHAM, „Infant statistical learning”, *Annual Review of Psychology* 69 (2018): 181–203, <https://doi.org/10.1146/annurev-psych-122216-011805>.

⁸ Natasha Z. KIRKHAM, Jessica A. SLEMMER és Scott P. JOHNSON, „Visual statistical learning in infancy: Evidence for a domain general learning mechanism”, *Cognition* 83, 2. sz. (2002): B35–B42, [https://doi.org/10.1016/S0010-0277\(02\)00004-5](https://doi.org/10.1016/S0010-0277(02)00004-5).

Más szóval, honnan származnak a részek? Kézenfekvő magyarázat, hogy a nyelv azért áll részekből, mert a közvetíteni kívánt jelentés természetesen ugyancsak részekből áll. Azaz a nyelv szerkezete abból a késztetésből fakad, hogy kifejezően beszélhessünk a világról, amely maga is strukturált.⁹ Jelen tanulmányban ugyanakkor egy ennél erősebb hipotézist vizsgálunk: az egy-
ségek pusztán a tanulhatóság hatására jelennek meg, még akkor is, ha nincs közvetíteni kívánt jelentés. Ez egy csábító hipotézis annak megválaszolására, hogy a csecsemők miért kezdik részekre bontani a nyelvet, mielőtt megtanulnák a részek jelentését.¹⁰ Ezt erősíti továbbá, hogy más megtanulható rendszerek (például a zene) is jelentés nélküli részek kombinációjára épülnek.

A nyelv statisztikai szerkezetének egy másik figyelemre méltó sajátossága a részek gyakorisági eloszlására vonatkozik. Ha megnézzük bármely nyelv szó-készletét, felűnő, hogy nem minden szó fordul elő egyforma gyakorisággal. A szavak gyakorisági eloszlása ennek megfelelően nagyon ferde: csupán kevés szó fordul elő nagyon sokszor, míg a legtöbb kifejezetten ritkán. A gyakoriság továbbá nem lineárisan csökken. Egy nyelv szavainak az előfordulási eloszlása inkább a hatványtörvény szerinti eloszlást követi úgy, hogy egy szó gyakorisága fordítottan arányos a gyakorisági rangsorban betöltött helyével (tehát a leggyakoribb szó kétszer olyan gyakori, mint a második, és háromszor, mint a harmadik – és így tovább). Ezt a szabályt először George Kingsley Zipf mutatta ki, ezért szokás Zipf-eloszlásként hivatkozni rá.¹¹ A szavak követik a Zipf-eloszlást (vagy egy kvázi-Zipf-eloszlást¹²) a beszélt nyelvekben és jelnyel-

⁹ Simon KIRBY, Hannah CORNISH és Kenny SMITH, „Cumulative cultural evolution in the laboratory: An experimental approach to the origins of structure in human language”, *Proceedings of the National Academy of Sciences* 105 (2008): 10681–10686, <https://doi.org/10.1073/pnas.0707835105>.

¹⁰ JUSCZYK, „How infants begin...”, 323–328; Anna FLÓ, Judit GERVAIN, Judit KISS, Judit KOVÁCS és Csaba PLÉH, „Newborns are sensitive to multiple cues for word segmentation in continuous speech”, *Developmental Science* 22 (2019): e12802, <https://doi.org/10.1111/desc.12802>.

¹¹ ZIPF, *Human Behavior*...

¹² Steven T. PIANTADOSI, „Zipf’s word frequency law in natural language: A critical review and future directions”, *Psychonomic Bulletin & Review* 21 (2014): 1112–1130, <https://doi.org/10.3758/s13423-014-0585-6>.

vekben,¹³ a gyermeknek szóló beszédben¹⁴ és a nyelvek különböző grammatikai kategóriáiban is.¹⁵ Azonban folyamatos vita zajlik az efféle eloszlások eredetéről és arról, hogy vajon ezek eredendő sajátosságai-e az emberi kogníciónak és kommunikációnak.¹⁶

Érdekes módon úgy tűnik, hogy a Zipf-eloszlás segíti a nyelvtanulást. A szavak szegmentálása egy mesterséges nyelvben javul, ha a tanulók ferde eloszlással szembesülnek egyenletes helyett (amelyben minden elem egyforma valószínűséggel jelenik meg).¹⁷ Sőt, a résztvevők akkor képesek a legjobban a szekvenciákat konstitutív elemeire bontani, amikor a részek a természetes nyelvekhez hasonló (az entrópiában megmutatkozó)¹⁸ ferde eloszlást követnek. A ferde eloszlásról az is kiderült, hogy elősegíti a vizuális statisztikai tanulást (ismétlődő vizuális hármassok felismerését egy összefüggő fo-

¹³ PIANTADOSI, „Zipf’s word frequency law in natural language...”, 1112–1130; Ilanit KIMCHI, Rachel STAMPS, Liam WOLTERS és Isaac ARNON, „Evidence of Zipfian distributions in three sign languages”, *Gesture* (elfogadva, megjelenés alatt); Ali MEHRI és Mohammad JAMAATI, „Variation of Zipf’s exponent in one hundred live languages: A study of the *Holy Bible* translations”, *Physics Letters A* 381, 31. sz. (2017): 2470–2477, <https://doi.org/10.1016/j.physleta.2017.05.061>; Ramon FERRER-I-CANCHO, „The variation of Zipf’s law in human language”, *European Physical Journal B* 44, 2. sz. (2005): 249–257, <https://doi.org/10.1140/epjb/e2005-00121-8>.

¹⁴ Ofir LAVI-ROTBAIN és Isaac ARNON, „Zipfian Distributions in Child-Directed Speech”, *Open Mind* 7 (2022): 1–30, https://doi.org/10.1162/opmi_a_00070.

¹⁵ PIANTADOSI, „Zipf’s word frequency law in natural language...”, 1112–1130; Carl BORSTELL, „Searching and utilizing corpora”, in *Signed Language Corpora*, szerk. Julie A. HOCHGESANG és Jordan FENLON, 90–127 (Washington, DC: Gallaudet University Press, 2022), <https://doi.org/10.2307/j.ctv2rcnfhc.9>.

¹⁶ PIANTADOSI, „Zipf’s word frequency law in natural language...”, 1112–1130; Ramon FERRER-I-CANCHO, Christian BENTZ és Cyrille SEGUIN, „Optimal coding and the origins of Zipfian laws”, *Journal of Quantitative Linguistics* 29, 2. sz. (2020): 165–194, <https://doi.org/10.1080/09296174.2020.1778387>; Edward GIBSON, Steven T. PIANTADOSI, Kyle MAHOWNEY, Evelina FEDORENKO, Roman FEDELENKO, Roger LEVY, Leon BERGEN, Martin T. SCHNUR és John T. HALE, „How efficiency shapes human language”, *Trends in Cognitive Sciences* 23, 5. sz. (2019): 389–407, <https://doi.org/10.1016/j.tics.2019.02.003>; Sarah SEMPLÉ, Ramon FERRER-I-CANCHO és Melissa L. GUSTISON, „Linguistic laws in biology”, *Trends in Ecology & Evolution* 37, 1. sz. (2022): 53–66, <https://doi.org/10.1016/j.tree.2021.08.012>.

¹⁷ Chigusa KURUMADA, Stephan C. MEYLAN és Michael C. FRANK, „Zipfian frequency distributions facilitate word segmentation in context”, *Cognition* 127, 3. sz. (2013): 439–453, <https://doi.org/10.1016/j.cognition.2013.02.002>; Ofir LAVI-ROTBAIN és Isaac ARNON, „The learnability consequences of Zipfian distributions in language”, *Cognition* 223 (2022): 105038, <https://doi.org/10.1016/j.cognition.2022.105038>.

¹⁸ Uo.

lyamban);¹⁹ a szituációkon átívelő szótanulást;²⁰ valamint az új grammatikai kategóriák²¹ és konstrukciók²² tanulását. A beszélőknek láthatólag kognitív preferenciája az efféle eloszlás: még kitalált főnevek egyenletes eloszlása is ferdévé válik, amikor a résztvevők egy történetet adnak tovább egymásnak.²³ Bár kevés tanulmány vizsgálja e mintázat hatását a nyelvtanulásra, ezek szintén a ferde eloszlás előnyeire világítanak rá számos nyelvi és nemnyelvi viszony esetében. Ezek az eredmények azt mutatják, hogy a Zipf-eloszlás elősegítheti a nyelvtanulást, mégis nyitva hagynak egy másik kérdést: miért alakul ki ez a tanulást megkönnyítő statisztikai törvény a nyelvben?

A nyelv mindkét visszatérő és szisztematikus sajátossága – az, hogy egységekből áll, és ezek meghatározott eloszlást követnek – a tanulhatóságban játszik fontos szerepet: segítik a tanulót a lényeges nyelvi egységek felismerésében és megjegyzésében. A tanulmányban azt mutatjuk be, hogy a nyelv statisztikai sajátosságai nemcsak elősegítik a tanulást, hanem ez utóbbi legfőbb oka létezésüknek. A nyelv statisztikai szerkezete tehát egyszerre oka és okozata a hatékony tanulásnak. Ez a látszólag körbenforgó érv egy olyan nyelvszemléletből ered, amely a kulturális átadás folyamatára épül. Ennek értelmében a generációk a korábbiak tanulásának eredményeként létrejött kimenetet sajátítják el. Más szóval a nyelv, amelyet a gyermek megpróbál szegmentálni, egy hasonló szegmentációs folyamat eredménye, amelyen nevelőik vagy a nyelvi közösség más tagjai már korábban keresztülmentek. És ugyanígy az ő tanulásuk során kialakult változat is nyelvi adattal fog szolgálni a jövőbeli generációk számára. Az ismétléses tanulás ilyen folyamata megfelelő magyarázatnak bizonyult a nyelv több szerkezeti sajátossága esetében

¹⁹ Ofir LAVI-ROTBAIN és Isaac ARNON, „Visual statistical learning is facilitated in Zipfian distributions”, *Cognition* 206 (2021): 104492, <https://doi.org/10.1016/j.cognition.2020.104492>.

²⁰ Anne T. HENDRICKSON és Amy PERFORs, „Cross-situational learning in a Zipfian environment”, *Cognition* 189 (2019): 11–22, <https://doi.org/10.1016/j.cognition.2019.03.005>.

²¹ Katherine D. SCHULER, Patricia A. REEDER, Elissa L. NEWPORT és Richard N. ASLIN, „The effect of Zipfian frequency variations on category formation in adult artificial language learning”, *Language Learning and Development* 13, 4. sz. (2017): 357–374, <https://doi.org/10.1080/15475441.2016.1263571>.

²² Jennifer K. BOYD és Adele E. GOLDBERG, „Input effects within a constructionist framework”, *The Modern Language Journal* 93, 3. sz. (2009): 418–429, <https://doi.org/10.1111/j.1540-4781.2009.00899.x>.

²³ Amir SHUFANIYA és Isaac ARNON, „A cognitive bias for Zipfian distributions?: Uniform distributions become more skewed via cultural transmission”, *Journal of Language Evolution* 7, 1. sz. (2022): 59–80, <https://doi.org/10.1093/jole/lzac005>.

is, úgy mint a kompozicionalitás,²⁴ a kettős tagolás,²⁵ a szabályosság²⁶ vagy a szavak szemantikája²⁷ kapcsán, valamint még az olyan nemnyelvi struktúrák esetében is, mint a dobminták²⁸ és a versformák kialakulása.²⁹ Általánosságban véve a több generáción át tartó ismétléses tanulás olyan nyelvek kialakulásához vezet, amelyek alkalmazkodnak a nyelvet továbbadók tanulási preferenciáihoz.³⁰

Érvelésünkben az ismétléses tanulás folyamata egy meghatározott tanulási stratégiával párosítva – amely az egésztől halad a részek felé – hozza létre a nyelv két statisztikai sajátosságát. Míg a nyelvelsajátítás tankönyvi leírása gyakran a résztől az egész felé haladó tanulást hangsúlyozza, amelyben a kisebb egységek nagyobb egészekké állnak össze (például a szótagok szavakká, a szavak mondatokká), egyre több bizonyíték van arra, hogy az egésztől a részek felé haladó stratégia – amelyben az egész később részekké bomlik fel – szintén fontos szerepet játszik,³¹ elsősorban a nyelvi struktúrák megtanulásában. A csecsemők ugyanis kezdetben nem tudják, hol vannak a szóhatárok, vagy hogy egyáltalán a szavak léteznek. Korai építőelemeiket szavak egyvelege és többszavas mondatok képezik, amelyek csak később bomlanak részekre (például: mi-ez [*what's-this*]). Ha a nyelvtanulást a nagyobb elemek felől kezdjük, amelyeket csak ezután szegmentálunk, könnyebb a szavak közötti önké-

²⁴ KIRBY, CORNISH és SMITH, „Cumulative cultural evolution in the laboratory...”, 10681–10686.

²⁵ Thomas VERHOEF, Simon KIRBY és Bart DE BOER, „Emergence of combinatorial structure and economy through iterated learning with continuous acoustic signals”, *Journal of Phonetics* 43 (2014): 57–68, <https://doi.org/10.1016/j.wocn.2014.02.005>.

²⁶ Uo.

²⁷ James W. CARR, Kenny SMITH, Jennifer CULBERTSON és Simon KIRBY, „Simplicity and informativeness in semantic category systems”, *Cognition* 202 (2020): 104289, <https://doi.org/10.1016/j.cognition.2020.104289>.

²⁸ Andrea RAVIGNANI, Teresa DELGADO és Simon KIRBY, „Musical evolution in the lab exhibits rhythmic universals”, *Nature Human Behaviour* 1, 1. sz. (2016): 0007, <https://doi.org/10.1038/s41562-016-0007>.

²⁹ Valeria DE CASTRO-ARRAZOLA és Simon KIRBY, „The emergence of verse templates through iterated learning”, *Journal of Language Evolution* 4, 1. sz. (2019): 28–43, <https://doi.org/10.1093/jole/lzy013>.

³⁰ Michael L. KALISH, Thomas L. GRIFFITHS és Simon KIRBY, „Iterated learning: Intergenerational knowledge transmission reveals inductive biases”, *Psychonomic Bulletin & Review* 14, 2. sz. (2007): 288–294, <https://doi.org/10.3758/BF03194066>; Simon KIRBY, Mike DOWMAN és Thomas L. GRIFFITHS, „Innateness and culture in the evolution of language”, *Proceedings of the National Academy of Sciences* 104 (2007): 5241–5245, <https://doi.org/10.1073/pnas.0608222104>.

³¹ Uo.; Isaac ARNON, *Starting Big – The Role of Multi-word Phrases in Language Learning and Use* (Stanford, CA: Stanford University, 2010).

nyes grammatikai viszonyok elsajátítása (ilyen például a nyelvtani nemmel rendelkező névelők és a főnevek közötti viszony egyes nyelvekben).³² Mivel a felnőttek már tudják, hogy léteznek szavak, és ismerik is anyanyelvük szó-készletét, kevésbé támaszkodnak a többszavas egységekre, ami hátrányos a grammatikai viszonyok tanulásakor. Empirikus bizonyíték szolgál arra, hogy a csecsemők már a korai szakaszban azonosítanak többszavas egységeket,³³ a felnőttek viszont ezt kevésbé teszik a nyelvtanulás során;³⁴ a nagyobb egységek előtérbe helyezése ugyanakkor elősegíti a nyelvtan elsajátítását.³⁵ Ezek a tanulmányok jól szemléltetik az egésztől a részek felé haladó tanulás jelenlétét a nyelvelsajátítás és a nyelvi struktúrák megismerése során.

A statisztikai tanulásról, az ismétléses tanulásról és az egésztől a részek felé haladó tanulásról szóló szakirodalmakra építve az a sejtésünk, hogy a nyelvi elemek eloszlásának mintázatai szintén a kulturális átadás során jönnek létre. Ha az egésztől a részek felé haladó tanulási stratégia generációkon át ismétlődik, akkor a kezdetben holisztikus nyelvből fokozatosan kialakulnak

³² Greville G. CORBETT, „Gender, grammatical”, in *Encyclopedia of Language & Linguistics*, szerk. Keith BROWN, 749–756 (Amsterdam: Elsevier, 2006²), <https://doi.org/10.1016/B0-08-044854-2/00191-7>.

³³ Isaac ARNON, Shane M. McCAULEY és Morten H. CHRISTIANSEN, „Digging up the building blocks of language: Age-of-acquisition effects for multiword phrases”, *Journal of Memory and Language* 92 (2017): 265–280, <https://doi.org/10.1016/j.jml.2016.07.004>; Barbora SKARABELA, Mitsuhiko OTA, Rory O’CONNOR és Isaac ARNON, „‘Clap your hands’ or ‘take your hands’?: One-year-olds distinguish between frequent and infrequent multiword phrases”, *Cognition* 211 (2021): 104612, <https://doi.org/10.1016/j.cognition.2021.104612>.

³⁴ Shane M. McCAULEY és Morten H. CHRISTIANSEN, „Computational investigations of multiword chunks in language learning”, *Topics in Cognitive Science* 9, 3. sz. (2017): 637–652, <https://doi.org/10.1111/tops.12258>; Noam HAVRON és Isaac ARNON, „Reading between the words: The effect of literacy on second language lexical segmentation”, *Applied Psycholinguistics* 38 (2017): 127–153, <https://doi.org/10.1017/S0142716416000138>.

³⁵ Isaac ARNON és Eve V. CLARK, „Why ‘brush your teeth’ is better than ‘teeth’—Children’s word production is facilitated in familiar sentence-frames”, *Language Learning and Development* 7 (2011): 107–129, <https://doi.org/10.1080/15475441.2010.505489>; Isaac ARNON és Michael RAMSCAR, „Granularity and the acquisition of grammatical gender: How order-of-acquisition affects what gets learned”, *Cognition* 122 (2012): 292–305, <https://doi.org/10.1016/j.cognition.2011.10.009>; Noam SIEGELMAN és Isaac ARNON, „The advantage of starting big: Learning from unsegmented input facilitates mastery of grammatical gender in an artificial language”, *Journal of Memory and Language* 85 (2015): 60–75, <https://doi.org/10.1016/j.jml.2015.07.003>; Noam HAVRON, Limor RAVIV és Isaac ARNON, „Literate and pre-literate children show different learning patterns in an artificial language learning task”, *Journal of Cultural and Cognitive Science* 2 (2018): 21–33, <https://doi.org/10.1007/s41809-018-0015-9>.

újrarendezhető részek.³⁶ Mindez egy idő után meg kell mutatkozzon a létrejövő nyelv statisztikai sajátosságaiban, vagyis az egységeken belüli következési valószínűségeknek nagyobbaknak kell lenniük, mint az egységhatárokon. Ezen kívül, ha a részek ferde eloszlása egyszerre oka és okozata a hatékonyabb tanulásnak, azt várjuk, hogy az eloszlás idővel egyre ferdebb lesz, és még inkább illeszkedik a nyelv hatványtörvényszerű működéséhez.

A hipotézis ellenőrzésére a nyelv kulturális evolúciójának egy kísérleti modelljét hívjuk segítségül, amelyben a résztvevők „nyelvet” tanulnak. A soron következő résztvevőnek reprodukálnia kell egy olyan „nyelvet”, amelyet egy korábbi résztvevő generált; a kimenet tehát bemenetté válik a következő résztvevő számára. A kísérletben a közvetített tartalom nem nyelvi jellegű, hanem színszekvenciákból áll. Mindegyik szekvencia egy megnyilatkozásnak felel meg, összességük pedig maga a megtanulandó „nyelv”. Alapvető fontosságú, hogy a kezdeti szekvenciahalmaz nem strukturált: mindegyik színkombináció egyformán lehetséges, azaz a szekvenciák nem ismétlődő részekből állnak. Kérdés, hogy a kezdeti sorozatot a résztvevők felruházzák-e az átadás folyamán a nyelvben fellelhető két statisztikai sajátossággal: részekből fog-e állni, amelyek egy meghatározott gyakorisági eloszlást követnek. A részek azonosításához a statisztikai tanulásról szóló szakirodalomra támaszkodunk (a részleteket lásd lentebb).

Fontos megjegyezni, hogy a feladat explicit módon egy egész szekvencia másolására irányul, mintsem kombinálható részek megtalálására; és úgy alakítottuk ki, hogy ez a másolás nehéz legyen, ami nyomást gyakorol a rendszerre, hogy megváltozzon és tanulhatóbbá váljon. Két okból adunk nem nyelvi feladatot. Egyrészt azért, mert ekkor sem a részeknek (amelyeknek ki kell alakulniuk), sem az egésznek (a szekvencia) nincs jelentése. Ez lehetővé teszi, hogy ellenőrizzük hipotézisünket, miszerint a statisztikailag elkülönülő részek, amelyek meghatározott ferde eloszlásba rendeződnek, a jelentéstől függetlenül alakulnak ki, illetve elősegítik a rendszer további elsajátítását. Másrészt a nem nyelvi ingerek, főként a jelentés nélküli egyedi elemek (és azok kombinációi), minimalizálják annak kockázatát, hogy a résztvevők egyszerűen áthelyezzék a saját nyelvük statisztikai szerkezetének ismeretét a kulturálisan átadott tartalmakra. Így tehát a kombinálható részek létrejötte a

³⁶ Alison WRAY, „Protolanguage as a holistic system for social interaction”, *Language & Communication* 18, 1. sz. (1998): 47–67, [https://doi.org/10.1016/S0271-5309\(97\)00033-5](https://doi.org/10.1016/S0271-5309(97)00033-5); Simon KIRBY, „Syntax without natural selection: How compositionality emerges from vocabulary in a population of learners”, in *The Evolutionary Emergence of Language: Social Function and the Origins of Linguistic Form*, szerk. Chris KNIGHT, 303–323 (Cambridge: Cambridge University Press, 2000), <https://doi.org/10.1017/CBO9780511606441.019>.

tanulók implicit preferenciáinak vagy elfogultságainak (*bias*) az eredménye lesz az ismétléses tanulás során, nem pedig az általuk bevezetett tudatos szabályokénak.

MÓDSZER

Kutatásunkban egy korábbi kísérlet eredményeit elemezzük újra,³⁷ amely szintén színszekvenciák ismétléses tanulását vizsgálta, hogy ellenőrizzük hipotézisünket, miszerint a nyelv statisztikai szerkezete a kulturális evolúció révén jön létre a nyelvtanulók ismétléses, az egésztől a részek felé haladó tanulása során. Az eredeti kísérletben azt vizsgáltuk, hogy vajon az absztrakt, ismétléses tanulás által közvetített szekvenciák generációnként tanulhatóbakká válnak-e. Nem elemeztük azonban a szekvenciák belső struktúráját, és a kísérletet sem ennek érdekében alakítottuk ki. Ennek az adathalmaznak az újraelemzésével tehát kifejezetten hatékonyan ellenőrizhetjük hipotézisünket, hiszen ez felszámolja a kutató elfogultságát, amely akkor keletkezik, ha egy kísérletet egy már előre remélt eredmény fényében tervez meg.

A kísérletet a „Simon” elnevezésű társasjáték online változatának mintájára alkottuk meg. A résztvevők felvillanó fényeket látnak egy kétszer kettes színtáblán (piros, zöld, sárga, kék). Közvetlenül ezután vissza kell idézniük ezt a fényszekvenciát a vonatkozó gombok megnyomásával, majd pontszám formájában visszajelzést kapnak arról, hogy milyen pontosan teljesítették a feladatot. Az eredeti játékhoz képest, amelyben a játékosok először csak egyetlen villanást látnak, majd kettőt, majd hármat stb., kísérletünkben a résztvevők mindig egy teljes szekvenciát látnak és idéznek fel.

A résztvevők átadási láncolatokba szerveződnek, ahol az n -edik résztvevő válaszait az $n + 1$ -ediknek kell lemásolnia (a szekvenciák sorrendje véletlenszerű). Minden résztvevő egy 60 elemből álló szekvenciahalmazt lát. A halmazt kétszer, két részletben mutatjuk meg, a szekvenciákat mindkétszer véletlenszerűen rendezve (a résztvevők nincsenek tudatában annak, hogy két részlet van, csak az egymást szünet nélkül követő szekvenciákat látják). A kiértékeléskor a második blokk válaszait rögzítjük és dolgozzuk fel. Minden láncolat első résztvevőjének egy olyan véletlenszerűen összerakott halmazt kell lemásolnia, amelyben egy szekvencia 12 villanásból áll, és mind a négy szín egyenletes el-

³⁷ Hannah CORNISH, Kenny SMITH és Simon KIRBY, „Systems from sequences: An iterated learning account of the emergence of systematic structure in a non-linguistic task”, in *Proceedings of the Annual Meeting of the Cognitive Science Society* 35 (2013): 340–345.

oszlással szerepel benne. A kezdeti véletlenszerű halmaz minden láncolatban más. A kiinduló szekvenciák hosszúsága azért 12 elem, hogy ne legyen se túl könnyű, se túl nehéz megjegyezni: szándékunk szerint a szekvenciák lemásolása elég nehéz ahhoz, hogy hibát is vétsenek a résztvevők, de nem túl nehéz ahhoz, hogy ne legyen reprodukálható. A lemásolt szekvenciák ezután a láncolat következő résztvevőjéhez kerülnek. Ha egy résztvevő 6 elemnél rövidebb szekvenciát hoz létre, akkor a forrásszekvenciát újra megmutatjuk az első blokk végén.

Minden láncolat 10 generáción keresztül folytatódik, azaz 660 szekvenciát elemzünk láncolatonként (egy véletlenszerű, 60 elemből álló kezdőhalmazt és 10 darab szintén 60 elemből álló sorozatot, amelyet a résztvevők generálnak). Hat ilyen láncolat eredményeit vizsgáltuk, mindegyikre ugyanazt a beállítást alkalmazva.³⁸ Ezáltal 3960 szekvenciát kapunk, generációk és láncolatok szerint rendezve (60 szekvencia $\times$ 10 generáció $\times$ 6 láncolat + 6 kezdeti 60 elemű halmaz, összesen 360 szekvencia). Megjegyzendő, hogy a feladat elején a résztvevők nem kapnak utasítást arra, hogy mintákat másoljanak, vagy hogy figyeljenek a halmaz egészének jellemzőire. A szekvenciákban felmerülő bármely szisztematikus rendszer tehát egy implicit tanulási torzítás eredményeként jöhet csak létre, amely hatással van a halmazok egészének kialakítására is. Ahogyan már említettük, a feladatot nehéznek szántuk, főleg a láncolatok legelején: azt várjuk, hogy a résztvevők sok hibát vétsenek a szekvenciák reprodukálásakor, különösen a korai generációk esetében. Kérdés, hogy a hibák elvezetnek-e egy strukturáltabb (és tanulhatóbb) rendszerhez.

EREDMÉNYEK

A teljes adathalmaz, illetve az ábrák létrehozásához és statisztikai elemzéséhez szükséges kódok elérhetők az alábbi repozitóriumban: https://github.com/smkirby/cultural_evolution_of_statistical_structure. Az 1. ábra a hibaarány generációnkénti csökkenését mutatja. A hibaérték kiszámítása a forrásszekvencia és a résztvevő által előállított szekvencia közötti normalizált Levenshtein-távolság felhasználásával történt – az 1-es érték teljes eltérést jelent a forrásszekvenciától, míg a 0-ás érték tökéletes egyezést jelöl. Egy lineáris ve-

³⁸ Az első négy láncolat elemzésének eredményei: CORNISH, SMITH és KIRBY, „Systems from sequences...”; a másik kettő eredményei: Simon KIRBY, Hannah CORNISH és Kenny SMITH, „Systems emerge: The cultural evolution of interdependent sequential behaviours in the lab”, in *Evolution of Language: Proceedings of the 10th International Conference* (2014): 463–464, https://doi.org/10.1142/9789814603638_0095.

gyes hatású modellt³⁹ alkalmaztunk a hibák alakulásának előrejelzésére. A modell figyelembe vette a generáció hatását mint rögzített hatást,⁴⁰ a láncok közötti különbségeket mint véletlen eltéréseket,⁴¹ generáció hatásának láncolatonkénti eltérését,⁴² az eredeti szekvencia hatását (a 0. generációban) mint véletlen eltérést.⁴³ (Próbálkoztunk azzal is, hogy az eredeti szekvencia hatását láncenként külön-külön modellezzük, de ez a modell matematikailag instabillá vált [ún. „singular fit”], ezért leegyszerűsítettük úgy, hogy véletlen interceptet használtunk.)⁴⁴ Generációnkénti szignifikáns hibaérték-csökkenést észleltünk ($\beta = -0.023$, $SD = 0.0032$, $p < 0.001$). Más szóval, a szekvenciák aszerint változnak, hogy a résztvevők egyszerűbben tudják lemásolni azokat. Ez megszokott eredménye az ismétléses tanulási kísérleteknek, amelyekben pusztán a generációk közötti átadásnak köszönhetően tanulhatóbakká válnak a továbbadott minták. A 2. ábra egy láncolat egy példaszekvenciáját és annak generációnkénti változását mutatja (emellett bemutatja az alább kifejtett szegmentációs módszerrel azonosítható alegységek kialakulását is). A könnyebben másolható szekvenciák nagy valószínűséggel maradandóbbak lesznek, míg a nehezebben másolhatók változni fognak a hibák miatt. Ennek köszönhetően a szekvenciák az idő elteltével egyre könnyebben másolhatókká válnak.

³⁹ A lineáris vegyes hatású modellek olyan statisztikai módszerek, amelyek lehetővé teszik, hogy egy elemzés egyszerre vegye figyelembe az általános törvényszerűségeket (például egy változó átlagos hatását) és az egyedi eltéréseket (például azt, hogy egyes csoportok vagy megfigyelési egységek eltérően viselkednek az átlagtól). A modell tulajdonképpen egy egyenletként értelmezhető, amelyben a különféle – rögzített és véletlenszerű – hatások szerepelnek. Ez az egyenlet segít megbecsülni, hogyan befolyásolják ezek a tényezők az empirikusan megfigyelt eredményeket. Jelen esetben például azt, hogy a generációk egymásra gyakorolt befolyása-e a résztvevők által vétett hibák számát, ha erre potenciálisan más tényezők is hatnak – *a szerk.*

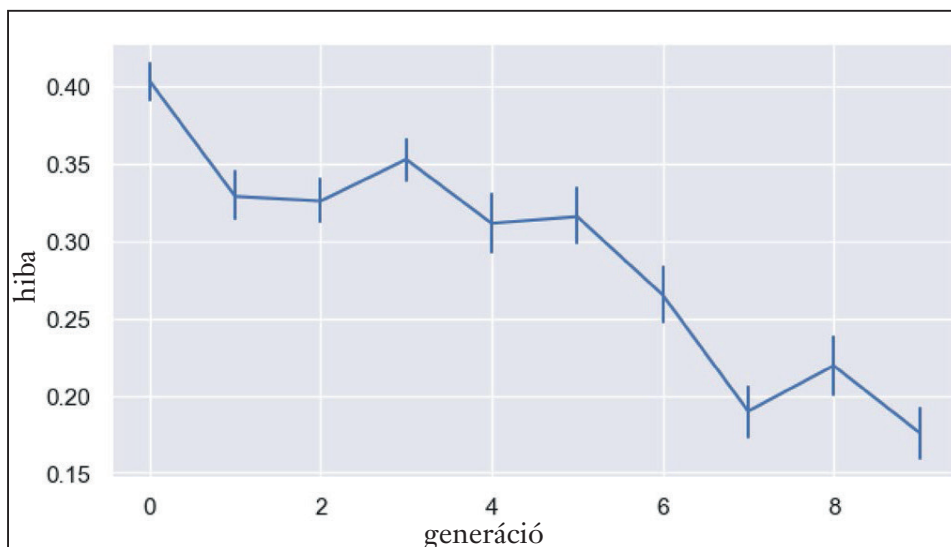
⁴⁰ Ez annak vizsgálatát jelenti, hogy hogyan befolyásolja a generációk száma a hibák mértékét – *a szerk.*

⁴¹ A modell megengedi, hogy minden lánc saját, kissé eltérő kiinduló hibaszinttel rendelkezzen – *a szerk.*

⁴² Tehát, hogy a generációk hatása ne legyen egyforma minden láncban – *a szerk.*

⁴³ Vagyis azt, hogy a kiindulópontként használt szekvencia önmagában is befolyásolhatja a hibákat – *a szerk.*

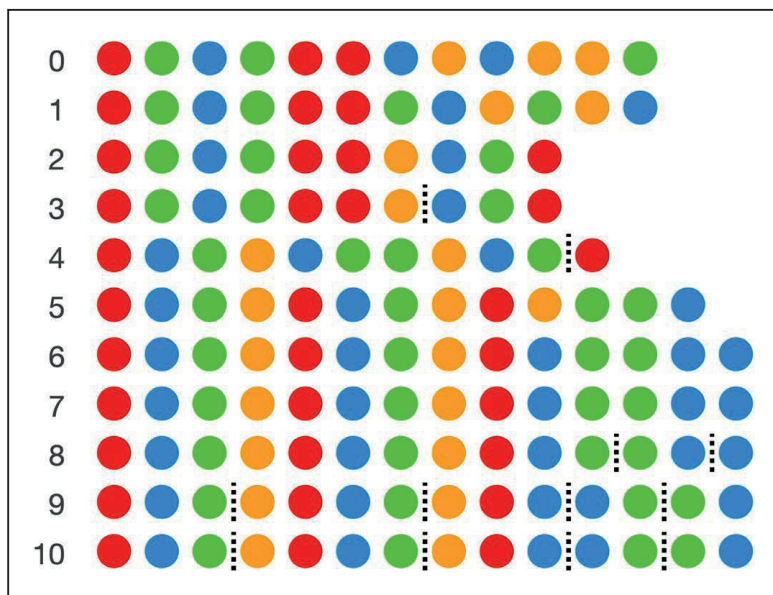
⁴⁴ Ez azt jelenti, hogy nem modellezték részletesen, hogyan viselkedik minden szekvencia minden láncban külön-külön, hanem csak azt vették figyelembe, hogy már az elején (a 0. generációban) lehettek különbségek a szekvenciák között abból a szempontból, hogy mennyire voltak hajlamosak hibákra – *a szerk.*



1. ábra. A hibaarány változása a kísérletben vizsgált generációk között. A hibaértéket átlagolva adjuk meg az egyes generációk különböző láncolataiban (a hibaérték minden szekvenciára külön-külön kiszámítva, majd az adott halmazhoz tartozó szekvenciák hibaértékei átlagolva, végül a hat láncolat értékei átlagolva). A hibasávok 95%-os konfidenciaintervallumokat jelölnek, amelyeket bootstrap (újramintavételezésen alapuló) módszerrel számoltunk ki. Jól láthatóan csökkenő tendencia figyelhető meg a hibaértékekben, ami azt mutatja, hogy a szekvenciahalmazok a későbbi generációkban könnyebben másolhatók a korábbi generációkhoz képest. A szekvenciák tehát tanulhatóbbakká váltak.

A SZEGMENTÁCIÓ KIALAKULÁSA

A szekvenciákat annak megfigyelésére is felhasználtuk, hogy vajon megjelennek-e bennük alegységek egyik generációról a másikra. Ne feledjük, hogy a statisztikai tanulás elméletei szerint a szótagok közötti következési valószínűség a szóhatárok jelzéseként funkcionál. Akkor történik szegmentálás, amikor csökken egy szótagszekvencia következési valószínűségének értéke – ez a csökkenés jelzi a szóhatárok lehetőségét, hiszen a szavakon belül a következési valószínűség általában nagyobb, mint a szavak között. Mivel a nyelvnek visszatérő szavai vannak, a szavakon belüli követkeзések kevésbé váratlanok, mint a szavak közöttiek. Felvethetjük, hogy hasonló mintázatok jelennek meg a mi adatbázisunkban is. A „Simon” játékban négy szín képezi az alapeleme-



2. ábra. Mi történik a kísérlet során egy meghatározott szekvenciával (a 60-ból) egy meghatározott láncolatban (a 6-ból)? A számértékek a kísérletben vizsgált generációkra vonatkoznak, a 0. generáció a kezdeti véletlenszerű sorozatot jelöli, az utána következő generációk pedig azok a szekvenciák, amelyeket a résztvevők generáltak.

Kezdetben a véletlenszerű szekvenciákat nehéz lemásolni, ezért több változtatás figyelhető meg, de generációk után a változtatások könnyebbé teszik a szekvenciák pontos ismétlését. A szaggatott vonalak az alább bemutatott szegmentációs módszerünkkel azonosított részek (amelyeken belül kifejezetten alacsony a következősi valószínűség) közötti határokat mutatják. Annak mértéke, hogy egy ehhez hasonló szekvencia könnyen vagy nehezen másolható, a halmaz többi 59 szekvenciájától is függ. Ugyanúgy, hogy mi számít szegmentációs határnak, az a teljes halmazra vonatkoztatott következősi valószínűség függvénye.

ket, a következősi valószínűség tehát az adott szín megjelenésének valószínűségét jelöli egy másik után. Ha létrejönnek alegységek, akkor bizonyos színek nagyobb valószínűséggel következnek egymásra, mint mások. Mindezek alapján a statisztikai tanulás szakirodalmára támaszkodva kidolgoztunk egy eljárást, amely képes azonosítani a következősi valószínűség csökkenését egy adott színszekvenciában, ami alegységek kialakulását jelezheti.

Az egymás utáni következősi valószínűségek hányadosán keresztül ugyanis számszerűsíthető a szegmentáció. Ehhez meg kell határoznunk magukat a valószínűségeket, amit két korábbi elem alapján teszünk meg, vagyis egy szín

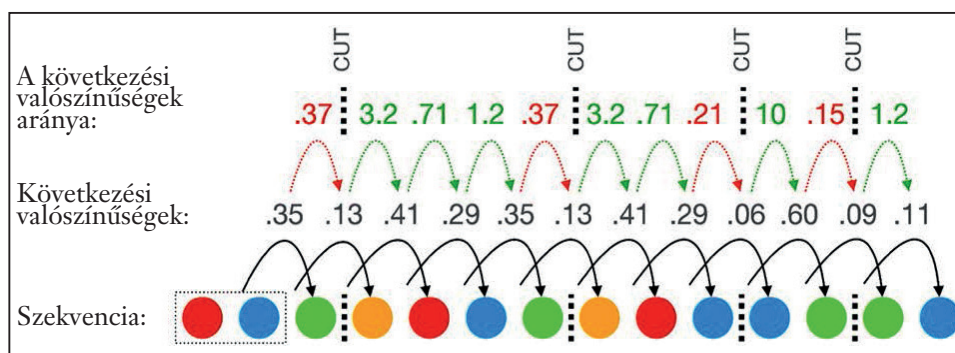
megjelenésének valószínűségét az azt megelőző két színhez viszonyítva számítjuk ki. Egy másik lehetőség az lett volna, ha csak az előző színt használjuk fel, de mivel csak négy színből áll a jelkészlet, ez az adatsor nagyon pontatlan becsléséhez vezetett volna. A teljes feltételes valószínűség elvét (vagyis az összes addigi színt a sorozatban) is használhattuk volna, ám az egyes generációkban rendelkezésre álló kis adatállomány (mindössze 60 szekvencia) miatt ez adatritkasági problémákhoz vezetne. Ezért a két megelőző szín figyelembevételével jelenthet kompromisszumot a két szélső érték között.

A következősi valószínűség legpontosabb becslést értékét egy adott generáció teljes, 60 elemből álló szekvenciahalmazában számítjuk ki egyetlen láncolat alapján. Tehát például a „piros-sárga-zöld-piros”, vagyis a „PSZP” szekvencia esetében megbecsüljük a „Z” elem következősi valószínűségét a „PS” után a teljes halmazhoz viszonyítva. Ehhez először megbecsüljük az „PSZ” elemcsoport előfordulási valószínűségét, majd elosztjuk a „PS” előfordulási valószínűségével. A valószínűségek mindig a teljes, 60 elemből álló szekvenciahalmaz függvényei. Aztán megbecsüljük a „P” következősi valószínűségét a „SZ” után. Ha az eredmény azt mutatja, hogy a „Z” nagy valószínűséggel következik „PS” után, de a „P” esetében ez nem mondható el, akkor jó okunk van rá, hogy a szekvenciát „PSZ–P” alegységekre szegmentáljuk.

Valahogyan el kellett döntenünk azt is, hogy a következősi valószínűség mikor elég alacsony ahhoz, hogy határt jelölhessünk ki. Egyik lehetőség, hogy a határt magára a következősi valószínűségre támaszkodva határozzuk meg, valahányszor egy bizonyos érték alá esik. Ezzel viszont nem lehet tetten érni a relatív csökkenést. Míg például egy 0,4-es következősi valószínűség nem túlságosan alacsony, egy 0,99-től 0,4-ig terjedő csökkenés szignifikáns lenne. Az ilyen típusú relatív csökkenések megragadására az egymást követő következősi valószínűségek közötti arányt (a két érték hányadosát) használjuk. Kiszámítjuk egy következősi valószínűség és az azt megelőző valószínűség arányát a szekvenciában (azaz a későbbit osztjuk el a korábbival), majd szegmentációs határt jelölünk ki, ha az arány szokatlanul alacsony (más szóval, amikor a következősi valószínűség hirtelen lecsökken). Az osztást követően 1 alatti értékek a következősi valószínűség csökkenését jelzik, míg az 1 fölötti arányok annak növekedését (vö. 3. ábra).

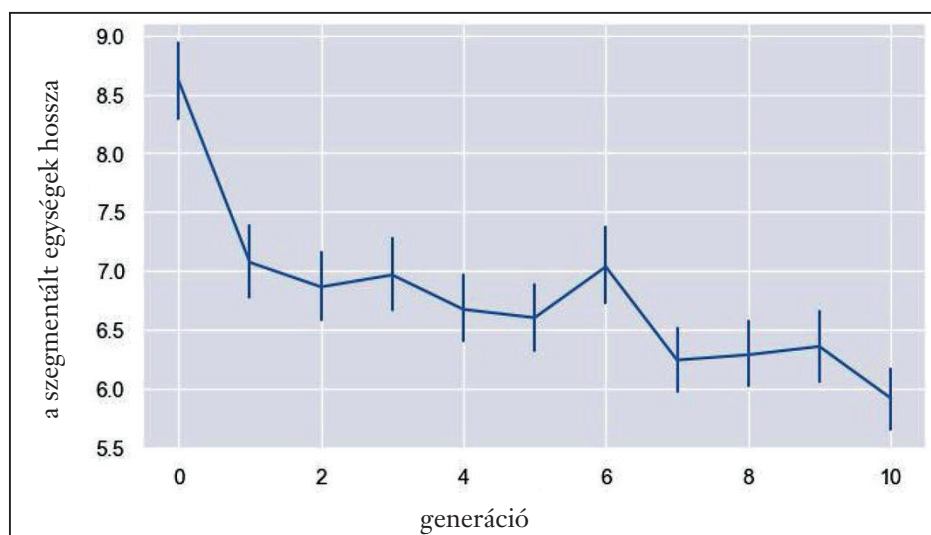
Ekkor már csak egy olyan módszerre van szükség, amellyel megítélhetjük, hogy egy következősi valószínűség-arány kellően alacsony-e ahhoz, hogy a sorozatot ezen a ponton meg lehessen szakítani. Ha túl alacsonyra szabjuk a küszöbértéket, akkor félő, hogy figyelmen kívül hagyjuk a tényleges szegmentációs határokat; ha pedig túl magasra, akkor szétválaszthatunk olyan ré-

szeket, amelyek inkább egybe tartoznának. Ennek megoldásaként a kezdeti véletlenszerű szekvenciákhoz fordultunk, mivel ezek nem rendelkeznek előre megadott struktúrával (azaz a színek egymásutánisága véletlenszerű és kiszámíthatatlan). Ezt a nemstrukturált halmazt használhatjuk annak megbecsülésére, hogy mennyire várható egy adott következési valószínűség csökkenése akkor, ha a szekvenciában nincs semmiféle belső rendszer. Ezáltal azonosíthatjuk a váratlanul nagy csökkenést jelző küszöbértéket. Ehhez kiszámítjuk a becsült következési valószínűség-arányok eloszlását mind a hat kezdeti, véletlenszerű sorozatban (minden láncolatból egyet), majd ezek összesített eloszlást hozzuk létre, amelynek 5%-os alsó régiója fogja képezni a küszöbértéket. Akkor szegmentálunk, ha a mért arány ennél (0,425) kisebb.



3. ábra. Példa a szekvenciák szegmentálására szolgáló eljárásra. A példa a kísérlet utolsó generációjának egyik szekvenciáját ábrázolja a megállapított szegmentációs határokkal. A szekvencia fölött azokat az értékeket tüntetjük fel, amelyek megmutatják, hogy egy adott szín milyen valószínűséggel következik az előtte lévő két szín után. Ezek az a láncolat teljes, 60 szekvenciából álló halmazára vonatkoztatva kerülnek kiszámításra. Egységhatárt tettünk oda, ahol jelentős csökkenés figyelhető meg a szekvencia következési valószínűségeiben. Ezt az egymásra következő valószínűségek kirívóan alacsony hányadosaként operacionalizáljuk. Az arányok (hányadosok) az ábra felső sorában találhatók. Például a következési valószínűség csökkenése 0,35-ről 0,13-ra kisebb a 0,425-ös küszöbértékünkénél, ami a szekvencia első zöld és első sárga egysége közötti határra utal. Piros színnel jelöltük, ha a következési valószínűségek hányadosa kisebb a küszöbértékünkénél, zölddel pedig, ha annál nagyobb. Megjegyzendő, hogy módszerünk egyik korlátja, hogy sosem helyezhetünk el határt a szekvencia harmadik színét megelőzve (mert a következési valószínűség az előző elempárra, míg az arányok mindig két egymás utáni következési valószínűségre épülnek).

Ezt az eljárást követve minden generáció esetében létrehozhatunk egy szegmentált részekből álló leltárat (lényegében egy lexikont). A 4. ábra bemutatja, hogyan csökken a szegmentált egységek hossza generációnként, azaz, hogy a szekvenciák egyre kisebb részekből tevődnek össze. Egy lineáris vegyes hatású modellt alkalmaztunk az egység-hossz alakulásának előrejelzésére. A modell tartalmazza a generáció hatását mint rögzített hatást,⁴⁵ valamint egy véletlen interceptet a láncokra,⁴⁶ és azt is megengedi, hogy a generáció hatása lánconként eltérő legyen (véletlen meredekség lánconként). Ez alapján szignifikáns generációnkénti csökkenés mutatható ki az egység-hosszokban ($\beta = -0.17$, $sd = 0.044$, $p = 0.012$). Ez még akkor is fennáll, ha nem vesszük számításba a kezdeti véletlenszerű szekvenciasorozatot ($\beta = -0.11$, $sd = 0.040$, $p = 0.041$).



4. ábra. Generációnkénti egység-hossz-változás (azaz az eljárásunk által szegmentált részek hosszának változása). A függőleges koordináták az összes egység átlagos hosszát jelölik egy meghatározott generációban, az összes láncolatra átlagolva. Az azonos hosszúságú egységeket többször is számítjuk az átlagolás során. A hibásávok 95%-os konfidenciaintervallumokat jelölnek, amelyeket bootstrap módszerrel számoltunk ki. Az egységek rövidebbek lesznek, ami generációnként egyre kisebb részek kiválását mutatja a nagyobb egészről.

⁴⁵ A modell azt vizsgálja, hogy általánosságban a generációk száma hogyan befolyásolja az egységek hosszát – *a szerk.*

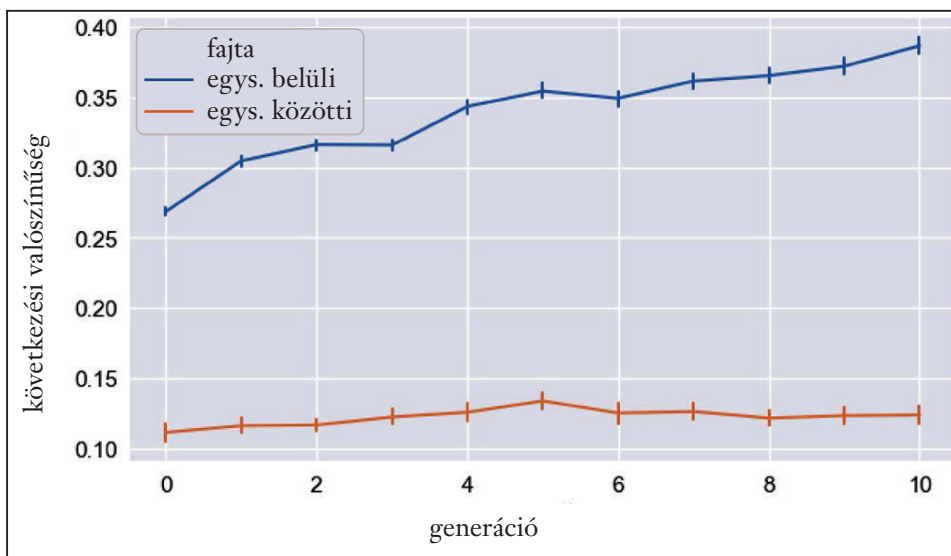
⁴⁶ Minden egyes lánc eltérő hosszúságú kiindulóegységekből állhat – *a szerk.*

Fontos továbbá, hogy az egységeken belüli következési valószínűségek generációnként nőnek, míg az egységek közöttiek kifejezetten alacsonyak maradnak. Ez azt jelenti, hogy a szekvenciahalmazok a kulturális evolúció során formálódnak, egyre egyértelműbb szegmentációs jelzéseket létrehozva (vö. 5. ábra). A következőképpen számoltuk ki ezt a két értéket: miután az eljárásunk szerint mindegyik halmazban szegmentáltuk a szekvenciákat, kategorizáltuk a következési valószínűségeket egységeken belüliekre és egységek közöttiekre. Például, ha szegmentáltuk a „PSZP”-t „PSZ–P”-re, akkor a „Z” következési valószínűsége „PS” után egyazon egységen belülinek számít, míg a „P” következési valószínűsége „SZ” után egységek közöttinek. Mivel a szekvenciákat a következési valószínűség csökkenésének függvényében szegmentáljuk, az egységek közötti követkeзések szükségszerűen kis valószínűségűek. Ezért fontosabb az, hogy az egységeken belüli következési valószínűségek viszont egyértelmű növekedést mutatnak, ami a szegmentált egységek statisztikai kohéziójának növekedését jelzi.

Ez esetben is egy lineáris vegyes hatású modellt alkalmaztunk a következési valószínűségek alakulásának becslésére. A modell rögzített hatásként tartalmazza a generációt, a következési valószínűség típusát (egységeken belül vagy között), valamint ezek kölcsönhatását. Emellett tartalmaz egy véletlen interceptet a láncokra,⁴⁷ és lehetővé teszi, hogy a generáció hatása lánconként eltérő legyen (véletlen meredekség). Az eredmény szignifikáns kölcsönhatást mutat a következési valószínűség típusa és a generációk között, azaz a generációnkénti valószínűségnövekedés nagyobb az egységeken belül, mint azok között. ($\beta = 0.010$, $sd = 0.00070$, $p < 0.001$).

Mindezek alapján azt mondhatjuk, hogy pusztán a szekvenciák generációnkénti átadásával megjelenik egy rendszerszintű statisztikai rendszer – annak ellenére, hogy a feladat nem egy struktúra megtanulása volt. A rendszer egyre erősebb szegmentációs jeleket fejleszt ki, ami az egységen belüli és az egységek közötti átmeneti valószínűségek növekvő különbségében mutatkozik meg. Más szavakkal, a részek az egészekből emelkednek ki.

⁴⁷ A láncok kiindulóegységeiben eltérők lehetnek a következési valószínűségek – *a szerk.*



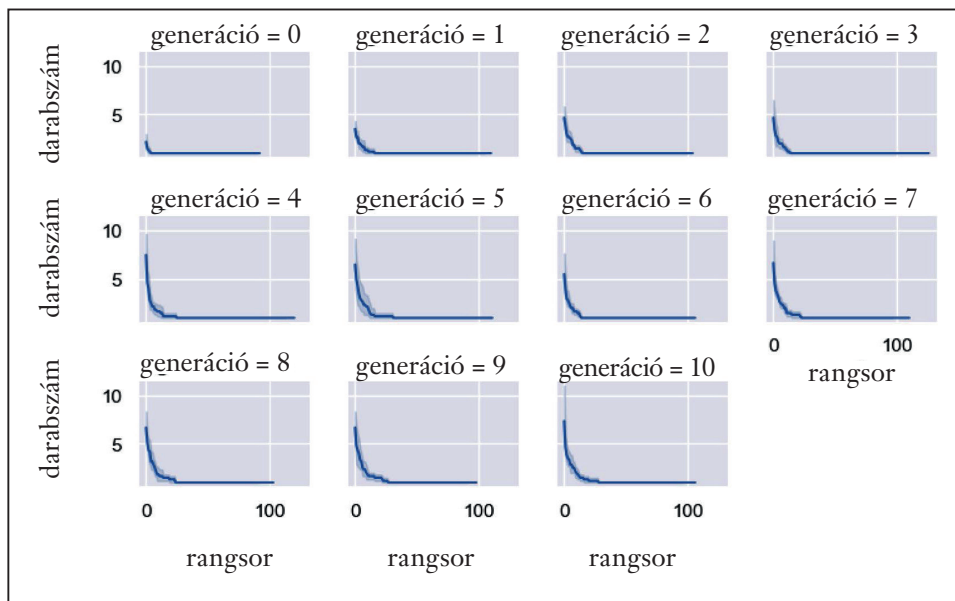
5. ábra. A következési valószínűség generációnkénti változása egységeken belül és egységek között. A hibasávok 95%-os konfidencia-intervallumokat jelölnek, amelyeket bootstrap módszerrel számoltunk ki. A szegmentációs módszerünk alkalmazása után mindegyik következési valószínűséget osztályozhatjuk aszerint, hogy egy egységen belül vagy két egység között helyezkedik el. A grafikon minden pontja azt az átlagot mutatja (egy generáció mindegyik láncolatára kiterjesztve), amely az egységen belüli és egységek közötti következési valószínűségekből lett kiszámítva.

A határok megállapításának módja miatt szükségszerűen kisebbek lesznek az egységek közötti következési valószínűségek, még a nulladik generációban is. Itt inkább az a lényeg, hogy az egységeken belüli következési valószínűségek növekednek.

A STATISZTIKAI STRUKTÚRA KIALAKULÁSA

Most, hogy kísérletünknek köszönhetően rendelkezésünkre áll minden generáció és láncolat szegmentált egységeinek leltára, feltehetjük a kérdést: hogyan változik az egységek gyakorisági eloszlása, amikor a szekvencialalmazokat ismétléses tanulás útján továbbítják? A kérdés megválaszolására a leltár összes egységének gyakoriságát kiszámítjuk minden generációra és láncolatra, majd a gyakoriságot az a gyakorisági rangsorban elfoglalt hely függvényében ábrázoljuk (vö. 6. ábra). Ez azt mutatja, hogy egy viszonylag konstans kezdeti állapothoz képest (majdnem az összes elem csak egyszer jelenik meg a kezdeti, véletlenszerű halmazban), egyre inkább ferde eloszlás alakul ki.

A leggyakoribb egységek 5–10-szer többször jelennek meg, mint a legkevésbé gyakoriak. Ferde eloszlás jön létre tehát pusztán a szekvenciák generációs átadásának eredményeképpen, az ismétléses tanulás során.

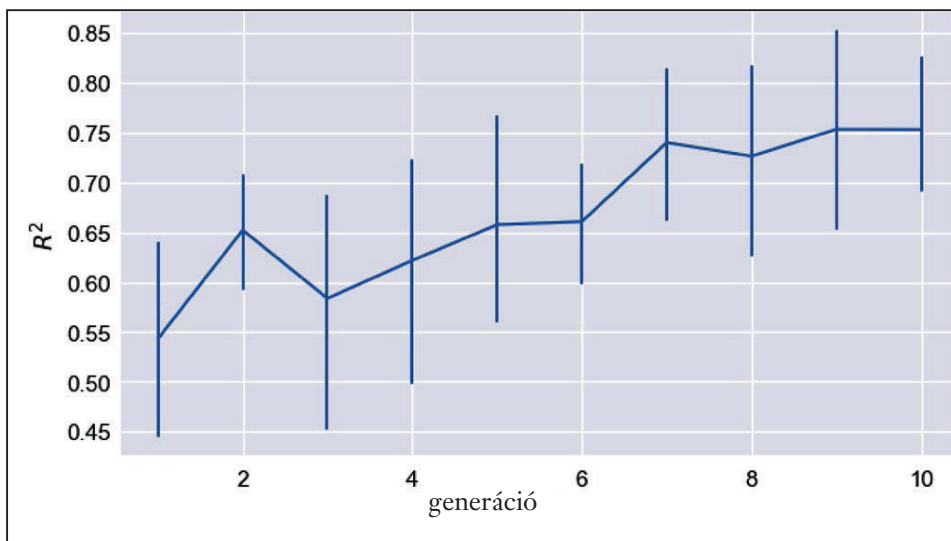


6. ábra. Az elemek eloszlása generációnként a szekvenciahalmazokban. A gyakoriságokat a gyakorisági rangsorban betöltött hely viszonyában ábrázoltuk. A láncolatok bootstrap módszerrel kiszámolt 95%-os konfidenciaintervallumai szürke mezőként láthatók az átlagok körül. A kezdeti viszonylag egységes eloszlásból egy ferde eloszlás alakul ki, kevés magas gyakoriságú és sok alacsony gyakoriságú elemmel.

Feltehetjük a kérdést, hogy az eloszlás mennyire áll közel a Zipf-féléhez, és hogy hogyan változik generációról generációra. Ehhez meg kell néznünk a gyakoriság és a rangsorban elfoglalt hely logaritmusának R^2 -ét.⁴⁸ Minden generáció minden láncolatához kiszámoltuk az R^2 -et. Látható, hogy a Zipf-eloszláshoz való illeszkedés generációnként nő (vö. 7. ábra). Ezt statisztikailag is alátámasztja egy lineáris vegyes hatású modell, amely az R^2 értéket jósolja

⁴⁸ Ha a gyakoriságokat és a rangsorban elfoglalt helyeket logaritmikus skálán ábrázoljuk, akkor a Zipf-eloszlás egy egyenes vonalat ad ki. Ezért egy egyenes vonalat illesztve az adatokra (lineáris regresszióval), kiszámítható, hogy az adatok mennyire szorosan követik ezt az egyenest. Ez az, amit az R^2 mutat meg: ha közel van az 1-hez, akkor az adatok szorosan követik a Zipf-féle mintázatot – *a szerk.*

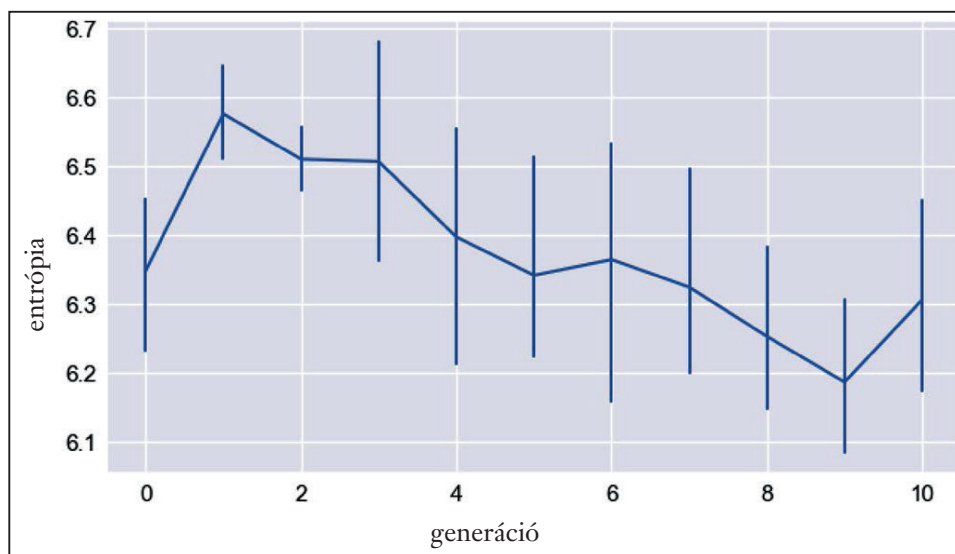
meg a generációk száma alapján. A modell tartalmaz egy rögzített hatást a generációra, valamint figyelembe veszi, hogy a láncok kiinduló szintjei különbözhetnek egymástól, és a generáció hatása is eltérhet lánconként. Ezek alapján az R^2 szignifikáns növekedését tapasztaltuk a generációk viszonylatában ($\beta = 0.022$, $sd = 0.0066$, $p = 0.020$).



7. ábra Az egységek logaritmizált gyakorisága és a logaritmizált rangsora közötti korreláció R^2 értékei (lásd a 6. ábrát). Az ábrán ezeket mutatjuk a generációk függvényében (a 0. generációt, azaz a kezdeti, véletlenszerű nyelvet kihagytuk, mivel több láncolat esetében az eloszlás teljesen lapos volt, így a korrelációs együttható nem számítható ki értelmesen). A hibasávok 95%-os konfidenciaintervallumokat jelölnek, amelyeket bootstrap módszerrel számoltunk ki. Idővel a szóeloszlások egyre inkább követik a Zipf-eloszlást. Ezt alátámasztja a növekvő lineáris kapcsolat a logaritmikus gyakoriság és a logaritmikus rangsor között.

Még jobban tudjuk számszerűsíteni a ferde eloszlás mértékét, ha a szekvencialmazok egységeinek entrópiáját számítjuk ki. Egyenletes eloszlás esetében ugyanis nagyobb az entrópia, mint a ferde eloszlásnál. Az entrópia generációnkénti csökkenését mutatja, még akkor is, ha figyelembe vesszük, hogy hány különböző egységet azonosított a szegmentálási eljárásunk. Statisztikailag ezt újfent egy lineáris vegyes hatású modellel teszteltük, amely az entrópiát jósolja a generációk és az eltérő egységek száma alapján. A modell

tartalmaz rögzített hatásokat a generációra és az egységek számára vonatkozóan, valamint számol a különböző láncok kiinduló szintjei közötti véletlenszerű eltérésekkel. Az egységek számát azért vontuk be, hogy biztosítsuk: a generáció hatása nem csak azért jelentkezik, mert kevesebb egység alacsonyabb entrópiát eredményez. (Egy bonyolultabb modell – amely lehetővé tenné, hogy a generáció hatása láncenként eltérjen – nem volt stabil, ezért egyszerűsített modellt használtunk.) Az eredmények szerint generációként szignifikánsan csökken az entrópia ($\beta = -0.033$, $sd = 0.0055$, $p < 0.0001$). Emellett, ahogyan várható volt, az egységek számának növekedésével az entrópia is nő ($\beta = 0.0028$, $sd = 0.00084$, $p = 0.0013$).



8. ábra. Az egységek entrópiája generációként ábrázolva. Az entrópiát mindegyik generációra és mindegyik láncolatra kiszámoltuk, majd az értékeket átlagoltuk.

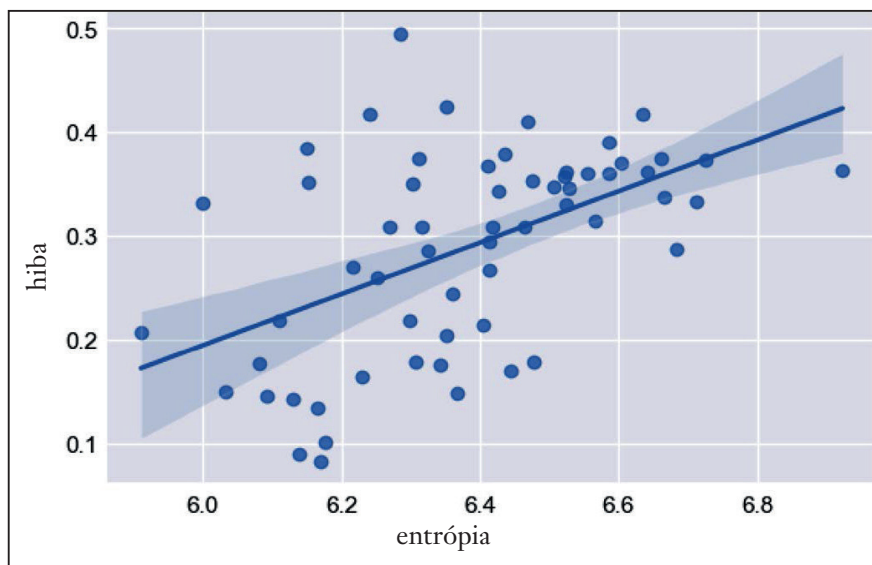
A hibásávok 95%-os konfidenciaintervallumokat jelölnek, amelyeket bootstrap módszerrel számoltunk ki. Megfigyelhető egy kezdeti entrópiánövekedés a halmaz méretének növekedése miatt (mivel a kezdeti véletlenszerű szekvenciahalmaznak nagyon kevés szegmentálható egysége van), de az egyre ferdébb eloszlásnak köszönhetően az egységek entrópiája is csökken generációról generációra.

A TANULHATÓSÁG MINT OK ÉS OKOZAT

Mindent összevetve az eredmények azt mutatják, hogy az ismétléses tanulás hozza létre a nyelv két szisztematikus sajátosságát: a részek elkülönülését és azok ferde eloszlását.

Bemutattuk, hogy ezeket a sajátosságokat az átadások eredményezik. De vajon be lehet-e bizonyítani, hogy a kialakuló statisztikai struktúra ténylegesen hozzájárul a rendszer tanulhatóságához? A 9. ábra azt mutatja, hogy milyen kapcsolat tételezhető a résztvevő által látott szókészlet entrópiája és az adott résztvevő által produkált hibaarány között. Eszerint világos kapcsolat áll fenn a statisztikai struktúra és a tanulhatóság között; a Person-féle korreláció értelmében a hiba és az entrópia pozitívan korrelál ($r = 0.52$, $p < 0.001$).

Az ismétléses tanulás során létre jövő statisztikai struktúra lehetővé teszi, hogy a megképződő rendszer egyre inkább tanulható legyen. A részek és azok ferde eloszlása tehát egyaránt oka a tanulhatóság növekedésének, és okozata a generációkon átívelő ismétléses tanulásnak.



9. ábra A hiba és az entrópia kapcsolata. Minden pont egy-egy résztvevőt jelöl; illetve ábrázoltuk az ezekre legjobban illeszkedő egyenest és a 95%-os konfidencia-intervallumot is. A kevésbé ferde eloszlású egységeket (azaz a nagyobb entrópiával rendelkező szekvenciahalmazokat) nehezebb megtanulni (azaz nagyobb hibaszámmal másolják le azokat).

AZ EREDMÉNYEK ÉRTELMEZÉSE

A nyelv kifejezőerejét – fajunk egyik legkülönlegesebb jellemvonását – kis részek nagyobb egészekké való kombinációja teszi lehetővé. A természet más kommunikációs rendszereihez képest a nyelv kombinálható részei azonban nem eleve adottak, hanem a gyermekek hallás után sajátítják el: ahhoz, hogy felfedezzék a részeket, el kell valahogyan különíteniük azokat egy tagolatlan beszédflowban. A részek nem egyforma eloszlásúak, hanem egy jellegzetes hatványtörvény szerinti eloszlást követnek kevés gyakori és sok ritka szóval (a gyakoriságok nem lineárisan csökkennek).

A tanulmány bemutatta, hogy a nyelv mindkét sajátossága – az, hogy részekből áll, amelyek meghatározott eloszlásúak – olyan kulturális átadás során is kialakulhat, amikor a tanulók egyszerűen sorozatok egészét reprodukálják. Egy nemnyelvi feladat segítségével kimutattuk, hogy a színszekvenciák az idő elteltével strukturáltabbá és tanulhatóbbá válnak. A strukturáltság és tanulhatóság növekedése a Zipf-eloszláshoz egyre jobban illeszkedő nyelvi részek kialakulásából adódik. Az elkülönülő részek meghatározása a színek közötti következtési valószínűségek felhasználásával történt – ugyanazzal a statisztikai mutatóval, amellyel a csecsemők is felismerik a szóhatárokat. Az eredmények szerint generációnként nőnek az egységeken belüli következtési valószínűségek; az eloszlásuk entrópiája időről időre csökken, és egyre inkább a Zipf-eloszlást követik; illetve az egységek entrópiája a hibaértékkel korrelál. Azt állítjuk, hogy a nyelv pontosan azért áll statisztikailag elkülöníthető részekből, amelyek egyre ferdebb eloszlásúak, mert a nyelvhasználók több generáción keresztül tanulják meg azt. A nyelv statisztikai struktúrája egyszerre oka és okozata a nagyobb tanulhatóságnak.

Hogyan jelenik meg mindez kutatásunkban? Eredményeink nem közvetlenül erre vonatkoznak, de felvázolnak egy lehetőséget. Kísérleti feladatunkban a résztvevők olyan szekvenciákat másolnak le, amelyeknek eleinte teljesen egynemű statisztikai struktúrája van: bármely szín követhet bármely más színt. A mintázatok másolásakor a résztvevők mégis bizonyos színeket gyakrabban párosítanak össze, mint másokat. Ez kezdetben pusztán hiba. Viszont emiatt a tanulók következő generációja számára a szekvenciák már statisztikai szempontból nem egyneműek. Így alakulnak ki a következtési valószínűség mintái, amelyekben az alacsony valószínűségeket olykor nagyok követik. A tanulók ezt az esetet a szekvencia alrészei közötti határvonal jelzésének tekintik. Ily módon jön létre a lehetséges egységek leltára, amelynek száma idővel növekedhet, egyre nagyobb különbségeket hozva létre az egységek közötti és belüli következtési valószínűségekben. Ez a gyakorisági előny generációról

generációra erősödik, létrehozva a nyelvtanulást elősegítő ferde eloszlást. Ebben a folyamatban „a gazdagabb még gazdagabb lesz” (*rich-gets-richer*), és az így kialakuló ferde eloszlás a következő valószínűségek hasznosíthatóságát növeli a határkijelölésben, ami az ismétlődő egységek még gyakoribb ismétléséhez, egyre ferdebb eloszláshoz, illetve pontosabb másolásához vezet. A folyamat ismétlődik és erősödik, az idő elteltével növekszik a ferde eloszlás mértéke, valamint az eloszlás korrelációja a nyelvtanulás sikerességével. Így tehát, még ha a nyelvtanulók olyan feladattal is állnak szemben, amelyben teljes szekvenciákat kell lemásolniuk (anélkül, hogy számítanának arra, hogy a szekvenciák alrészekből állnak, amelyek ráadásul a különböző szekvenciákban visszatérnek), egy idő után elkezdenek kialakulni elkülöníthető részek. Ezek a részek később újrendezhetőek, hogy a generációkon átívelő továbbadás szempontjából optimalizált szekvenciahalmaz jöjjön létre.

Az egésztől a részekig haladó stratégia látszólag kontrainuitív ahhoz a verzióhoz képest, amikor először különálló egységeket tanulnak meg a nyelvhasználók, és ezeket kombinálják újra. Maga a környezet indokolja azonban e stratégia létezését: a részeket ugyanis nem szegmentált egységként ismerjük meg. A nyelv sokféle ehhez hasonló példával szolgál: a fonémák a teljes szóalakból vezethetőek le,⁴⁹ a morféma többmorfémás szavakból,⁵⁰ a szavak pedig többszavas egységekből.⁵¹ Kutatásunk fényében a részek létezésének másik oka, hogy az egésztől a részek felé haladó stratégia valóban több szempontból megkönnyíti a nyelvtanulást, és elősegíti a részek gyakran önkényes viszonyainak megtanulását. A névelő-főnév illeszkedése például jobban tanulható, ha először egyetlen egységekkel kezdődik a tanulás, amelyeket csak aztán bontanak elemeire a nyelvhasználók, és nem külön-külön a névelőt és a főnevet tanulják meg a későbbi kombinálás érdekében.⁵² Az egésztől a részek felé haladó tanulás az egyén fejlődésének más aspektusaiban is segíthet. A csecsemők kezdeti gyenge látásélessége például, amely megakadályozza

⁴⁹ Shelley L. VELLEMAN és Marilyn M. VIHMAN, „Whole-word phonology and templates”, *Language and Speech* 32 (1989): 149–170, [https://doi.org/10.1044/0161-1461\(2002/002\)](https://doi.org/10.1044/0161-1461(2002/002)).

⁵⁰ Dorit RAVID és Alon MALENKY, „Awareness of linear and nonlinear morphology in Hebrew: A developmental study”, *First Language* 21, 61. sz. (2001): 25–56, <https://doi.org/10.1177/014272370102106102>.

⁵¹ Isaac ARNON, „The Starting Big approach to language learning”, *Journal of Child Language* 48, 5. sz. (2021): 937–958, <https://doi.org/10.1017/S0305000921000386>.

⁵² ARNON és RAMSCAR, „Granularity and the acquisition of grammatical gender...”, 292–305; SIEGELMAN és ARNON, „The advantage of starting big...”, 60–75.

őket a részletek észlelésében, a jelek szerint az átfogó térbeli elemzést támogatja, ami többek között az arcfeldolgozás szempontjából is kedvező.⁵³

Tekintve, hogy kísérletünkbe tudatosan nemnyelvi feladatot építettünk be, a memóriajátéknak nincs közvetlen köze a nyelvhez. A színszekvencia nem hasonlít semmilyen beszélt vagy írott nyelvhez, és a feladat nem úgy volt keretezve, hogy a résztvevők ehhez hasonló mintázatokat keressenek vagy ismerjenek fel. Továbbá sem a teljes szekvenciáknak, sem a kialakuló részeknek nincsen jelentésük. Ennek ellenére rendszerszintű statisztikai struktúra jön létre, amely látványosan hasonlít a nyelv szerkezetére. Ez arra utal, hogy a nyelvben lévő rendszer is ugyanígy, a kulturális evolúció és az ismétléses tanulás általános folyamatán keresztül jött létre, semmint a nyelvspecifikus tanulási preferenciák által. Várhatóan mindenütt ferde eloszlású szekvenciákat találunk, ha ezek szintén a kulturális átadás révén alakulnak, és ha ezeket a tanulók először egészként sajátítják el, majd ezután fedezik fel bennük a részeket. Például a zene is ilyen kulturálisan átadott mintázat, ezért itt is a Zipf-eloszláshoz való illeszkedést feltételezhetjük. Erre van is bizonyíték: a zenei elemek – többek között a hangmagasság, a dallamfutamok és a konszonancia – pontosan követik a Zipf-eloszlást.⁵⁴ Az ének is kulturálisan átadott mintázat, és szintén kitett az egésztől a részek felé haladó tanulásnak, hiszen a dalokat gyakran a maguk teljességében tanuljuk és értjük meg. Itt is igaz, hogy a világ különböző kultúráiban a ritmus- és dallamelemek eloszlása a hatványtörvényt követi.⁵⁵

Messzebbre tekintve feltehetjük a kérdést, hogy más fajok kommunikációs rendszerében is találhatók-e hasonló ferde eloszlások. Ha a tanulhatóság egyszerre oka és okozata a ferde eloszlásnak, ez várhatóan kialakulhat más kulturálisan átadott rendszerek esetében is, főként ha azok rész-egész viszonyokból állnak. Gondolhatunk néhány olyan faj kommunikációs módjára, amelyek rendelkeznek ezekkel a sajátosságokkal, ilyenek például a delfinek, a bálnák vagy az énekesmadarak. Érdekes módon a felnőtt palackorrú delfi-

⁵³ Lina VOGELSANG, Jürgen SCHMIDHUBER, Jonathan C. FLACK, Michael P. TKACH és Sarah B. HRDY, „Potential downside of high initial visual acuity”, *Proceedings of the National Academy of Sciences* 115, 44. sz. (2018): 11333–11338, <https://doi.org/10.1073/pnas.1800901115>.

⁵⁴ Bill MANARIS, Chaoyi ROMERO, Panayotis P. PANAYIDES, Dan McNAB és Joan M. KISER, „Zipf’s law, music classification, and aesthetics”, *Computer Music Journal* 29, 1. sz. (2005): 55–69, <https://doi.org/10.1162/comj.2005.29.1.55>.

⁵⁵ Samuel A. MEHR, Manvir SINGH, Dean KNOX, Daniel T. BOWLING, Cody M. MOSER és Patrick E. SAVAGE, „Universality and diversity in human song”, *Science* 366, 6468. sz. (2019): eaax0868, <https://doi.org/10.1126/science.aax0868>.

nek füttyrepertoárja követni látszik a hatványtörvényt,⁵⁶ ahogyan a törpeposzáta (egyfajta énekesmadár) dalrepertoárjának szótagjai is.⁵⁷ Ugyanígy a hosszúszárnýú bálna énekének elemeiről is kimutatták, hogy a nagyban követik a ferde eloszlást.⁵⁸ Ezek az eredmények arra utalnak, hogy a nem emberi fajok esetében is az emberi kommunikációhoz hasonló folyamatok mennek végbe. Ebből kiindulva további kutatások pontosabban feltárhatják majd a ferde eloszlás nyelvtanulást elősegítő hatását a nem emberi fajok esetében is.

Kutatásunk több különböző hagyományt fog össze a kulturális evolúció, a statisztikai tanulás, a fejlődépszichológia és a nyelvészet területeiről. Bemutattuk, hogyan lehet ezeket együttesen alkalmazni egy olyan kísérlet létrehozásához, amely hozzájárul az emberi nyelv alapvető sajátosságainak és ezek eredetének megértéséhez. Ez nem jöhetett volna létre a különböző nézőpontok vegyítése nélkül – ilyen az ismétléses tanulás a kulturális evolúció területéről vagy a következási valószínűség elemzése a statisztikai tanulás területéről. Úgy gondoljuk, hogy mindez bizonyítja az interdiszciplinaritás szükségességét a nyelv eredetének és evolúciójának empirikusan megalapozott elméleti kutatásában.

ADATBÁZIS ELÉRHETŐSÉGE

A teljes adathalmaz, illetve az ábrák létrehozásához és statisztikai elemzéséhez szükséges kódok elérhetők az alábbi repozitóriumban: https://github.com/smkirby/cultural_evolution_of_statistical_structure.

⁵⁶ Brenda McCOWAN, Sarah F. HANSER és Laurel R. DOYLE, „Quantitative tools for comparing animal communication systems: Information theory applied to bottlenose dolphin whistle repertoires”, *Animal Behaviour* 57, 2. sz. (1999): 409–419, <https://doi.org/10.1006/anbe.1998.1000>; Ryo SUZUKI, James R. BUCK és Peter L. TYACK, „The use of Zipf’s law in animal communication analysis”, *Animal Behaviour* 69, 1. sz. (2005): F9–F17, <https://doi.org/10.1016/j.anbehav.2004.08.004>.

⁵⁷ Antonio M. PALMERO, Javier ESPELOSÍN, Paolo LAIOLO és Juan C. ILLERA, „Information theory reveals that individual birds do not alter song complexity when varying song length”, *Animal Behaviour* 87 (2014): 153–163, <https://doi.org/10.1016/j.anbehav.2013.10.026>.

⁵⁸ Jessica A. ALLEN, Eleanor C. GARLAND, Rochelle A. DUNLOP és Michael J. NOAD, „Network analysis reveals underlying syntactic features in a vocally learnt mammalian display, humpback whale song”, *Proceedings of the Royal Society B* 286, 1917. sz. (2019): 20192014, <https://doi.org/10.1098/rspb.2019.2014>.

A fordítás az alábbi szöveg alapján készült: IbnaL ARNON és Simon KIRBY.
„Cultural evolution creates the statistical structure of language”. *Scientific Reports* 14 (2024): 5255 (article number).

Fordította: *László István*