

Pozitív polaritású elemek a magyarban

A vagy és a többiek*

2. rész

2. Kísérlet

2.1. Módszer

2.1.1. Résztevők. A kísérletünkben részt vevő 36 alanynak (11 nő és 25 férfi, 20–43 év közöttiek) „klasszikus” PPE-ket, illetve diszjunkciót tartalmazó tagadó mondatok elfogadhatóságát kellett megítélniük. A résztvevőket a Prolific online toborzó platformon keresztül vontuk be a kísérletbe, és a kísérlet elvégzéséért egy általunk előre meghatározott, fix összegű juttatást kaptak. A jelentkezők személyes körülményeiből fakadó potenciális zavaró tényezők kizárása céljából az alábbi szűrő feltételeket alkalmaztuk a kísérleti alanyok bevonásához: a résztvevőknek Magyarországon született magyar állampolgároknak kellett lenniük, akik az adatfelvétel ideje alatt és azt megelőzően életvitelszerűen Magyarországon éltek. További elvárás volt, hogy a jelentkezők egynyelvű környezetben felnövő, nyelvi rendellenességekkel nem rendelkező, monokulturális egyénnek vallják magukat. A kísérletben való részvételhez elvárt minimum végzettségként a középiskolai érettségi vizsgát határoztuk meg. A résztvevőknek a Prolificre történt regisztrációját, valamint az oldal keretein belül meghirdetett és ismertetett kísérletre való jelentkezését a tájékoztatáson alapuló beleegyező nyilatkozat megadásának tekintettük. A résztvevőket tájékoztattuk a kísérlet céljáról és menetéről, továbbá arról is, hogy adataikat bizalmasan kezeljük, és kizárólag statisztikai elemzésben használjuk fel.

2.1.2. Eljárás. Az elfogadhatósági ítéletek mérését a Google Űrlapok platformon egy ötpontos Likert-típusú skála segítségével valósítottuk meg. A kísérleti feladat leírását a kérdőív elején, írásban mutattuk be a résztvevőknek. A kérdőív minden oldalának alsó részén egy megítélendő stimulus mondatot, a mondat felett, idézőjelek között pedig egy célértelmezést olvastak a résztvevők. A résztvevőknek a skála megfelelő elemének kiválasztásával azt kellett megjelölniük, hogy mennyire könnyen tudják a mondatot a felette látható idézőjeles parafrázis szerint érteni.

Minden oldalon először az idézőjeles célértelmezés jelent meg először, majd kis idő múlva a stimulus mondat is. A parafrázis és a mondat bemutatási sorrendjének megválasztását a mondat potenciális kétértelműsége motiválta. Ha a célértelmezés helyett a stimulus mondatot prezentáltuk volna előbb, nem lett volna kizárható, hogy a résztvevők az esetek legalább egy részében e potenciálisan két-

* Köszönettel tartozunk anonim bírálóinknak hasznos kérdéseikért és megjegyzéseikért. Tanulmányunk a PPKE-BTK-KUT-23-2 számú projekt keretében a Pázmány Péter Katolikus Egyetem Bölcsész- és Társadalomtudományi Kar támogatásával készült. A tanulmány 1. részét lásd MNy. 2024: 428–436.

értelmű mondatnak nem az idézőjeles parafrázis szerinti jelentését, hanem a másik logikailag lehetséges jelentését érik el először. Egy esetlegesen a célértelmezéstől eltérő kezdeti értelmezés önmagában megnehezíthette volna a mondatnak az idézőjeles parafrázis szerinti értelmezését, mely így voltaképpen a mondat első benyomást követő átértelmezését kívánta volna meg a résztvevőktől.⁷

A skála használatát, valamint a feladat megértését megkönnyítendő a feladatot bemutató instrukció részeként négy példát helyeztünk el, amelyek eltérőek voltak a tekintetben, hogy a magyar anyanyelvűek általában mennyire könnyen tudják az adott példamondatot a felette megjelenő parafrázis szerint érteni. Ezen illusztratív példák pusztán a feladat jobb megértését szolgálták: a résztvevőknek róluk még nem kellett ítéleteket hozniuk. A példák után – a korábbi szakirodalom alapján – szöveges visszajelzést adtunk arra vonatkozóan, hogy a magyar anyanyelvűek többsége mennyire könnyen tudja az adott példamondatot az idézőjeles értelmezésnek megfelelően érteni, továbbá, hogy ez az általános ítélet a skála mely elem(ei)nek feleltethető meg. A résztvevők a feladatleírás és az illusztratív példák után három gyakorló példán keresztül ismerkedhettek meg a feladattal, melyek után szintén visszajelzést adtunk a magyar anyanyelvűeknek a gyakorló mondatokkal kapcsolatos általános ítéleteiről. A gyakorlás fázisában a résztvevőknek a skála használatával már meg kellett ítélniük a mondatokat. A gyakorló példák és az őket követő kísérlet során az alanyok egyesével haladtak végig a stimulusokon.

2.1.3. Kísérleti anyag. Kísérletünk kritikus mondatai mondattagadást, valamint háromféle PPE egyikét (ezeket a továbbiakban PPE-típusnak nevezzük): vagy az *egy* névelővel bevezetett főnévi szintagmát, vagy *vala*-névmást, vagy pedig diszjunkciót tartalmazó összetett mondatok voltak. A PPE mindig a jelöletlennek tekinthető alany–ige–tárgy szórendet követő, tárgyi szerepű beágyazott mondaton belül, a tagmondat tárgyi mondatrészében jelent meg. A klasszikus PPE-ket tartalmazó kritikus mondatok esetében e tárgyi mondatrészi szerepben vagy egy az *egy* névelővel bevezetett határozatlan főnévi szintagma, vagy egy határozatlan *vala*-névmás állt, míg a diszjunktív kritikus mondatok két határozott, diszjunkció által mellérendelt tárgyi szerepű főnévi kifejezést tartalmaztak.

A tagadás PPE-khez viszonyított lokalitása alapján kétfajta kritikus mondat-típust alkottunk: az egyikben a tagadás is és a PPE is a pozitív polaritású főmondatba beágyazott tárgyi tagmondaton belül szerepelt, míg a másikban a PPE-k a tagadó főmondatba ágyazott pozitív polaritású tárgyi mellékmondatban helyezkedtek el. Az előbbi mondattípusban (melyet ez alapján LOKÁLIS kondíciónak nevezünk) a tagadás tehát a PPE lokális tartományában állt, az utóbbiban pedig (melyet pedig

⁷ A feladat a kísérletes szemantikai szakirodalomban bevett igazságfeltétel-megítélési (truth value judgement) feladat egyik változatával (vö. KRIFKA 2011: 254) rokonítható, amennyiben a résztvevőknek végső soron azt kellett megítélniük, hogy igaz-e a stimulus mondat, amennyiben az elsőként bemutatott parafrázis által leírt helyzet áll fenn. BOTT és RADÓ (2007) egy a miénkhez hasonló, kérdés-válasz párokat használó kísérleti feladatról megmutatták, hogy az annak révén gyűjthető anyanyelvi beszélői ítéletek konzisztenciája a más módszerek által kiváltott ítéletekétől nem különbözik szignifikánsan. BOTT és RADÓ feladatában a kérdés játszotta a parafrázis szerepét, a válasz pedig a stimulus mondatét. Köszönjük anonim bírálónk kérdését, melynek nyomán a kísérleti módszer leírását a jelen jegyzettel kiegészítettük.

NEM LOKÁLIS kondícióként címkézünk) azon kívül. A LOKÁLIS kondícióban mindhárom PPE-típus esetén a tagadás a beágyazott tagmondaton belül megelőzte az ígét és a PPE-t tartalmazó tárgyat. A LOKÁLIS és NEM LOKÁLIS kondíciók mondatai csak a tagadásnak a PPE-hez viszonyított elhelyezkedésében különböztek, a stimulusok lexikai tartalma ettől eltekintve megegyezett. A kísérletben minden résztvevő találkozott a lent (8)–(13)-ban bemutatott hatféle kritikus stimulus (két lokalitási kondíció * három PPE-típus) mindegyikével, tehát mind a PPE-típus, mind a lokalitás belső, résztvevőn belüli tényezőként szerepelt a kísérletben.

- (8) LOKÁLIS VAGY: *Azt hiszem, hogy Péter nem szereti a mandarint vagy a banánt.*
- (9) NEM LOKÁLIS VAGY: *Nem hiszem, hogy Péter szereti a mandarint vagy a banánt.*
- (10) LOKÁLIS EGY: *Azt hiszem, hogy a javítás során a szerelő nem cserélt ki egy alkatrészt.*
- (11) NEM LOKÁLIS EGY: *Nem hiszem, hogy a javítás során a szerelő kicserélt egy alkatrészt.*
- (12) LOKÁLIS VALAKI: *Azt hiszem, hogy a romok közül a tűzoltók nem mentettek ki valakit.*
- (13) NEM LOKÁLIS VALAKI: *Nem hiszem, hogy a romok közül a tűzoltók kimentettek valakit.*

Az idézőjelek között megjelenő parafrázált célértelmezések mindhárom PPE-típus esetében és mindkét lokalitási kondícióban a stimulus mondat felszíni hatóköri értelmzésének, azaz a PPE tagadáshoz viszonyított szűk hatóköri olvasatának feleltek meg. A parafrázisokat a stimulus mondatokhoz hasonló felépítésű, ugyanakkor hatókörileg egyértelmű mondatok képezték:

- (14) Parafrázis – LOKÁLIS VAGY: *Azt hiszem, hogy Péter sem a mandarint, sem a banánt nem szereti.*
- (15) Parafrázis – NEM LOKÁLIS VAGY: *Nem hiszem, hogy Péter akár a mandarint, akár a banánt szereti.*
- (16) Parafrázis – LOKÁLIS EGY: *Azt hiszem, hogy a javítás során a szerelő egy alkatrészt sem cserélt ki.*
- (17) Parafrázis – NEM LOKÁLIS EGY: *Nem hiszem, hogy a javítás során a szerelő akár egy alkatrészt is kicserélt.*
- (18) Parafrázis – LOKÁLIS VALAKI: *Azt hiszem, hogy a romok közül a tűzoltók senkit sem mentettek ki.*
- (19) Parafrázis – NEM LOKÁLIS VALAKI: *Nem hiszem, hogy a romok közül a tűzoltók akárkit is kimentettek.*

A kritikus stimulusok kísérleti kondíciók közötti eloszlását az alábbi, 1. táblázat mutatja be.

1. táblázat

A kritikus stimulusok kísérleti kondíciók közötti eloszlása

		LOKALITÁS (belső tényező)	
		LOKÁLIS	NEM LOKÁLIS
PPE-TÍPUS (belső tényező)	VALAKI	5 stimulus	5 stimulus
	EGY	5 stimulus	5 stimulus
	VAGY	5 stimulus	5 stimulus

A stimulusokat két listába rendeztük, melyek összesen harminc (mindhárom PPE-típusból mindkét lokalitási kondícióban öt-öt) kritikus, tíz kontroll-, valamint tíz töltelékstimulust tartalmaztak. A listák kizárólag abban különböztek egymástól, hogy a stimulusok álvéletlen sorrendje egymás fordítottja volt a két listában.

Kétfajta, a kritikus mondatokhoz hasonló felépítésű kontrollstimulust használtunk: az egyik típusban jólformált mondatokhoz grammatikailag engedélyezett és könnyen elérhető célértelmezéseket (jó kontroll; JK), míg a másik típusban grammatikailag nem engedélyezett, nem elérhető célértelmezéseket (rossz kontroll; RK) párosítottunk. E kontrollstimulusok egyrészt viszonyítási alapot szolgáltatottak a kritikus mondatok elfogadhatóságának vizsgálatához, másrészt ezeket a stimulusokat a figyelmetlen résztvevők kiszűréséhez is felhasználtuk.

A JK stimulusok mondatai egy pozitív polaritású főmondatból és egy ebbe legalább egyszeresen beágyazott, szintén pozitív polaritású tagmondatból álltak. Ezek a mondatok a kritikus mondatokhoz hasonlóan kétértelműek voltak: esetükben a kétértelműség abból eredt, hogy a beágyazott tárgyi szerepű határozatlan főneves kifejezés megengedte mind a specifikus, mind a nem specifikus olvasatot. Míg a specifikus olvasat a tárgyi kifejezéssel társítható egzisztenciális kvantornak a tárgyat vonzó intenzionális ígéhez viszonyított tág logikai hatókörével, addig a nem specifikus olvasat az ahhoz viszonyított szűk hatókörével írható le (lásd (20)). A JK mondatokhoz tartozó idézőjeles célértelmezések a szűk hatókörű értelmmezést, azaz a tárgyi szerepű főnevek nem specifikus olvasatát kódolták (lásd (21)).

(20) JK: *Azt hiszem, hogy Klári keres egy kék tollat.*

(21) Parafrázis – JK: *Azt hiszem, hogy Klári kék tollat keres.*

A (22)-ben illusztrált RK mondatok egy pozitív polaritású főmondatból és egy beágyazott tagadó mellékmondatból épültek fel. A beágyazott mondatokban a *mindegyik* pozitív univerzális kvantor által bevezetett tárgyi szerepű főnévi kifejezés az ige mögött foglalt helyet. A magyarban a pozitív univerzális kvantoros kifejezések nem vehetnek fel hatókört közvetlenül a mondattagadás felett (pl. É. KISS 2002). Ezért az RK mondatok esetében csak a „nem > mindegyik” olvasat volt elérhető. A kísérlet során az RK mondatokhoz azonban a grammatikailag nem engedélyezett, fordított hatóköri ’mindegyik > nem’ (= ’egyik sem’) értelmezést társítottuk (lásd (23)).

(22) RK: *Azt hiszem, hogy Péter nem nézte meg mindegyik filmet.*

(23) Parafrázis – RK: *Azt hiszem, hogy Péter egyik filmet sem nézte meg.*

A kontrollstimulusokhoz hasonlóan a tölteléksztimulusoknak is két típusát használtuk a kísérletben. A többi stimulushoz hasonlóan a tölteléksztimulusok is egy (esetükben pozitív polaritású) főmondatból és egy tárgyi szerepű beágyazott tagmondatból épültek fel. Mindkét típusú tölteléksztimulus logikailag potenciálisan kétértelmű grammatikus magyar mondatokat tartalmazott. Az egyik típusú tölteléksztimulusban a mondathoz társított, a felszíni hatóköri viszonyoknak megfelelő célértelmezés grammatikailag engedélyezett, és várakozásunk szerint közepesen elfogadható volt (közepes tölteléksztimulus; KT). A másik típusban a megadott, szintén egyenes hatóköri viszonyokat kifejező célértelmezés grammatikailag nem volt engedélyezett, ennek megfelelően várakozásunk szerint a mondat nem volt a parafrázis szerint értelmezhető (rossz tölteléksztimulus; RT).

Amint azt (24) illusztrálja, a KT mondatok alany–ige–tárgy szórendű, tagadó beágyazott mondatainak tárgya minden esetben a *két* számnév által módosított főnevet tartalmazó főnévi kifejezés volt. Ezek a mondatok a beágyazott mondatban szereplő tagadás és a tárgyi szerepű, egzisztenciálisan értelmezett számneves kifejezés két lehetséges hatókörértelmezéséből fakadóan voltak kétértelműek (’nem > létezik két’; ’létezik két > nem’). A parafrázis, mely alapján a résztvevőknek ezeket a mondatokat meg kellett ítélniük, a *legfeljebb egy* módosított számnevet tartalmazó főnévi kifejezés kódolta a felszíni hatókörnek megfelelő ’nem > létezik két’ értelmezést (lásd (25)).

(24) KT: *Azt hiszem, hogy Julcsi nem olvasott végig két cikket.*

(25) Parafrázis – KT: *Azt hiszem, hogy Julcsi legfeljebb egy cikket olvasott végig.*

A (26)-ban bemutatott RT mondatok tárgy–alany–ige szórendű, tagadó beágyazott mondatainak tárgya minden esetben egy univerzálisan kvantifikált főnevet tartalmazó tárgyi szerepű főnévi kifejezés volt. Az RT mondatok logikailag potenciális kétértelműsége a beágyazott mondatban található kvantoros kifejezés és a tagadás két potenciális hatóköri olvasatából fakadt (’nem > minden’; ’minden > nem’). mondatokhoz párosított parafrázisok tagadó egyeztetéses formával fejezték ki a felszíni hatókörnek megfelelő ’minden > nem’ értelmezést (lásd (27)).

(26) RT: *Úgy gondolom, hogy minden filmet a filmkritikusok sem láttak.*

(27) Parafrázis – RT: *Úgy gondolom, hogy a filmkritikusok sem láttak egy filmet sem.*

A – kritikus stimulusokhoz részben hasonló, de azokkal az eredmények elemzésekor közvetlenül nem összehasonlított, a kísérlet faktoriális elrendezésén kívül eső – tölteléksztimulusok egyszerre több célt szolgáltak. Egyrészt segítettek a feladatban található mondatok megadott célértelmezés szerinti értelmezésének nehézségi szintjeit kiegyensúlyozni. Másrészt növelték a megítélendő stimulusok változatosságát, és ezzel hozzájárultak a résztvevők figyelmének fenntartásához. Végül: az esetlegesen előforduló figyelmetlen résztvevők kiszűréséhez felhasználtuk az RT mondatokra adott ítéleteket is.

A három gyakorló stimulusmondat megítélése során a résztvevők három különböző nehézségű értelmezéssel találkoztak. A gyakorlómondatoknak a velük társított parafrázisoknak megfelelő értelmezése az egyik gyakorlómondat esetében könnyű, a másokban nehéz, a harmadikban grammatikailag lehetetlen volt.

2.2. Predikciók. A két kutatási kérdésünkre adható lehetséges válasz-kombinációk közül először azt az esetet tárgyaljuk, amelyben eredményeink mindkét kérdésre pozitív választ adnak. Ebben az esetben mind az általunk használt „klasszikus” PPE-k (*egy, vala-névmás*), mind a *vagy* a szakirodalom szerinti PPE-szerű viselkedést mutatják – azaz a lokális, velük egy tagmondatban lévő tagadás hatókörében nem értelmezhetőek, míg a hozzájuk képest nem lokális, főlérendelt tagmondatban lévő tagadás jelenlétében elfogadhatóak. Ez esetben azt várjuk, hogy (A) PPE-típustól függetlenül a LOKÁLIS kondíció mondatainak megítélése statisztikailag szignifikánsan alacsonyabb lesz a NEM LOKÁLIS mondatokra adott értékeknél. E hatás mértékét tekintve azt várjuk, hogy (B) a különbség mindhárom PPE-típusban legalább közepes méretű lesz. A közepes hatásméret alsó küszöbértékét az elfogadhatósági ítéletek szakirodalmában is alkalmazott 0,5 Cohen-féle *d* értékben határozzuk meg.⁸

A LOKÁLIS kondíció elvárt alacsony elfogadhatósági szintjét a NEM LOKÁLIS kondícióval való összehasonlításra túl úgy is jellemezhetjük, hogy összehasonlítjuk az RK, illetve a JK mondatokéval. A LOKÁLIS kondíció mondatai és az RK mondatok között nem feltétlenül várunk szignifikáns különbséget, tekintve, hogy mindkét stimulustípusban a hipotézis szerint grammatikailag nem engedélyezett olvasat szerepel. Azonban a korábbi szakirodalom sem empirikus adatokat nem mutat be, sem elméleti jóslatokat nem tesz arra vonatkozóan, hogy a PPE-k pontosan mennyire elfogadhatatlanok egy a tagadáshoz hasonló lokális anti-engedélyező elem közvetlen hatókörében, miközben az elfogadhatósági ítéletek gyakorlatából tudható, hogy egyes grammatikailag nem engedélyezett, ezért alacsony elfogadhatóságú formák vagy jelentések elfogadhatósági szintjei között szignifikáns különbségek lehetségesek. Mivel kísérletünk kritikus és kontrollmon-

⁸ A Cohen-féle *d* két átlag különbségének egy – többek között a pszichológia tudományterületén elterjedt – standardizált mérőszáma, mely két kondíció átlagának és a szórásnak az arányát fejezi ki. COHEN (1988) nyomán a 0,2 és 0,5 közötti különbségeket kis méretű hatásnak, a 0,5 és 0,8 közöttieket közepes méretű hatásnak, a 0,8-nál nagyobbakat pedig nagy méretű hatásnak nevezik.

datai szerkezetükben és lexikai tartalmukban sem voltak teljesen párhuzamosak, ezért az RK és LOKÁLIS mondatok összehasonlításával kapcsolatban a következő – viszonylag szigorúnak tekinthető – jóslatot fogalmazzuk meg. Várakozásunk szerint (C) a három PPE-típus LOKÁLIS mondatai vagy nem fognak statisztikailag szignifikáns módon különbözni az RK mondatoktól, vagy ha mégis szignifikáns különbség mutatható ki, akkor ez a különbség kis hatásméretű lesz (Cohen-féle $d < 0,5$). A JK és a LOKÁLIS mondatok összehasonlításaival kapcsolatban pedig azt várjuk, hogy (D) a két kondíció mondatai között nagy méretű (Cohen-féle $d > 0,8$), statisztikailag szignifikáns különbség fog mutatkozni.

A három PPE-t egymással összehasonlítva azt várjuk, hogy mindhárom PPE-típus elfogadhatósága egymáshoz hasonló lesz mind a LOKÁLIS, mind a NEM LOKÁLIS kondícióban, így (E) nem várunk statisztikailag szignifikáns különbséget sem a LOKÁLIS, sem a NEM LOKÁLIS mondatok elfogadhatósági értékei között. Továbbá azt is jósoljuk, hogy (F) a három PPE-típushoz tartozó LOKÁLIS kondíciók mondatai közel megegyező mértékben fognak különbözni az RK mondatoktól, és hasonlóképpen a JK mondatoktól is (melyet az egyes PPE-típusokban a LOKÁLIS vs. RK, illetve a LOKÁLIS vs. JK összehasonlítások által feltárt különbségek hatásméretei 95%-os konfidenciaintervallumainak vizsgálata révén ellenőrzünk).

Várakozásainkat az alábbiakban foglaljuk össze:

- (A) A LOKÁLIS kondíció megítélése mindhárom PPE-típus esetében statisztikailag szignifikánsan alacsonyabb lesz a NEM LOKÁLIS kondíciónál.
- (B) A LOKÁLIS és a NEM LOKÁLIS kondíciók különbsége mindhárom PPE-típus esetében legalább közepes méretű lesz.
- (C) A LOKÁLIS és az RK kondíció között vagy nem lesz szignifikáns különbség, vagy legfeljebb kis hatásméretű (Cohen-féle $d < 0,5$) szignifikáns különbség lesz mindhárom PPE-típus esetében.
- (D) A JK és a LOKÁLIS kondíciók között mindhárom PPE-típus esetében nagy méretű (Cohen-féle $d > 0,8$) szignifikáns különbség lesz.
- (E) A három PPE-típus között nem lesz statisztikailag szignifikáns különbség sem a LOKÁLIS, sem a NEM LOKÁLIS kondícióban.
- (F) A LOKÁLIS mondatok a három PPE-típus esetében egymáshoz hasonló mértékben fognak különbözni a RK mondatoktól és a JK mondatoktól is.

Ha a kísérletben gyűjtött adatok kielégítik e várakozásokat, akkor tanulmányunk 1. részének végén megfogalmazott mindkét kutatási kérdésre pozitív válasz adható. Ez a kimenetel azt támasztaná alá, hogy azt általunk vizsgált mindhárom elem – így tehát a diszjunkció is – a PPE-k közé sorolható. Amennyiben a magyar jól reprezentálja a +PPE nyelveket, akkor ez az eredmény azt valószínűsíti, hogy az 1. részben tárgyalt, a +PPE nyelvek diszjunkciójának PPE-státuszát

érintő, nyelvközi és egyes nyelveken belüli ellentmondások forrása az adott különböző kísérletek módszertani eltéréseiben keresendő.

Természetesen az is előfordulhat, hogy az adatok valamelyik – vagy akár mindkét – kérdésünk tekintetében negatív válaszra vezetnek. A diszjunkciónak az 1. részben bemutatott ellentmondásos viselkedésből kiindulva nem kizárt, hogy a „klasszikus” PPE-k a fent bemutatott várakozásoknak megfelelően az elvárt PPE-viselkedést mutatják, a diszjunkció azonban nem, ehelyett az általa kiváltott ítéletek kimutathatóan különböznek a „klasszikus” PPE-kre kapott ítéletek mintázataitól. Ezt a lehetséges kimenetelt a jelenlegi PPE-elméletek önmagukban nem magyaráznák, így kísérletünk további, az okok feltárását célzó kutatások felé nyitna utat. A diszjunkció által mutatott mintázatnak a „klasszikus” PPE-kétől való eltéréseinek függvényében ez az eredmény adott esetben akár kétséget is ébreszthet afelől, hogy a *vagy* valóban a PPE-k közé tartozik-e.

Tekintve, hogy a lokális tagadás hatókörében álló „klasszikus” PPE-k elfogadhatatlanságáról nem áll rendelkezésre kellő mennyiségű korábbi szakirodalmi adat, az a lehetőség sem kizárt, hogy ugyan a *vagy* és a vizsgált „klasszikus” PPE-k azonos módon viselkednek – ti. ugyanazokat a várakozásainkat teljesítik –, azonban sem a *vagy*, sem az *egy*, sem a *vala*-névmás nem a PPE-ktől elvárt mintázatot mutatják. Reális eshetőség például, hogy mindhárom általunk vizsgált PPE LOKÁLIS kondíciója viszonylag magas értékelést kap. Ez felvetné, hogy a PPE-k valójában mégsem teljesen értelmezhetetlenek, legfeljebb kisebb-nagyobb mértékben degradáltak vagy diszpreferáltak egy lokális anti-engedélyező közvetlen hatókörében.⁹ Ez a kimenetel szembeállítaná őket az NPE-kkel, melyek az őket engedélyező operátorok hiányában nem csupán relatíve degradáltak, hanem szélsőségesen alacsony elfogadhatóságúak (l. pl. DRENHAUS et al. 2005; LIU–SOEHN 2009; MAYER et al. 2019; SCHAEBBICKE et al. 2021 kísérleti eredményeit).

2.3. Eredmények. A kísérletet teljesítő 36 résztvevő közül azoknak az adatait vetettük alá statisztikai elemzésnek, akik az alábbi két szűrőfeltétel közül legalább az egyiknek megfeleltek: 1. a JK, RK és RT kondíciók összesen tizenöt, egyértelműen magasnak illetve alacsonynak várt elfogadhatóságú mondatának megítélése során legfeljebb három, az elvárttal ellentétes ítéletet adtak (a JK mondatok esetén az 1-es és 2-es, az RK és RT mondatok esetén a 4-es és 5-ös értékeléseket tekintettük az elvárttal ellentétesnek); 2. a JK mondatokra adott értékeléseiknek átlaga magasabb volt, mint 3, illetve az RK és RT mondatok átlagos értékelése – mindkét mondat típusban külön-külön – alacsonyabb volt, mint 3. E szűrőfeltételek alkalmazása után hat főt kellett kizárnunk.¹⁰

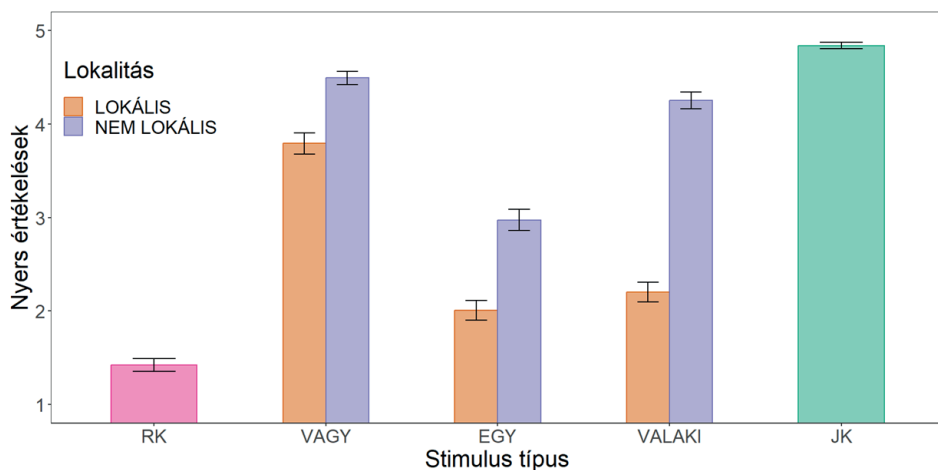
⁹ Ez az eredmény látszólag rokonítható volna LIU és SOEHN (2009) PPE-ket tartalmazó német mondatok elfogadhatósági ítéletekre épülő kísérletének egyes részeredményeivel. LIU és SOEHN kísérlete – mely általánosságban megerősíti, hogy a tagadás hatókörében álló PPE-k átlagos elfogadhatósága alacsony – egyes, a szerzők által PPE-ként kategorizált, szintaktikailag összetett idiomatikus kifejezésekről azt mutatja ki, hogy nem teljesen elfogadhatatlanok tagadás hatókörében. Ezt azonban a szerzők azzal magyarázzák, hogy az adott idiómákat a kísérlet résztvevőinek egy része feltehetőleg nem ismerte: ezt támasztja alá, hogy ugyanezeknek az idiómáknak az elfogadhatósága pozitív környezetekben sem volt jelentősen magasabb.

¹⁰ A weben keresztül végzett, nem felügyelt kísérletek a figyelmetlen kitöltés nagyobb veszélyével járnak, mint a személyes jelenléttel végzett kísérletek (MEADE–CRAIG 2012). A résztvevők

A kritikus és kontroll kondíciókra kapott nyers ítéletek átlagait az 1. ábra mutatja be. A statisztikai elemzés során – mely így összesen 30 résztvevő (11 nő és 19 férfi, 20–43 év közöttiek) adatait tartalmazta – a nyers ítéleteket az összes stimulusra adott válaszuk alapján résztvevőnként z-transzformáltuk.¹¹ A két kontroll kondíció (JK és RK), valamint a hat kritikus kondíció nyers és z-transzformált átlagait (szórásukkal együtt) a 2. táblázat tartalmazza.

1. ábra

A kritikus és kontroll kondíciókra kapott nyers ítéletek átlagai standard hibával



2. táblázat

Az egyes kondíciók nyers ítéleteinek és z-értékeinek átlagai (szórással)

Kondíciók		Nyers ítéletek	Z-értékek
KONTROLL	RK	1,42 (0,84)	-1,02 (0,52)
	JK	4,84 (0,40)	1,10 (0,32)
VAGY	LOKÁLIS	3,79 (1,40)	0,45 (0,81)
	NEM LOKÁLIS	4,49 (0,86)	0,87 (0,52)
EGY	LOKÁLIS	2,00 (1,31)	-0,65 (0,80)
	NEM LOKÁLIS	2,97 (1,40)	-0,06 (0,80)
VALAKI	LOKÁLIS	2,20 (1,30)	-0,55 (0,76)
	NEM LOKÁLIS	4,25 (1,11)	0,71 (0,66)

kizárási aránya az ilyen nyelvészeti kísérletek esetében nem ritkán 10-20% közötti (pl. SPROUSE 2011). Köszönjük az egyik anonim bírálónk erre vonatkozó kérdését.

¹¹ A z-transzformáció az egyes résztvevők skálahasználatának különbségeit normalizálja azáltal, hogy a nyers értékek helyett az adott résztvevő összes válaszáának átlagától számított, a szórás egységében kifejezett távolságát adja meg.

A statisztikai elemzést az R (R Core Team 2022) program 1 me 4 csomagjával (BATES et al. 2015) végeztük. A kritikus mondatok z -értékeit egy lineáris kevert modellben elemeztük. A *car* csomag (FOX–WEISBERG 2019) anova függvénye alkalmazásával végzett visszalépéses elimináció eredményeként kapott modell egymással interakcióban lévő független változóként tartalmazta a PPE-típust és a lokalitást. A modell a következő random tényezőket tartalmazta: a kísérleti stimulusok mint metszéspont, a kísérleti személyek mint metszéspont, valamint a PPE-típus és a lokalitás tényezők interakciója mint a metszésponthoz tartozó random meredekség. A modell a PPE-típus ($\chi^2 = 43,940$; $p < 0,0001$), a lokalitás ($\chi^2 = 47,574$; $p < 0,0001$), valamint e két tényező interakciójának ($\chi^2 = 15,758$; $p < 0,001$) szignifikáns fő hatását mutatta ki. A fő hatások további vizsgálata céljából az *lsmeans* csomag (LENTH 2016) megegyező nevű függvényével Holm–Bonferroni-korrekciót tartalmazó páronkénti összehasonlításokat végeztünk.

A páronkénti összehasonlítások tanúsága szerint a lokalitás tényező szignifikáns fő hatása abból származott, hogy mindhárom PPE esetében a NEM LOKÁLIS mondatok elfogadhatósága szignifikánsan jobbnak bizonyult a LOKÁLIS mondatok elfogadhatóságánál (LOKÁLIS VAGY vs. NEM LOKÁLIS VAGY: $t = -2,443$; $p = 0,0197$; Cohen-féle $d = 0,62$; LOKÁLIS EGY vs. NEM LOKÁLIS EGY: $t = -3,558$; $p = 0,0011$; Cohen-féle $d = 0,74$; LOKÁLIS VALAKI vs. NEM LOKÁLIS VALAKI: $t = -7,502$; $p < 0,0001$; Cohen-féle $d = 1,77$).

A PPE-típus tényezőjének szignifikáns fő hatása mögött az állt, hogy a három PPE elfogadhatósága sem a LOKÁLIS, sem a NEM LOKÁLIS kondícióban belül nem alakult egyformán. A LOKÁLIS kondícióban belül a VAGY és az EGY ($t = -5,772$; $p < 0,0001$; Cohen-féle $d = 1,37$), valamint a VAGY és a VALAKI ($t = 5,070$; $p < 0,0001$; Cohen-féle $d = 1,27$) elfogadhatósági szintjei is statisztikailag szignifikáns módon különböztek egymástól, míg csupán az EGY vs. VALAKI összehasonlítás nem mutatott ki szignifikáns különbséget ($t = -0,625$; $p = 0,53$; Cohen-féle $d = 0,13$). A NEM LOKÁLIS kondícióban sem volt azonos a három PPE elfogadhatósága: a VAGY és az EGY ($t = -5,676$; $p < 0,0001$; Cohen-féle $d = 1,37$), illetve az EGY és a VALAKI ($t = -4,347$; $p < 0,0001$; Cohen-féle $d = 1,00$) összehasonlítás esetében is szignifikáns különbség mutatkozott. Ebben a kondícióban is csak egyetlen olyan összehasonlítás volt, amelyben a statisztikailag szignifikáns különbség hiánya a mondatok hasonló elfogadhatóságára utalt: VAGY vs. VALAKI ($t = 0,973$; $p = 0,33$; Cohen-féle $d = 0,26$).

A PPE-típus és a lokalitás interakciója abból származott, hogy bár mindhárom PPE esetében a NEM LOKÁLIS mondatok szignifikánsan magasabb értékeléseket kaptak a LOKÁLIS mondatoknál, a NEM LOKÁLIS kondícióban az EGY elfogadhatósága (átlag: 2,97) a másik két PPE-hez viszonyítva (VAGY: 4,49 VALAKI: 4,25) számottevően alacsonyabb volt, a LOKÁLIS kondícióban pedig a VAGY mondatai (átlag: 3,79) jelentősen magasabb értékeléseket kaptak, mint az EGY (átlag: 2,00) és a VALAKI (átlag: 2,20) mondatai.

A LOKÁLIS mondatok az RK, illetve a JK stimulusokkal történő összehasonlításához egy a kritikus és a kontrollmondatok z -értékeit tartalmazó lineáris kevert modellt használtunk. A kritikus stimulusok vizsgálatkor alkalmazott elrendezést, azaz a PPE-típus és a lokalitás külön tényezőként való kezelését ez a modell nem

tette lehetővé, mivel a kontrollstimulusok nem rendelkeztek kétféle lokalitással. Ezért ebben a modellben a kétféle kontroll (JK, RK) és a hatféle kritikus stimulus (háromféle PPE-típus * kétféle lokalitás) egyetlen változó (stimulustípus) összesen nyolc szintjét alkotta.¹² A *c a r* csomag (FOX–WEISBERG 2019) anova függvényével végzett visszalépéses elimináció után kapott modell független változóként a stimulustípust, egyetlen random tényezőként pedig a kísérleti stimulusokat mint metszéspontot tartalmazta. A modell a stimulustípus szignifikáns fő hatását mutatta ki ($\chi^2 = 515,42$; $p < 0,0001$). A LOKÁLIS vs. RK, illetve JK összehasonlítások vizsgálatának céljából az *l s m e a n s* csomag (LENTH 2016) azonos nevű függvényének segítségével Holm–Bonferroni-korrekción tartalmazó páronkénti összehasonlításokat végeztünk.¹³ Ezek eredményei azt mutatták, hogy az RK mondatok mindhárom PPE-típus LOKÁLIS mondataitól szignifikáns módon különböztek (RK vs. LOKÁLIS VAGY: $t = -11,255$; $p < 0,0001$; RK vs. LOKÁLIS EGY: $t = -2,794$; $p = 0,0436$; RK vs. LOKÁLIS VALAKI: $t = -3,580$; $p = 0,0090$). A LOKÁLIS és a JK stimulusok páronkénti összehasonlításai azt mutatták, hogy mind a LOKÁLIS VAGY ($t = 5,056$; $p = 0,0002$), mind a LOKÁLIS EGY ($t = 13,516$; $p < 0,0001$), mind a LOKÁLIS VALAKI ($t = 12,730$; $p < 0,0001$) kondíciók stimulusaira adott értékelések szignifikánsan különböztek a JK mondatok megítélésétől.

A hatásméretek 95%-os konfidenciaintervallumának kiszámításához az *M B E S S* csomag (KELLEY 2007b) *ci.smd* függvényét használtuk. A három PPE-típus LOKÁLIS vs. RK hatásméreteinek összehasonlításai mindhárom esetben nagyobbak bizonyultak a várt, legfeljebb 0,5 Cohen-féle d mértékű különbségnél. Míg a LOKÁLIS EGY és az RK (Cohen-féle $d = 0,53$, 95%-os konfidenciaintervallum: 0,31–0,77), valamint a LOKÁLIS VALAKI és az RK (Cohen-féle $d = 0,72$, 95%-os konfidenciaintervallum: 0,48–0,95) különbségek hasonló méretűek voltak – és így hatásméretük konfidenciaintervallumai átfedtek –, addig a LOKÁLIS VAGY vs. RK különbség (Cohen-féle $d = 2,16$, 95%-os konfidenciaintervallum: 1,87–2,44) annyival nagyobbak bizonyult a másik két különbségnél, hogy hatásmérete konfidenciaintervallumának nem volt a másik két összehasonlítás konfidenciaintervallumaival átfedő tartománya.

A LOKÁLIS vs. JK kondíciók összehasonlítása során kapott hatásméretek mindhárom esetben jelentős mértékben meghaladták az előzetesen várt legalább 0,8 Cohen-féle d mértékű különbséget. Azonban – a fenti LOKÁLIS vs. RK összehasonlításokkal párhuzamosan – amíg a LOKÁLIS EGY vs. JK (Cohen-féle $d = 2,88$, 95%-os konfidenciaintervallum: 2,56–3,20) és a LOKÁLIS VALAKI vs. JK (Cohen-féle $d = 2,83$, 95%-os konfidenciaintervallum: 2,51–3,15) különbségek mérete hasonlóknak bizonyult – így hatásméretük konfidenciaintervallumai átfedést mutat-

¹² A kritikus mondatok a kontrollokat nem tartalmazó korábbi modellben történő statisztikai elemzését egyszerűen a kísérleti elrendezés minél szorosabb követése, másrészt a kontrollokkal kiegészített modell páronkénti összehasonlításainak magas számából adódó túlkorrigálás veszélyének kiküszöbölésére való törekvés motiválta.

¹³ A kontrollstimulusokat is tartalmazó modellen elvégzett páronkénti összehasonlítások a kritikus mondatok összehasonlításakor ugyanazokat az összehasonlításokat találták statisztikailag szignifikánsnak, mint a fentebb bemutatott, csak a kritikus stimulusokat tartalmazó modell páronkénti összehasonlításai.

tak – addig a LOKÁLIS VAGY és a JK összehasonlítás (Cohen-féle $d = 1,07$ 95%-os konfidenciaintervallum: 0,83–1,31) által feltárt különbség mérete annyival kisebb volt, hogy a hatásméret konfidenciaintervalluma nem fedett át a másik két PPE és a JK mondatok különbségének konfidenciaintervallumaival.

(Folytatjuk.)

SURÁNYI BALÁZS – GULÁS MÁTÉ
HUN-REN Nyelvtudományi Kutatóközpont
Pázmány Péter Katolikus Egyetem