

El repertorio de la crianza

Cuatro códigos consonánticos, los clusters que forman y los que no: NANA · MAMA · PAPA · TATA en 5.501 lenguas

Alejandro Toledo Martínez

Nota de alcance. Este trabajo mide un fenómeno del nivel de la adquisición y la articulación —las *nursery forms*—, no una firma genealógica profunda. Se hizo, entre otras razones, como prueba de que el instrumento recupera lo que ya se sabe antes de pedirle nada nuevo. Los resultados que aquí se dan como propios son tres: el catálogo cuantificado de clusters, el resultado negativo sobre la criatura, y la constatación de que no existe corpus comparativo del registro afectivo para la inmensa mayoría de las lenguas tratadas.

Resumen

Se quitan las vocales a 100.507 formas léxicas de Lexibank (5.501 lenguas, 298 familias) y se sustituye cada consonante por su clase de Dolgopolsky. Sobre ese esqueleto se pregunta si los conceptos del núcleo de crianza —madre, padre, bebé, abuela, abuelo— se surten de un repertorio de cuatro códigos: N (nana, ana), M (mama, ama), P (papa, apa, baba) y T (tata, ata).

La respuesta es afirmativa y con margen amplio. De las palabras formadas por una sola clase de consonante, el cuarteto se lleva el 80,9% [79,8; 82,0] en la crianza frente al 48,0% [47,3; 48,8] del vocabulario corriente, cuando el azar sobre diez clases daría 40%; la diferencia es de 32,9 puntos [31,5; 34,1], con una h de Cohen de 0,71, y se repite dentro de 40 de 45 familias comparadas consigo mismas (Wilcoxon pareado $p = 8,4 \cdot 10^{-10}$, delta de Cliff 0,78). El cuarteto no solo sube: expulsa. La clase velar K, la más frecuente entre los monoclase del léxico corriente (18,1%), cae al 6,7%.

El desglose por tamaño de efecto obliga a corregir el titular, y la corrección se da aquí y no enterrada más abajo. El cuarteto no sube en bloque. M, N y P pasan del 30,9% al 63,0% y se llevan casi todo el efecto; T no se mueve (17,2% → 17,9%, $h = 0,02$). T no es una clase que la crianza prefiera: es una clase que la crianza no expulsa, mientras sí expulsa a K, S y R. Lo que hay que explicar es una evitación, no una preferencia.

El resultado no depende de la tabla de clases. Repetido bajo siete lecturas —Dolgopolsky tal como lo trae Lexibank, Dolgopolsky recalculado con LingPy, SCA en dos lecturas, ASJP, rasgos articulatorios de PanPhon y IPA crudo sin fusionar nada (201 consonantes distintas, donde el azar da 8% al cuarteto)—, la diferencia crianza – control se mantiene entre +31 y +34 puntos. Tampoco depende de las familias grandes: quitando enteros el indoeuropeo y el austronesio quedan +30,7 puntos, y sorteando una lengua por familia, de modo que todas pesen igual, quedan +31,8 [26,1; 37,0], positiva en el 100% de 500 sorteos.

El repertorio es compartido y no repartido: los cuatro códigos aparecen en los cinco conceptos, con inclinaciones distintas —el femenino hacia las nasales, el masculino hacia las oclusivas— y solapamientos de perfil de hasta 0,92. Se cataloga cada cluster con su tamaño, su dispersión familiar y su huella geográfica, y se distingue con una cifra —el número de familias que toca— entre los clusters del repertorio (50–60 familias) y los clusters de herencia de rama (4–7 familias).

Dos resultados son negativos y se dan por tales. Primero, la lectura de los cuatro códigos como dos ejes cruzados (nasal/oclusiva × labial/dental) produce una razón de momios global de 2,01 [1,52;

2,67] que se desvanece a 1,01 [0,65; 1,56] al estratificar por familia. Segundo, y contra la expectativa, no hay cluster de la criatura: con once conceptos, 4.563 lenguas, la lengua como unidad de conteo y los intentos igualados, la criatura queda diez puntos por debajo de su propio vocabulario corriente (28,6% frente a 38,7%, $h = -0,21$), en 24 de 27 familias, y sin más reduplicación que el control. La razón es estructural y se lee en el largo del código: las lenguas no nombran a la criatura con una palabra primitiva, la derivan.

Queda en pie, sin embargo, un efecto de transferencia: cuando la criatura sí usa el repertorio, tiende a tomar el código de los padres de su misma lengua (18,7% frente a 14,5% esperado, $p = 0,0005$, $h = 0,11$ — real y pequeño). El código no le pertenece: le llega prestado.

Finalmente se documenta una laguna de corpus que condiciona todo lo anterior. Lexibank recoge el término de referencia, no el de trato. De sus 1.740.092 registros, nueve llevan marca de vocativo o de término de trato, y esos nueve indican que el vocativo es una palabra distinta y más próxima al repertorio. Las cifras de este trabajo son, por tanto, un piso y no un techo.

1. El problema y lo que ya se sabía

Que la palabra para 'madre' se parezca en lenguas sin parentesco es una de las observaciones más antiguas de la lingüística comparada, y también una de las mejor explicadas. Murdock (1959) recogió los paralelos entre términos parentales en centenares de lenguas; Jakobson (1960) dio la explicación que sigue vigente: la nasal labial es lo primero que un lactante produce con la boca ocupada en mamar, y el adulto recoge ese balbuceo y lo devuelve como nombre. Ferguson (1964) describió el registro entero —el *baby talk*— como un subsistema con su fonología, su morfología y un léxico especializado de veinticinco a sesenta ítems.

Lo que aquí se añade no es la explicación, que ya existe, sino tres cosas que no estaban medidas.

Primera: el tamaño y la forma de los clusters. Decir «se parecen en muchas lenguas» no es lo mismo que decir «este código agrupa 639 lenguas de 53 familias y este otro 136 lenguas de 7». La segunda formulación distingue un repertorio de una herencia, y esa distinción se hace con una cifra.

Segunda: qué pasa con la criatura. La literatura describe el registro dirigido al niño; no consta que se haya medido si la palabra *para* el niño participa del mismo repertorio. Aquí se mide y sale que no.

Tercera: dónde se acaba el corpus. El registro afectivo —lo que la madre dice al niño en brazos— no está en las bases de datos comparativas, y conviene decirlo con números.

Una advertencia de método que atraviesa todo el trabajo. Un fenómeno tan visible como *mama / papa* invita a la apofenia: siempre habrá lenguas que encajen. Por eso cada afirmación se contrasta contra un bloque de comparación fijado antes de calcular, y las que dependen del tamaño de las familias se rehacen tomando la familia —y no la lengua— como unidad.

1.1 Tres palabras, definidas antes de usarlas

Tres términos sostienen el argumento y cada uno se usa con más de un sentido en la literatura vecina. Se fijan aquí, una vez, y después no se estiran.

Código. Un objeto puramente operativo: la cadena que queda al quitarle las vocales a una forma y sustituir cada consonante por su clase de sonido. *mama* → MM, *ama* → M, *madre* → MTR. Un código es una representación de una palabra, nada más. No afirma nada sobre el significado, sobre la historia ni sobre la mente del hablante.

Cluster. Un conjunto de lenguas que le dan el mismo código al mismo ser. Nada más. Es un objeto estadístico, definido por extensión: N-femenino son las 639 lenguas, de 53 familias, cuya palabra para

madre o abuela tiene código N. Un cluster no es una familia histórica, ni una protoforma, ni una regla: las lenguas de dentro pueden no tener relación entre sí, y las de fuera no son excepciones a nada. En todo este trabajo, cada vez que se reporta un cluster van con él dos cifras: cuántas lenguas y cuántas familias independientes. La segunda es la que separa un cluster de convergencia de uno de herencia, y el apartado 6.1 hace esa separación explícita.

Repertorio. La afirmación propiamente dicha, y la única de las tres que podría ser falsa. Un repertorio es un inventario pequeño de códigos que (i) un número grande de lenguas sin relación entre sí emplea para el mismo conjunto reducido de conceptos, (ii) está *compartido* entre esos conceptos en vez de asignar un código a cada uno, y (iii) es *cerrado*: no solo usa de más a sus miembros, sino que usa de menos a las alternativas. Cada una de las tres partes se mide por separado más abajo: (i) en 4.3, por familias; (ii) en 4.4, por solapamiento de perfiles; (iii) en 4.1, por lo que le pasa a K y a S. Un hallazgo que dijera solo que los conceptos de crianza usan más nasales que la media no establecería un repertorio en este sentido; y un hallazgo de que cada concepto tiene su código privado falsaría la parte (ii).

1.2 De qué nivel es la explicación

Al lector se le debe una declaración llana de qué clase de afirmación se está haciendo, porque los mismos datos podrían vestirse de cuatro teorías distintas.

Esta es una afirmación sobre articulación y adquisición, leída en el registro tipológico. El mecanismo es el de Jakobson y no está en discusión: los primeros cierres del lactante producen un conjunto estrecho de consonantes, quien lo cuida los oye como interpelación, y la forma resultante se fosiliza en el léxico. Aquí no se propone ningún mecanismo nuevo. Lo nuevo es la medida: cuán grande es la convergencia resultante, qué clases recluta y cuáles rechaza, y si es un inventario o son cinco.

Explícitamente no es una afirmación sobre genealogía profunda. Ningún cluster de los que se reportan aquí se ofrece como prueba de descendencia común, y el apartado 6.1 existe precisamente para apartar los clusters que sí son heredados. Tampoco es una afirmación sobre el significado: los códigos no significan nada, y ninguna lectura de M como «materno» o de T como «paterno» está pretendida ni respaldada.

Dónde se sitúa exactamente el fenómeno entre la fonética y la pragmática es una pregunta abierta de verdad, y el apartado 9 sostiene que no puede zanjarse con los corpus que hoy existen, porque el registro donde el fenómeno de hecho vive —el trato, no la referencia— es justamente el que ninguna base de datos recoge.

2. Datos y método

2.1 El corpus

Lexibank (List *et al.* 2022) reúne listas léxicas estandarizadas de más de cinco mil lenguas con segmentación fonética uniforme y enlace a Concepticon. Se leyó el `forms.csv` completo —1.740.092 registros, de los que 1.729.060 traen material fonético— y se extrajeron 100.507 formas de 33 conceptos en 5.501 lenguas y 298 familias.

Una advertencia sobre la vía de acceso, porque afecta a la reproducibilidad. Las tablas derivadas que este proyecto mantiene descartan las formas con menos de dos clases distintas, lo que elimina justamente *mama*, *ama*, *nana* y *papa*. Para el concepto MOTHER eso significa 1.530 formas y 96 familias en la tabla derivada, frente a 4.420 formas y 189 familias en el fichero crudo. Todo lo que sigue se calcula sobre el crudo.

2.2 El código

De cada forma se eliminan las vocales y cada consonante se sustituye por su clase de Dolgopolsky (1964), un agrupamiento de sonidos por probabilidad de cambio fonético que Lexibank ya trae calculado mediante LingPy:

clase	sonidos	clase	sonidos
P	p, b, f	M	m
T	t, d, θ	N	n, ŋ, ɲ
S	s, z, ʃ	R	l, r
K	k, g, x, q	W	w, v
H	h, ʔ	J	j

Así *mama* → MM, *ama* → M, *nene* → NN, *ana* → N, *madre* → MTR.

Conviene advertir dos asignaciones que no son las que la intuición articulatoria sugeriría, porque el criterio de Dolgopolsky no es el punto de articulación sino la estabilidad histórica. La fricativa *v* no va con *f* sino a *W*, junto a *w*; y las africadas no van con las sibilantes: *ts*, *dz* y *tʃ* caen en *K*, mientras *tʃ* y *dʒ* caen en *T*. Se ha comprobado que la asignación es unívoca en el corpus —cada segmento a una sola clase en el 100% de sus ocurrencias— y que las africadas afectan al 3,69% de las formas de los cinco seres frente al 5,63% de las de control, de modo que la decisión no favorece la hipótesis.

Se usa deliberadamente un inventario publicado y ajeno en vez de uno propio. El inventario queda congelado antes de mirar los datos y no puede ajustarse a lo que se quiere ver, que es la prótesis más barata contra la apofenia composicional. Esa defensa no basta por sí sola, porque un inventario publicado puede ser sencillamente el inventario equivocado; por eso el apartado 7 repite la prueba central bajo cinco codificaciones más, una de ellas sin fusionar nada.

Cuando hace falta un grado más de grano grueso se usa la firma {...}: el conjunto de clases presentes, sin orden ni repetición. *M*, *MM* y *MMM* comparten firma {*M*}.

2.3 Los bloques de comparación

Fijados antes de calcular:

- Crianza (5 conceptos): MOTHER, FATHER, BABY, GRANDMOTHER, GRANDFATHER.
- Otro parentesco (7): CHILD, BROTHER, SISTER, SON, DAUGHTER, UNCLE, AUNT. Si el patrón es de la crianza y no del parentesco en general, aquí debe bajar.
- Vocabulario corriente (14): WATER, STONE, FIRE, SUN, EYE, HAND, TREE, TOOTH, BLOOD, NAME, MOON, DOG, HEAD, BONE. Si el patrón saliera también aquí, sería un artefacto del método.
- Campo de la criatura (11), para el apartado 8: los anteriores más BOY, GIRL, GRANDCHILD, CHILD (DESCENDANT), DESCENDANTS, YOUNG MAN, YOUNG WOMAN.

2.4 Guardarraíles

Cuatro precauciones, todas necesarias y una de ellas decisiva.

Una fila por lengua. Las lenguas con varias formas para el mismo concepto se cuentan una vez.

La familia como unidad. El austronesio aporta 956 lenguas al corpus y el indoeuropeo 139. Contar lenguas sobrerrepresenta a las familias grandes, así que toda afirmación fuerte se rehace por familia.

Control de longitud. Una palabra corta tiene más probabilidad de salir monoclasa. Donde la comparación puede contaminarse por eso, se restringe el nulo a las formas del mismo número de consonantes, o se condiciona a las que ya son monoclasa.

Estratificación. Una asociación agregada puede existir sin existir dentro de ningún estrato. Se comprueba con Mantel–Haenszel (1959) siempre que la afirmación sea sobre pares dentro de una misma lengua. En el apartado 5 esta precaución tumba un resultado que sin ella habría pasado.

2.5 Incertidumbre y magnitud

Toda proporción de titular lleva su intervalo de Wilson (1927) al 95%, y todo contraste entre dos proporciones lleva un intervalo de Newcombe (1998) sobre la diferencia y la h de Cohen (1988) como tamaño de efecto, con la lectura habitual de 0,2 como pequeño, 0,5 como medio y 0,8 como grande. Las tablas de contingencia llevan la V de Cramér (1946). La comparación pareada por familias lleva la delta de Cliff (1993), la correlación biserial de rangos pareada y un intervalo por bootstrap sobre familias. Las razones de momios llevan intervalo de Woolf; las estratificadas, intervalo de Robins–Breslow–Greenland (1986) y test de homogeneidad de Tarone (1985). Las tablas completas están en `out/effects.md`.

3. Un concepto, todo el mundo

Antes del repertorio completo, el caso más simple: un solo concepto.

MOTHER en 4.420 formas, 3.787 lenguas y 189 familias. Quitadas las vocales, el mundo se reparte en 471 códigos distintos, pero dos grupos concentran la mayor parte.

La firma {N} reúne 794 lenguas de 65 familias y {M} 624 de 67. Juntas, 1.418 lenguas —el 37,4% del corpus— en 106 familias independientes. La tercera firma, {NT}, cae ya a 182.

Que sean nasales no es un efecto de que las nasales abunden. Contra el vocabulario entero de Lexibank, y comparando solo con palabras del mismo número de consonantes:

código	% en 'madre'	% en el léxico	enriquecimiento	IC 95%	a igual longitud
N	14,5%	1,57%	×9,2	[8,6; 10,0]	×3,3
M	11,2%	1,03%	×10,8	[9,9; 11,9]	×3,9
NN	6,1%	0,61%	×9,9	[8,8; 11,3]	×8,3
MM	5,3%	0,20%	×26,1	[22,8; 30,0]	×21,8
T	3,1%	2,08%	×1,5	[1,3; 1,8]	×0,5
P	2,1%	1,72%	×1,2	[1,0; 1,5]	×0,4
K	1,9%	2,84%	×0,7	[0,5; 0,9]	×0,2

Las oclusivas no es que no estén enriquecidas: están por debajo de lo esperado, y el intervalo de K excluye el 1 por arriba. 'Madre' no es una palabra corriente con nasales de más; es una palabra de la que las oclusivas se retiran.

El control interno es el más limpio posible: la palabra de al lado, en las mismas lenguas.

En esas 3.656 lenguas la firma {N} es 20,4 veces más frecuente en 'madre' que en 'padre' [14,8; 28,2]; {P} es 6,9 veces más frecuente en 'padre' [5,7; 8,3] y {T} 4,2 veces [3,5; 5,0].

Cada familia elige, y elige distinto: el turco va al 83% a {N} (*ana*, con una entropía normalizada de 0,18, casi monolítica), el tungús al 84%, el dogón al 75%; el sino-tibetano al 59% a {M} (*ama*); el indoeuropeo, en cambio, tiene como código más frecuente MTR —*mater*, el compuesto y no la nasal sola— con solo el 29%; y el pama-ñungano no tiene ningún código por encima del 6%.

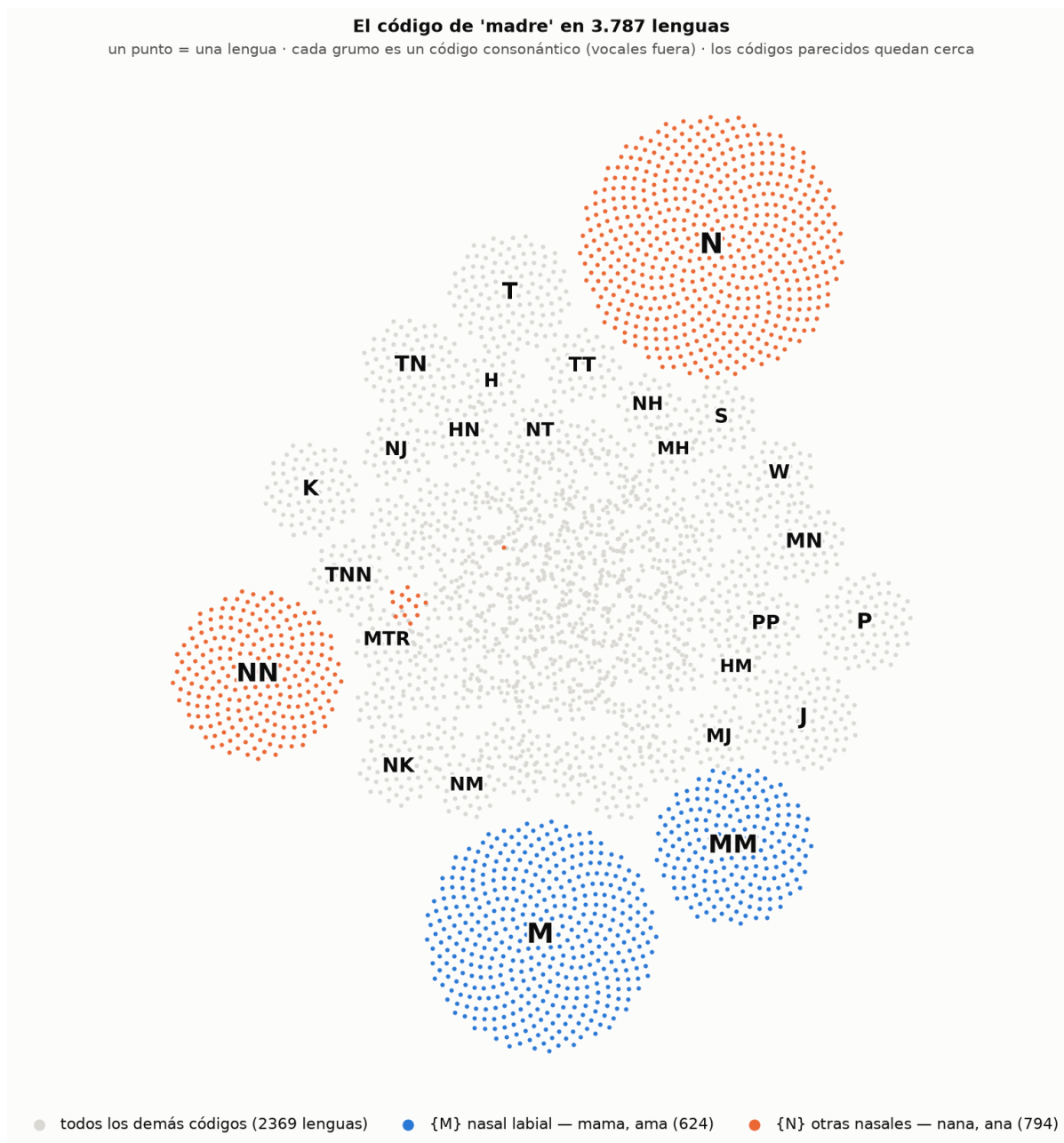


Figure 1: La nube de códigos de 'madre'. Un punto es una lengua; cada grumo, un código consonántico; los códigos parecidos quedan cerca (distancia de edición y escalado multidimensional). Los cuatro grumos mayores son nasales.

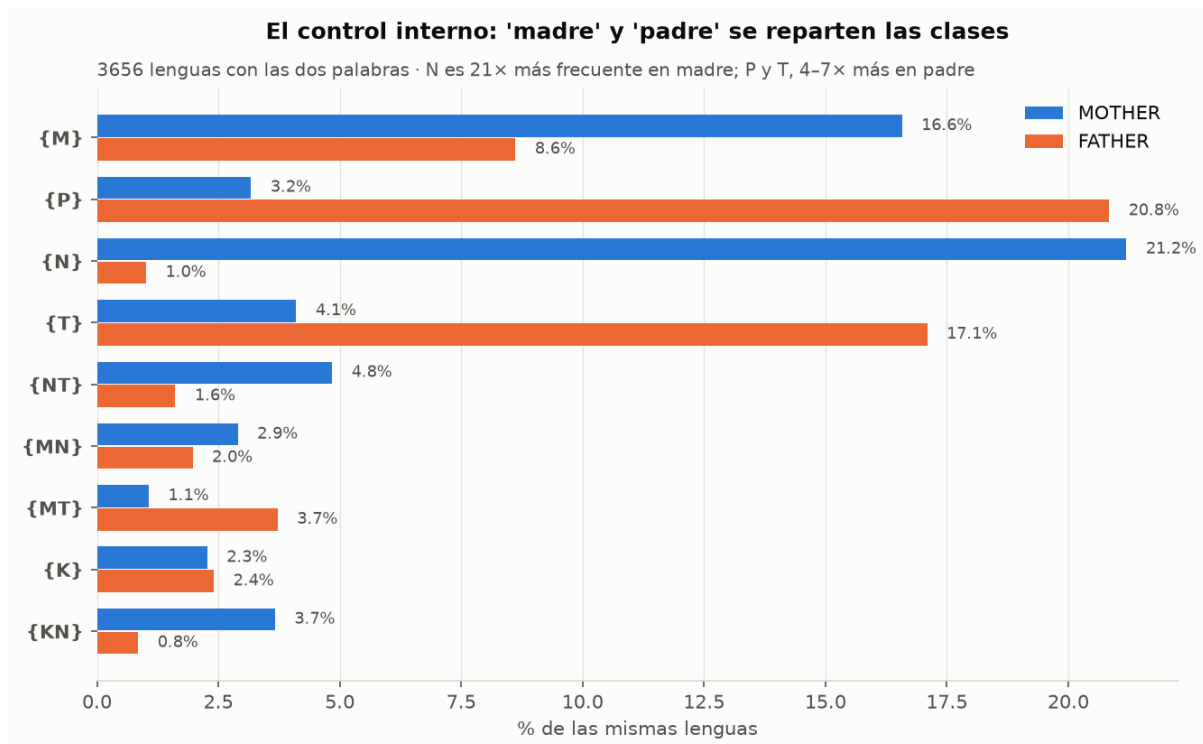


Figure 2: MOTHER contra FATHER en las mismas 3.656 lenguas. Si el procedimiento fabricara el patrón, las dos palabras darían lo mismo. El reparto es casi complementario.

4. El cuarteto

Se generaliza ahora la pregunta a los cinco conceptos de crianza y a las cuatro clases. Un código cuenta como núcleo si su firma es de una sola clase y esa clase está en N, M, P o T.

4.1 De las palabras de una sola consonante, ¿cuál?

Esta es la prueba central, y es la que neutraliza la longitud: condicionada a que la palabra ya sea monoclase, la pregunta es solo *qué clase* elige.

El cuarteto se lleva el 80,9% [79,8; 82,0] en la crianza, el 50,6% [48,1; 53,0] en el resto del parentesco y el 48,0% [47,3; 48,8] en el vocabulario corriente. La diferencia frente al vocabulario corriente es de 32,9 puntos [31,5; 34,1], con una *h* de Cohen de 0,71 —un efecto grande en la lectura convencional. Y el desplazamiento es doble: K cae del 18,1% al 6,7% y S del 10,7% al 2,4%. La crianza no elige cuatro clases al azar entre diez: elige estas y aparta las otras.

Clase por clase, sobre las palabras monoclase de cada bloque:

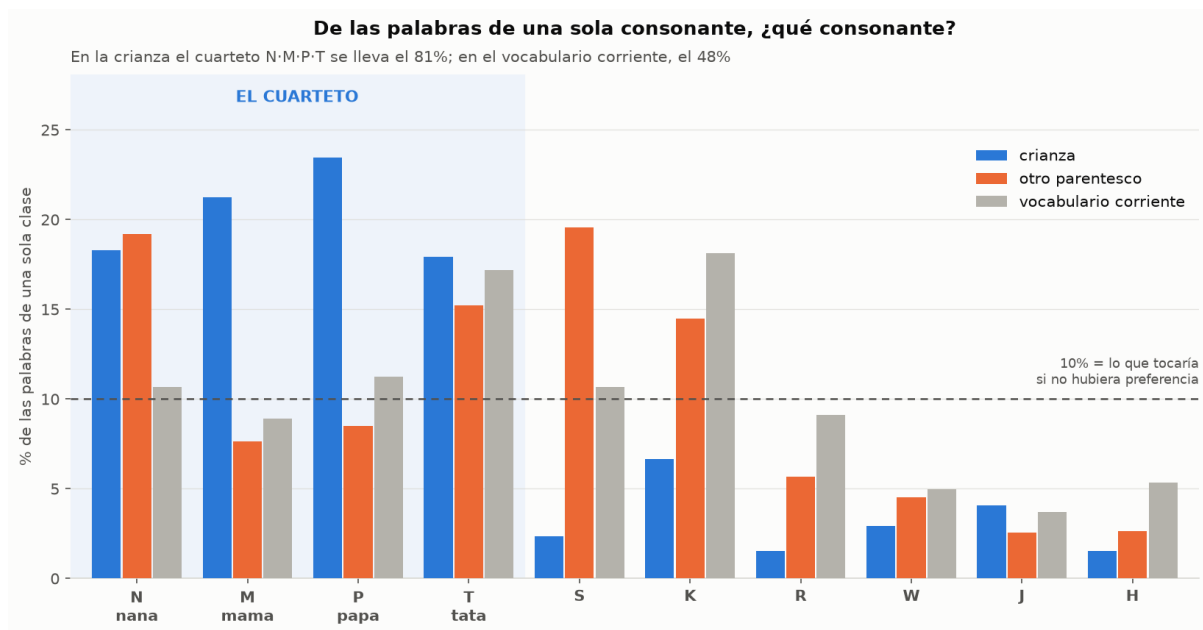


Figure 3: De las palabras formadas por una sola clase de consonante, qué clase es. Con diez clases, el azar daría 10% a cada una y 40% al cuarteto.

clase	crianza	IC 95%	corriente	IC 95%	<i>h</i> de Cohen
P	23,5%	[22,3; 24,6]	11,2%	[10,8; 11,7]	+0,33
M	21,2%	[20,1; 22,4]	8,9%	[8,5; 9,4]	+0,35
N	18,3%	[17,2; 19,4]	10,7%	[10,2; 11,2]	+0,22
T	17,9%	[16,9; 19,0]	17,2%	[16,6; 17,7]	+0,02
K	6,7%	[6,0; 7,4]	18,1%	[17,6; 18,7]	−0,36
J	4,1%	[3,6; 4,7]	3,7%	[3,5; 4,0]	+0,02
W	2,9%	[2,5; 3,4]	5,0%	[4,7; 5,3]	−0,11
S	2,4%	[2,0; 2,8]	10,7%	[10,2; 11,1]	−0,36
R	1,5%	[1,2; 1,9]	9,1%	[8,7; 9,6]	−0,37
H	1,5%	[1,2; 1,9]	5,3%	[5,0; 5,7]	−0,22

Sobre la tabla completa de tres bloques por diez clases, la *V* de Cramér vale 0,231 ($\chi^2 = 2.570$, *gl* = 18): una asociación moderada para una tabla de este tamaño, y concentrada casi entera en seis de las diez filas.

Puestos en una sola escala los 26 conceptos de los tres bloques —los siete del campo de la criatura se incorporan en el apartado 8—, padre y madre encabezan la lista.

Conviene señalar aquí un límite de la métrica más ingenua. La proporción de palabras que caen en el cuarteto *sin condicionar* no separa limpiamente los bloques: FIRE alcanza el 24,1% y supera a GRANDFATHER (21,5%), sencillamente porque 'fuego' también es una palabra corta. Lo que sí separa es qué clase se elige: de las monoclase de FIRE, el 67% son del cuarteto; de las de MOTHER, el 83%.

4.2 Un trío que sube y un cuarto que sobrevive

La tabla clase por clase obliga a corregir el titular, y conviene decirlo con todas las letras porque la cifra agregada lo tapa.

El cuarteto no sube en bloque. M, N y P pasan del 30,9% de las monoclase corrientes al 63,0% de las de crianza, y de ahí sale casi la totalidad de los 32,9 puntos. T no se mueve: 17,2% en el vocabulario

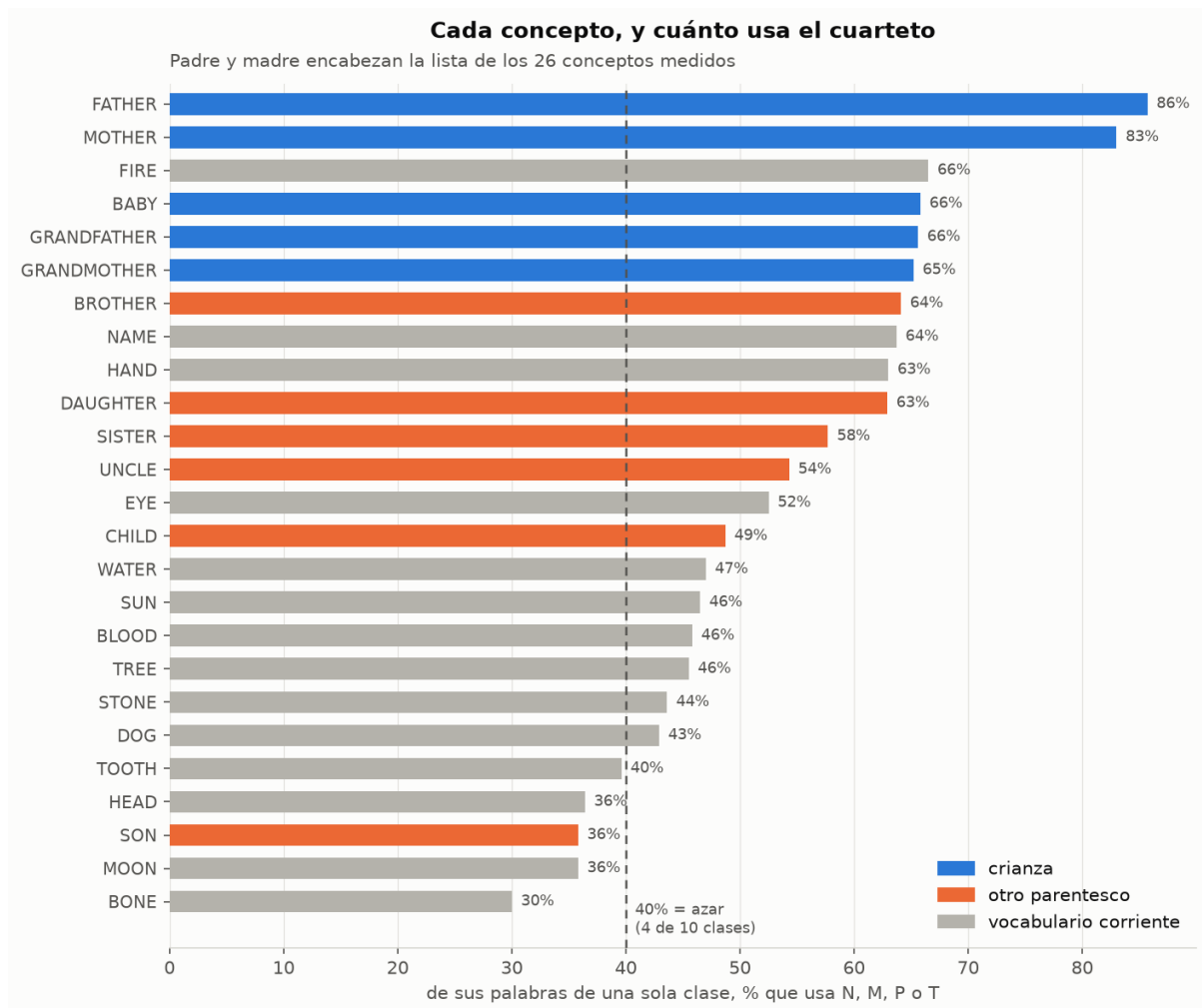


Figure 4: Los 26 conceptos medidos, ordenados por la proporción de sus palabras monoclase que usa el cuarteto.

corriente, 17,9% en la crianza, $h = 0,02$. Con esta evidencia, T no es una clase que el campo de la crianza prefiera. Es una clase que el campo no consigue expulsar, mientras sí expulsa a K ($h = -0,36$), S ($-0,36$), R ($-0,37$) y H ($-0,22$).

El reencuadre importa para saber qué hay que explicar. La pregunta no es «¿por qué la crianza elige T y no K?» —no elige ninguna de las dos— sino «¿por qué la crianza evita K cuando no evita T?». El apartado 10.3 vuelve sobre esto y descarta el candidato más barato.

El cuarteto, entonces, es un inventario real con una costura por dentro: tres clases que el campo recluta y una cuarta que solo tolera. Las dos cosas forman parte del hallazgo.

4.3 La prueba que no depende del tamaño de las familias

Cada familia se compara consigo misma: de sus palabras monoclasa, qué proporción usa el cuarteto en los conceptos de crianza y qué proporción en los de control.

Sobre las 45 familias con al menos diez formas monoclasa en los dos bloques: 78,5% frente a 49,5%, una diferencia media de +29,0 puntos con intervalo bootstrap sobre familias de [+22,0; +35,8] y mediana de +29,3. La crianza va por delante en 40 de 45 familias (89%, IC 76–96%; test de signos $p = 7,9 \cdot 10^{-8}$; Wilcoxon pareado $p = 8,4 \cdot 10^{-10}$). La correlación biserial de rangos pareada vale 0,91 y la delta de Cliff 0,78: no es que unas pocas familias tiren mucho, es que casi todas empujan en la misma dirección. Ninguna familia grande arrastra el resultado, y el apartado 7.2 las va quitando una a una para enseñarlo.

4.4 El repertorio es compartido, no repartido

Es un punto importante y fue una corrección al diseño inicial. La hipótesis no dice que a cada concepto le toque un código en exclusiva, sino que estos conceptos se surten del mismo repertorio.

Lo que cambia entre conceptos es la inclinación, no el inventario: el femenino hacia las nasales (N+M = 84% en MOTHER), el masculino hacia las oclusivas (P+T = 80% en FATHER) y bebé sin preferencia apreciable (25/31/27/17), que es el caso más puro de repertorio compartido. Los perfiles se solapan mucho —BABY·GRANDMOTHER 0,92, FATHER·GRANDFATHER 0,91— y el único par realmente contrastado es MOTHER·FATHER, con 0,36.

Dentro de una misma lengua el repertorio funciona además como sistema de contrastes: en las 1.193 lenguas donde madre y padre tienen ambas código núcleo, el 94% [92,4; 95,1] usa un núcleo distinto para cada una. No es asignación fija —el 17% de los padres-núcleo son M, *mama*— sino una inclinación fuerte dentro del mismo inventario.

5. Un resultado negativo: el eje labial/dental

Los cuatro códigos invitan a una lectura elegante: dos ejes cruzados, nasal/oclusiva separando madre de padre y labial/dental emparejándolos, de modo que la lengua que dice *mama* dijera *papa* y la que dice *nana* dijera *tata*.

Sobre las lenguas con madre nasal y padre oclusiva, la tabla agregada la respalda:

	padre P (papa)	padre T (tata)
madre M (mama)	285	156
madre N (nana)	178	196

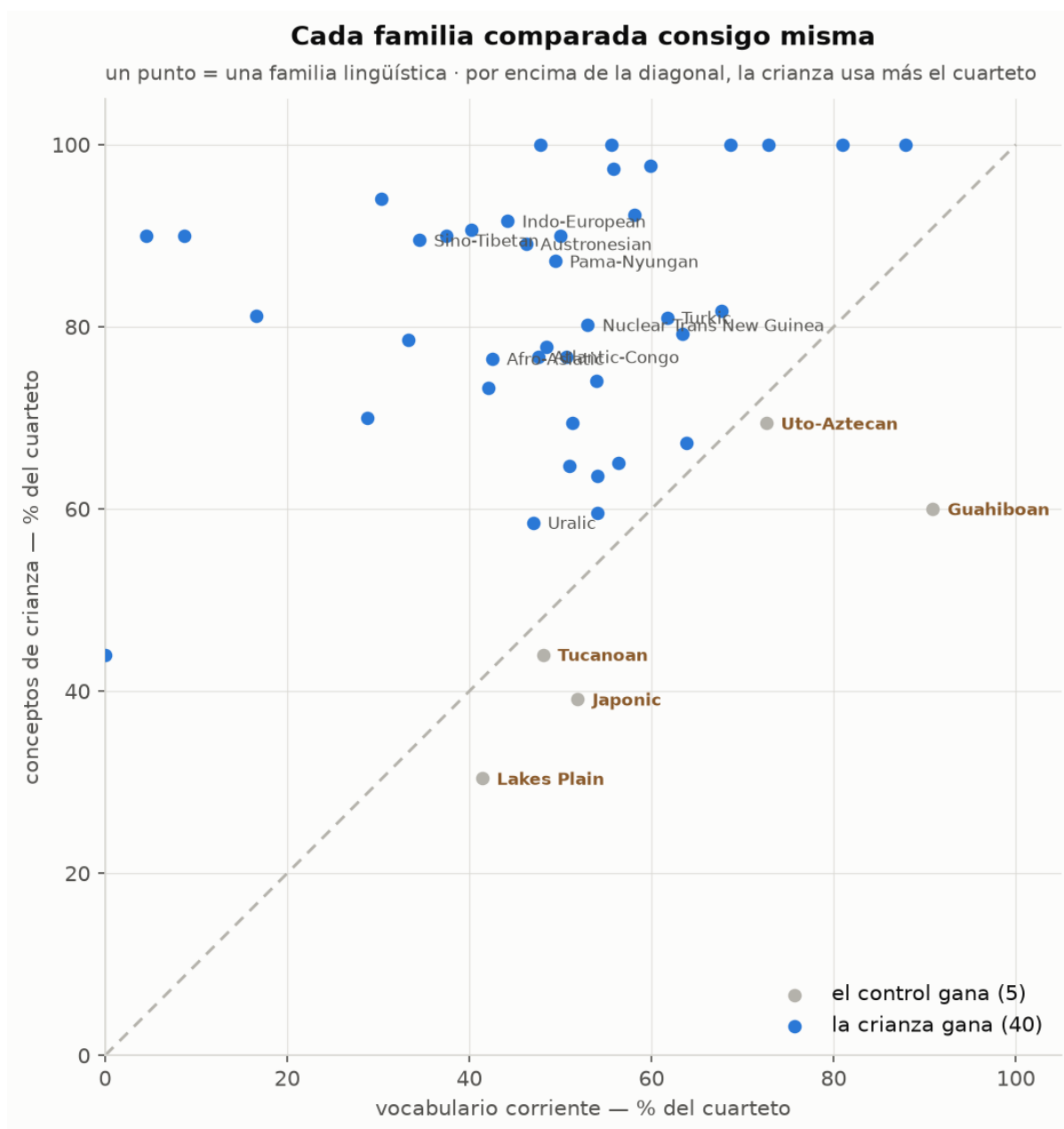


Figure 5: Cada familia comparada consigo misma. Un punto es una familia; por encima de la diagonal, la crianza usa más el cuarteto que el vocabulario corriente de esa misma familia.

Un repertorio compartido, no un código por concepto

los cinco conceptos se surten del mismo cuarteto; lo que cambia es la inclinación — madre hacia las nasales, padre hacia las oclusivas, bebé sin preferencia

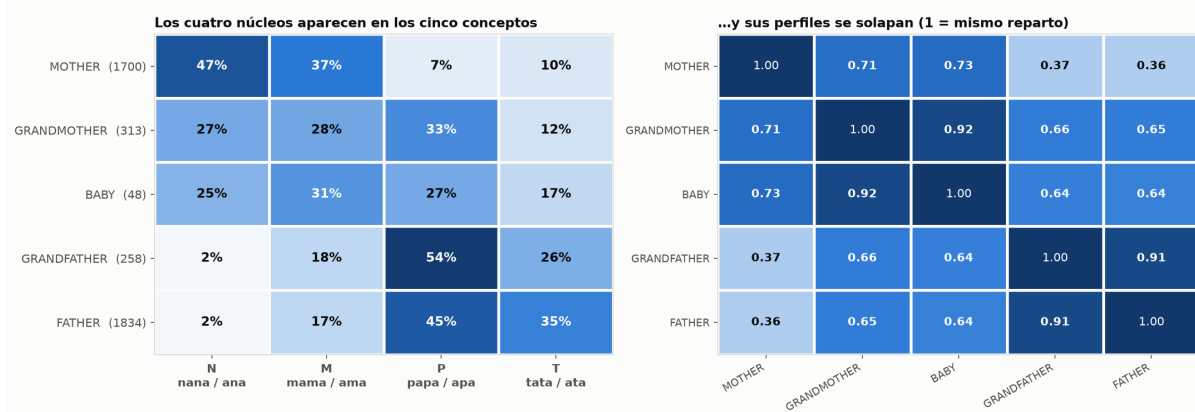


Figure 6: Izquierda: los cuatro núcleos aparecen en los cinco conceptos. Derecha: solapamiento de los perfiles (1 = mismo reparto).

Razón de momios 2,01, intervalo de Woolf [1,52; 2,67] ($p = 1,3 \cdot 10^{-6}$, Fisher). Quien dice *mama* tiene el doble de posibilidades de decir *papa* que *tata*.

La cifra no sobrevive al control. Estratificando por las doce familias con al menos diez casos, la razón de Mantel–Haenszel baja a 1,01, con intervalo de Robins–Breslow–Greenland de [0,65; 1,56] y $p = 0,97$. El intervalo no es simplemente ancho alrededor de una estimación positiva: está casi centrado en el 1. Las razones de momios por familia van de 0,08 a 3,38 con mediana 0,79 —repartidas a los dos lados de la ausencia de asociación— y el test de homogeneidad de Tarone sobre los siete estratos sin casillas vacías da $p = 0,02$: las familias difieren de verdad entre sí.

Esa combinación es la firma exacta de un efecto de composición. No hay una asociación común escondida tras la variación entre familias; hay familias que prefieren el par *mama/papa* y familias que prefieren *nana/tata*, y al juntarlas aparece una correlación que dentro de casi ninguna existe.

El eje nasal/oclusiva es sólido. El eje labial/dental no está probado, y se deja consignado como tal. El intervalo estratificado excluye razones por encima de 1,56, de modo que un efecto del tamaño del 2,01 agregado queda descartado; los efectos menores no, y con doce familias útiles el test no tiene potencia para intentarlo. La afirmación honesta no es «es falso» sino «este corpus no lo sostiene, y tampoco podría detectar una versión débil».

6. El catálogo de clusters

Se invierte ahora la pregunta —no «qué código le toca a este concepto» sino «este código, a qué seres agrupa lenguas»— y se amplía el repertorio a las combinatorias: cualquier código formado solo con M, N, P y T, incluidos MN (*mana*), NT (*nata*) o PT. Los cinco conceptos se agrupan por el ser al que señalan: femenino (madre + abuela), masculino (padre + abuelo) y criatura (bebé).

No se busca universalidad. Un cluster de 639 lenguas en 53 familias es un hecho aunque el resto del mundo haga otra cosa; lo que se cataloga es su tamaño, su dispersión y su geografía.

Con las combinatorias dentro, el repertorio cubre el 49,8% [48,9; 50,8] de las palabras de los cinco seres, frente al 22,8% [22,5; 23,1] del vocabulario corriente y el 17,0% del léxico entero. No es una tasa de cobertura universal: es un contraste de densidad, y es lo que hace que los clusters sean clusters y no ruido de fondo.

Ordenados por cuánto especializa cada código, sale un gradiente limpio:

El catálogo entero: cuántas lenguas en cada cluster			
filas = código consonántico (· puro, ° combinatoria) · columnas = el ser al que acerca			
P ·	129	1	586
N ·	603	5	31
M ·	433	4	182
T ·	131	7	315
MM ·	261	11	181
TT ·	63	1	383
PP ·	90	12	329
NN ·	250	5	15
TM °	19		136
MN °	55	1	48
TN °	78	1	15
TMN °	1		85
NT °	39	2	26
TNN °	50	1	7
NM °	38	1	8
PT °	14	2	29
PN °	7	1	30
MT °	22	3	10
TP °	15	1	17
NP °	10	2	14
	femenino	criatura	masculino

Figure 7: El catálogo completo: cuántas lenguas en cada cluster. Filas, el código consonántico; columnas, el ser al que acerca.

código	se inclina a	fuerza	lenguas	lectura
N	femenino	94%	639	el más especializado del repertorio
NN	femenino	93%	270	acompaña a N
TT	masculino	86%	447	reduplicado y muy marcado
P	masculino	82%	716	el más transversal: toca 61 familias
PP	masculino	76%	431	acompaña a P
M	femenino	70%	619	nasal y labial a la vez
T	masculino	70%	453	marcado, con presencia femenina real
MM	femenino	58%	453	el bisagra

El eje que ordena el repertorio no es labial/dental —ya se ha visto que ese no aguanta— sino nasal frente a oclusiva. Y hay una explicación estructural para los extremos del gradiente: M y MM son nasales, lo que los empuja al femenino, pero labiales, lo que los emparenta con P; esa doble pertenencia es justamente lo que los deja en medio, mientras N, que no tiene la ambigüedad, no baja del 94%.

Merece subrayarse el caso de la labial para el ser femenino —el tipo *baba*, con su prolongación esclava en *babuška*—, porque no es una anomalía suelta: P·femenino reúne 129 lenguas de 28 familias y PP·femenino otras 90 de 16. Más de doscientas lenguas de treinta y tantas familias usan la labial para la madre o la abuela. El repertorio no está asignado.

6.1 Repertorio o herencia: una cifra los separa

Los clusters del repertorio tocan cincuenta o sesenta familias cada uno: N·femenino 53, P·masculino 61, M·femenino 51. Frente a ellos, TM·masculino reúne 136 lenguas pero solo siete familias —127 de ellas austronesias, el reflejo de *tama*—, TMN 85 lenguas en cuatro familias y TNN·femenino 50 en dos, el reflejo de *tina*. Son grandes en lenguas y minúsculos en ramas: eso no es el repertorio actuando, es una herencia viajando por una rama. La cifra que los separa es el número de familias, y conviene aplicarla antes de interpretar cualquier cluster.

7. Robustez: la codificación, las familias y lo que la reducción destruye

Todo lo anterior descansa en una operación: quitar las vocales, sustituir cada consonante por su clase de Dolgopolsky. De ahí salen tres objeciones, y a cada una se le responde aquí con una cifra en vez de con un argumento.

7.1 Seis codificaciones alternativas, un mismo resultado

La prueba central —de las palabras ya formadas por una sola clase, qué proporción usa el cuarteto— se recalculó bajo seis lecturas alternativas: Dolgopolsky recalculado desde el IPA con LingPy, SCA en dos lecturas, ASJP, rasgos articulatorios e IPA crudo sin fusionar. El cuarteto se define cada vez de la misma manera, por articulación y no por la letra que le toque: nasal labial, nasal no labial, oclusiva labial, oclusiva coronal no sibilante. Cada sistema lo traduce a su alfabeto. Cuanto más fina la codificación, más clases tiene y menos regala el azar.

codificación	clases	azar	crianza	parentesco	corriente	diferencia	familias a favor
Dolgopolsky (Lexibank)	10	40%	80,9%	50,6%	48,0%	+32,9	40/45 8,4·10 ⁻¹⁰

codificación	clases	azar	crianza	parentesco	corriente	diferencia	familia a favor
Dolgopolsky (LingPy)	10	40%	80,9%	50,5%	47,9%	+33,0	40/45 $6,5 \cdot 10^{-10}$
SCA, estricto	15	27%	79,6%	50,2%	45,9%	+33,7	40/45 $9,6 \cdot 10^{-8}$
SCA, por lugar	15	40%	81,9%	52,0%	50,3%	+31,6	40/45 $1,2 \cdot 10^{-9}$
ASJP	35	31%	83,0%	48,8%	51,4%	+31,6	42/46 $4,0 \cdot 10^{-9}$
IPA crudo, sin fusionar rasgos articulatorios	201	8%	79,5%	46,8%	47,7%	+31,8	35/40 $1,4 \cdot 10^{-8}$
	16	25%	83,8%	50,4%	52,1%	+31,7	46/49 $2,0 \cdot 10^{-11}$

Tres de estas filas merecen comentario. SCA estricto (List 2012) separa f/v de p/b y θ/ð de t/d, de modo que el cuarteto se hace a propósito más pequeño de lo que Dolgopolsky permite; la diferencia se ensancha en vez de estrecharse. ASJP (Brown *et al.* 2008) está cerca de un sistema de transcripción, con 35 clases atestiguadas. Los rasgos articulatorios se calculan con PanPhon (Mortensen *et al.* 2016) directamente desde lugar y modo, sin ninguna tabla histórica de por medio; bajo esa codificación η es una nasal dorsal y no cuenta para el cuarteto, que es la lectura más estricta de la hipótesis y la que puntúa más alto.

La fila decisiva es la sexta. Con IPA crudo, sin fusionar absolutamente nada —201 consonantes base distintas, donde el azar daría 8% al cuarteto—, el bloque de crianza sigue poniendo el 79,5% de sus palabras de una sola consonante sobre los sonidos del cuarteto, frente al 47,7% del vocabulario corriente. En esa cifra no interviene ningún sistema de clases. Mirado sin tabla alguna, el ranking de las doce palabras de una sola consonante más frecuentes lo enseña directo:

#	crianza	%	corriente	%
1	m	21,2%	t	10,5%
2	n	17,8%	m	9,0%
3	t	12,6%	n	8,7%
4	b	10,6%	k	8,5%
5	p	8,9%	s	6,9%
6	d	5,3%	p	5,3%
7	j	4,1%	l	4,5%
8	k	3,9%	d	4,2%
9	w	2,6%	h	4,1%
10	s	1,6%	j	4,1%
11	h	1,1%	w	4,0%
12	tʃ	1,0%	b	3,9%

Los seis primeros puestos de la columna de crianza son m, n, t, b, p, d. La tabla de Dolgopolsky no los puso ahí.

7.2 Sin las familias grandes

escenario	lenguas	familias	crianza	corriente	diferencia
corpus completo	5.501	298	80,9% [79,8; 82,0]	48,0% [47,3; 48,8]	+32,9
sin indoeuropeo	5.197	297	80,6%	48,2%	+32,4

escenario	lenguas	familias	crianza	corriente	diferencia
sin austronesio	4.523	297	79,6%	48,3%	+31,3
sin ninguna de las dos	4.219	296	79,2%	48,5%	+30,7

Quitando por turno cada una de las diez familias mayores, la diferencia queda entre +29,3 (sin sino-tibetano) y +34,0 (sin tupí). Y el control más duro disponible: sortear una lengua por familia, de modo que una familia de 956 lenguas y una de una sola pesen exactamente igual, 500 veces. La diferencia sale +31,8 puntos, percentiles 2,5–97,5 [+26,1; +37,0], positiva en el 100% de los sorteos.

7.3 Cuánta información destruye la reducción

Quitar las vocales y fusionar consonantes acerca palabras que no estaban tan cerca: *mater*, *mother* y *mutter* acaban compartiendo MTR. Cuánto se pierde se contesta con dos cifras por escalón: cuántos bits de entropía sobreviven, y cuántas veces dos palabras distintas acaban idénticas.

escalón	formas distintas	entropía (bits)	bits que quedan	colisión, crianza	colisión, corriente	razón
forma IPA completa	50.413	14,85	100%	0,12%	0,011%	×10,8
esqueleto de consonantes (IPA)	19.153	11,43	77%	1,04%	0,262%	×4,0
código de Dolgopolsky	6.388	8,81	59%	2,30%	0,871%	×2,6
firma (sin orden)	665	6,83	46%	4,53%	1,593%	×2,8

La pérdida es real y grande: cuando una forma se ha convertido en código de Dolgopolsky se ha ido el 41% de la entropía, y 50.413 formas distintas se han quedado en 6.388. Pero las dos columnas de la derecha dicen algo que la objeción no anticipa. La reducción perjudica más al bloque de control. La razón entre la tasa de colisión de la crianza y la del vocabulario corriente *baja* en cada escalón —de ×10,8 en la forma IPA completa a ×2,6 en el código de Dolgopolsky— porque el vocabulario corriente es el que más gana al fusionarse.

Lo mismo aparece en las fusiones falsas. De cada cien parejas que el código de Dolgopolsky declara idénticas, 55 tienen esqueletos de consonantes IPA distintos en el bloque de crianza y 70 en el de control. La fusión es más generosa con 'piedra' que con 'madre'.

Así que la reducción no fabrica el contraste; si algo hace, lo erosiona. En el nivel de la forma IPA completa, antes de quitar una sola vocal, dos palabras de crianza sacadas al azar ya coinciden diez veces más a menudo que dos corrientes.

8. La criatura: un resultado negativo y firme

El punto más frágil de la hipótesis inicial era el tercero: que la criatura de brazos —*nene*, *bebé*— participara del repertorio. Con el concepto BABY solo se quedaba sin material: 601 lenguas y un código medio de 3,50 consonantes. Se rehízo con tres cambios de diseño.

La unidad de conteo pasa a ser la lengua, no el concepto. La pregunta correcta no es «¿BABY usa el repertorio?» sino «¿esta lengua tiene *alguna* palabra de la criatura en el repertorio?». Si 'bebé' es un compuesto pero existe *nene* como CHILD, la lengua dispone de la palabra.

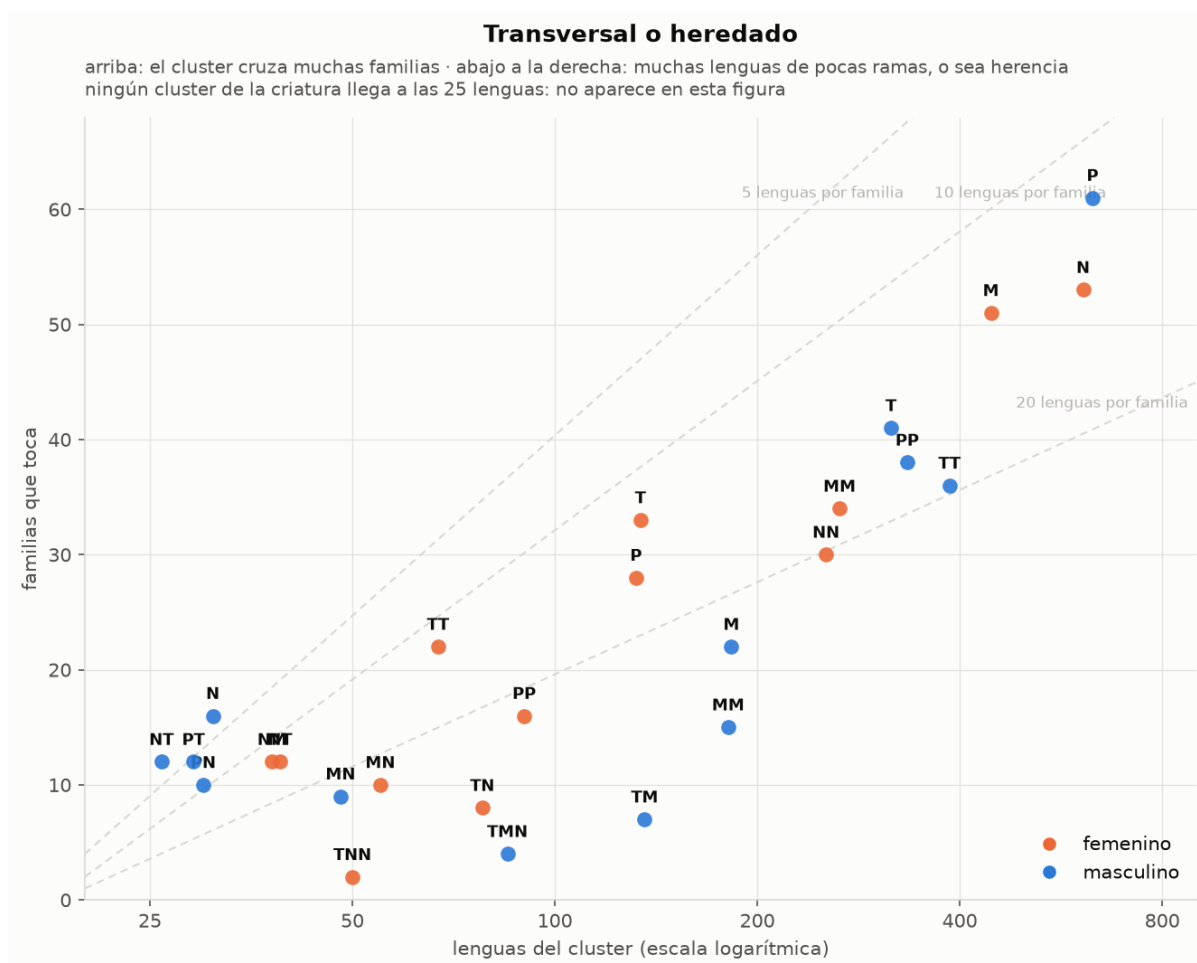


Figure 8: Transversal o heredado. Arriba, clusters que cruzan muchas familias; abajo a la derecha, muchas lenguas de pocas ramas.

Los seis clusters mayores, sobre el mapa

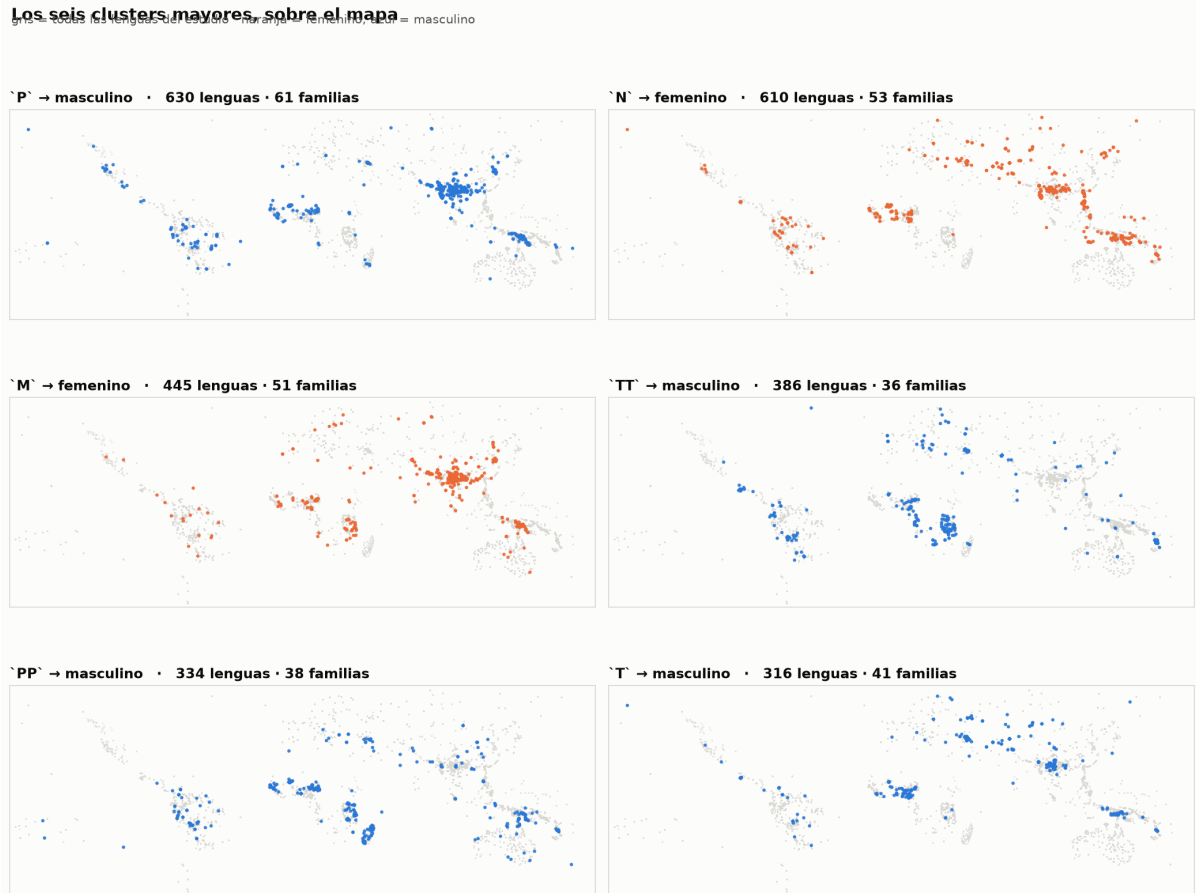


Figure 9: Dónde vive cada uno de los seis clusters mayores. En gris, todas las lenguas del estudio.

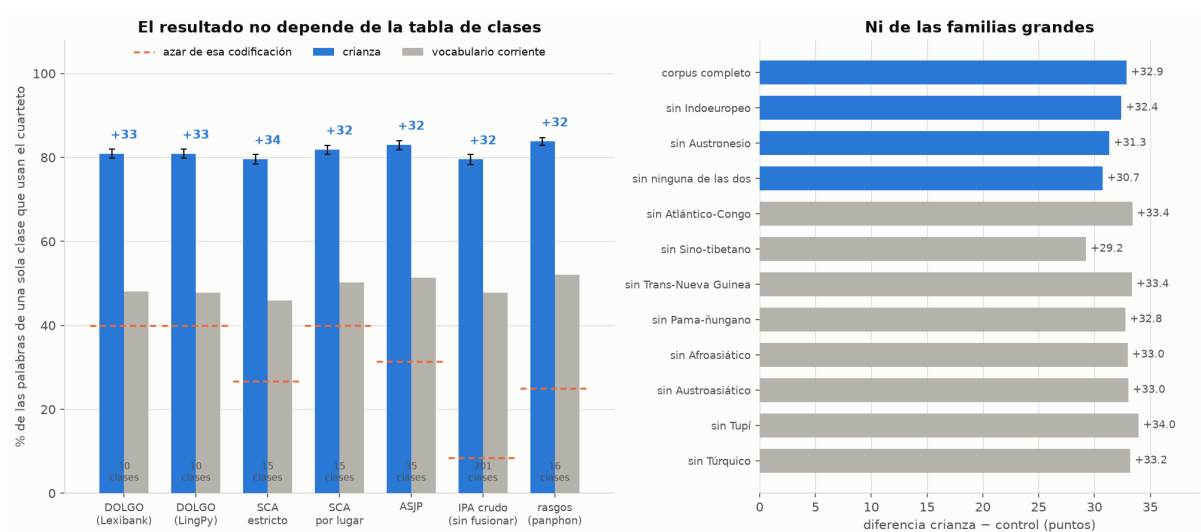


Figure 10: Izquierda: la cifra central bajo siete lecturas, cada una con su nivel de azar; las barras de error son intervalos de Wilson al 95%. Derecha: la misma cifra quitando familias enteras del corpus.

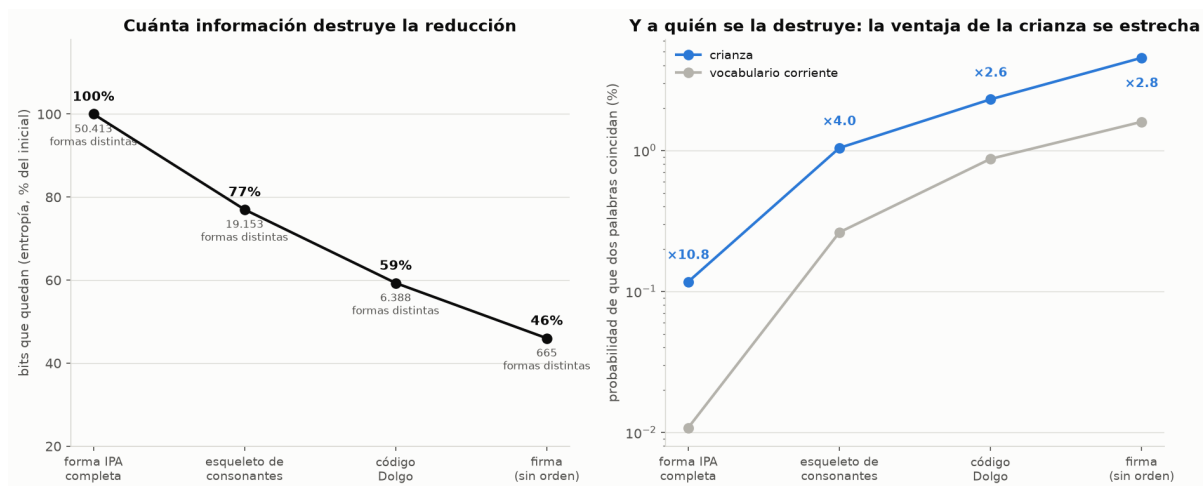


Figure 11: Izquierda: entropía que queda en cada escalón de la reducción. Derecha: probabilidad de que dos palabras del mismo bloque coincidan, en cada escalón, en la crianza y en el vocabulario corriente.

Los intentos se igualan. Una lengua con seis palabras de criatura tiene seis oportunidades de caer en el repertorio; se le dan las mismas seis al control, sorteando seis de sus conceptos corrientes, trescientas veces. Sin esta corrección la métrica sube sola.

El campo se abre en tres capas, porque 'niño' es una edad y 'hijo' es una relación: criatura de brazos (BABY, CHILD, CHILD·DESC), niño/niña (BOY, GIRL, GRANDCHILD) y descendencia (SON, DAUGHTER, DESCENDANTS, YOUNG MAN, YOUNG WOMAN).

El resultado es negativo en las tres capas. Sobre 4.563 lenguas, 28,6% [27,3; 29,9] frente al 38,7% del control igualado (percentiles 2,5–97,5 del remuestreo: 37,7–39,9): diez puntos por debajo, con una h de Cohen de $-0,21$. Los dos intervalos no se rozan por ningún lado: es un negativo firme, no un empate mal medido. Y no es falta de datos —BOY tiene 1.778 lenguas y 228 familias, más que 'madre'.

No hay cluster de la criatura

once conceptos · 4.563 lenguas · la lengua como unidad y los intentos igualados contra el vocabulario corriente de cada lengua

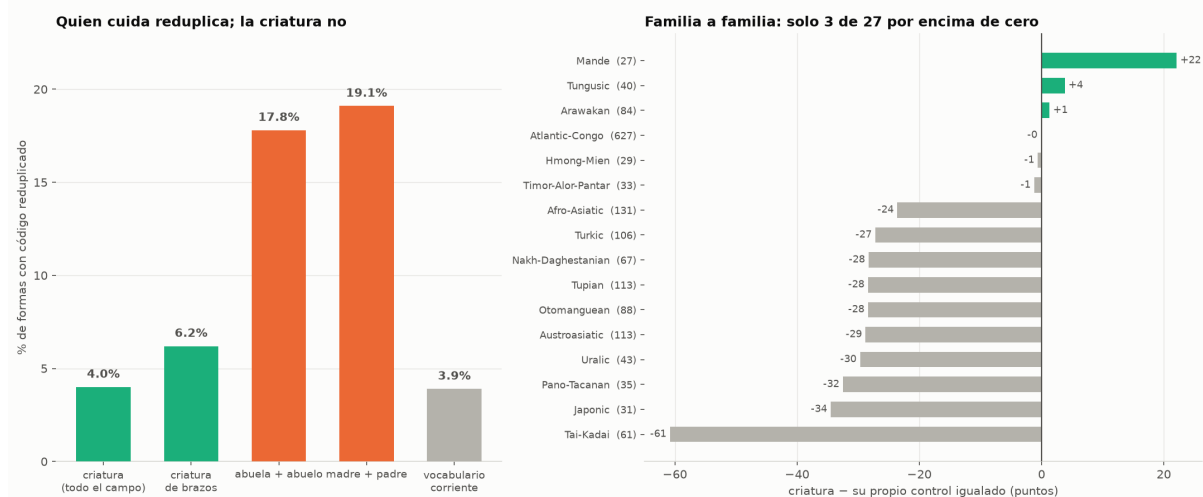


Figure 12: Izquierda: la reduplicación por grupos. Derecha: cada familia contra su propio control igualado.

Tampoco es un fenómeno indoeuropeo que el corpus diluya: el indoeuropeo queda noveno de 27 familias y también por debajo de su propio control (45% contra 49%). Solo 3 de 27 familias tienen más

repertorio en la criatura que en su vocabulario corriente, y la única con efecto claro es el *mandé* (+22 puntos, 27 lenguas).

Se comprobó además una hipótesis alternativa: que lo propio de la criatura no fuera el cuarteto sino la reduplicación, con la consonante que fuera. Tampoco. Madre y padre reduplican el 19,1% y los abuelos el 17,8% —cinco veces más que el vocabulario corriente—, mientras la criatura se queda en el 4,0% contra el 3,9% del control.

La razón está en el largo del código. Madre y padre rondan 1,9 consonantes; la criatura va de 2,6 a 4,2. Las lenguas no nombran a la criatura con una palabra primitiva: la derivan —‘el pequeño de’, ‘el hijo de’, ‘el que nace’— y una palabra derivada no cabe en un código de dos consonantes. Dicho de otro modo: el repertorio de la crianza es de quien cuida, no de quien es cuidado.

8.1 Lo que sí queda: la transferencia

Hay, sin embargo, un efecto real y en el lugar exacto que predice la intuición psicodinámica —que los padres reciclan sus propios términos hacia el niño: *papi*, *papito*, *nene*.

Cuando la criatura sí usa el repertorio, tiende a usar el mismo código que los padres de su misma lengua: sobre 949 lenguas, 18,7% de coincidencia [16,3; 21,3] frente al 14,5% esperado al barajar los códigos de criatura entre lenguas dentro de cada familia (2.000 permutaciones; percentiles 2,5–97,5 del nulo: 12,9–16,1; $p = 0,0005$). La h de Cohen vale 0,11 —real, y pequeño, y así conviene leerlo. El código no le pertenece a la criatura: le llega prestado desde quien la nombra. La transferencia existe; lo que no existe es un cluster autónomo.

9. La laguna: no hay corpus del registro afectivo

Este apartado no reporta un resultado sino una ausencia, y es probablemente lo más importante que este trabajo tiene que decir sobre sus propios límites.

Todo lo medido hasta aquí es el término de referencia —‘la madre’, ‘el niño’—, que es lo que una lista comparativa recoge por diseño. Lo que la madre efectivamente *dice* a la criatura en brazos pertenece a otro registro, el vocativo y afectivo, y no está documentado de forma comparable para la inmensa mayoría de las lenguas tratadas aquí.

La magnitud de la laguna se puede dar con una cifra. De los 1.740.092 registros de Lexibank —no de la muestra de este trabajo, sino del repositorio entero— exactamente nueve llevan alguna marca de vocativo o de término de trato. Nueve de un millón setecientos mil. Y esas nueve bastan para mostrar que no se trata de un matiz sino de otra palabra:

lengua	término de referencia	término de trato
huntergatherer-258 (MOTHER)	<i>deni</i> → TN	«term of address is mama» → MM
huntergatherer-353 (FATHER)	<i>jum-me</i> → KM	<i>pa</i> : (vocativo) → P
huntergatherer-353 (MOTHER)	<i>manko</i>	<i>ma</i> : (vocativo) → M
huntergatherer-163 (MOTHER)	<i>niatija</i>	«originally vocative, reanalyzed as referential»

En los tres primeros casos el término de trato cae en el repertorio y el de referencia no. El cuarto es aún más elocuente: un vocativo que se congeló en término de referencia, es decir, el mecanismo

mismo por el que estas palabras entran al léxico.

La consecuencia metodológica es directa: las cifras de este trabajo son un piso y no un techo. Se ha medido sistemáticamente el término que menos favorece la hipótesis, porque es el único que el corpus recoge.

Se exploró si Wiktionary podía suplir la laguna. Los volcados de Kaikki (Ylönen 2022) conservan el registro en el texto de la glosa —*mamá* aparece con «a term of affection for a woman»— lo que permite cruzar un patrón de registro con uno de referente. El barrido de 105 volcados locales devolvió 366 entradas en 23 lenguas, de las que 182 apuntan a la criatura.

No sirve para tipología, y conviene decir por qué con precisión. De las 105 lenguas disponibles localmente, 104 son indoeuropeas —el volcado se había descargado para un estudio de esa familia— y el inglés aporta por sí solo el 41% de los aciertos. Además son palabras de diccionario (*poppet*, *honnepon*, *chiquitín*), no habla espontánea. Como señal, y sin nulo comparativo, la primera consonante de esos 182 términos cae en el cuarteto el 60,9% de las veces, con P dominante (41,9%); se consigna como pista, no como resultado.

La fuente adecuada existe y es otra: CHILDES / TalkBank (MacWhinney 2000), con 48 lenguas y 436 corpus de acceso público, formado por transcripciones de habla real entre adultos y niños. El experimento pendiente está bien definido: extraer los vocativos que la madre dirige al niño, codificarlos por el mismo procedimiento y compararlos con el término de referencia de esa misma lengua en Lexibank, en un contraste pareado por lengua. Los cuatro casos de la tabla anterior predicen la dirección del resultado.

10. Discusión

Lo que estos datos describen no es un universal ni un parentesco, sino un repertorio en el sentido fijado en 1.1: un inventario pequeño de códigos consonánticos que un número grande de lenguas sin relación entre sí emplea para los mismos seres, con inclinaciones estadísticas y no con reglas.

Las tres propiedades que lo definen se midieron por separado, y las tres se sostienen.

Es un inventario cerrado y expulsivo. Cuatro clases de diez, con K, S, R y H desplazadas activamente, y el desplazamiento es tan grande en tamaño de efecto como el reclutamiento. Eso es lo que distingue el hallazgo de «las palabras de crianza tienen más nasales».

Es compartido y no repartido. Los cinco seres se surten del mismo inventario, con solapamientos de perfil de hasta 0,92. Lo que los distingue es una inclinación —nasal para el femenino, oclusiva para el masculino— que admite excepciones masivas: doscientas lenguas llaman a la madre con labial. Un cluster no es una regla.

Discrimina por gradiente y no por casilla. De N con el 94% a MM con el 58% hay un continuo, y su extremo bajo tiene una explicación fonológica interna: M participa a la vez de la nasalidad y de la labialidad, y por eso no distingue.

10.1 La convergencia, y por qué el hallazgo no es Jakobson repetido

La objeción natural a todo lo anterior es que se trata de un corolario de lo ya sabido. Si el balbuceo produce labiales y corales, y quien cuida oye el balbuceo como interpelación, entonces una acumulación translingüística de M, N y P en los términos parentales no es un descubrimiento sino una inevitabilidad: el equivalente fonológico de observar que las cosas pesadas caen.

La objeción tiene parte de razón y conviene conceder esa parte. La *dirección* del efecto la predicen Jakobson (1960) y todo lo que vino después: la descripción de Ferguson (1964, 1978) del *baby talk*

como subsistema, los trabajos de Locke (1983) y Vihman (2014) sobre lo que el lactante puede efectivamente producir en el primer año, y el modelo *frame/content* de MacNeilage y Davis (2000), donde la oscilación mandibular con la lengua quieta da como salidas por defecto del balbuceo los cierres en los labios y en la cresta alveolar. Nadie necesitaba este trabajo para esperar *mama*.

Lo que no estaba predicho es la forma. Cuatro puntos de los datos no son corolarios de la explicación articulatoria.

El rechazo es más fuerte que la preferencia. Una explicación articulatoria predice que los sonidos fáciles estén sobrerrepresentados. No predice que los *difíciles* sean expulsados de un campo semántico mientras siguen perfectamente disponibles en palabras de la misma longitud y la misma forma — véase 10.3.

El inventario está compartido, no repartido. Una presión articulatoria convergente actuando por separado sobre cinco conceptos produciría cinco distribuciones parecidas, que es lo que se observa; pero no da ninguna razón para el sistema de contrastes *dentro de la lengua* del apartado 4.4, donde el 94% de las lenguas que usan códigos núcleo para los dos padres usa uno *distinto* para cada uno. Eso es un hecho léxico sobre paradigmas, no un hecho articulatorio sobre bocas.

El único lugar donde la explicación articulatoria predice el efecto con más claridad es justo donde falla. Si el repertorio fuera simplemente balbuceo capturado, la criatura —el ser más próximo al balbuceo— debería ser su usuaria más densa. La criatura queda diez puntos por debajo de su propio vocabulario corriente (apartado 8). El repertorio es de quien cuida y no de quien es cuidado, y esa asimetría no la genera una historia articulatoria por sí sola.

Las magnitudes son nuevas. «Suena parecido en muchas lenguas» y «80,9% frente a 48,0%, sosteniéndose en 40 de 45 familias y también en IPA crudo» son afirmaciones distintas, y la segunda es la que se puede comparar con el fenómeno siguiente. Blasi *et al.* (2016) hicieron el movimiento análogo para los sesgos sonido–significado en vocabulario básico, y Wichmann *et al.* (2010) para ASJP; este trabajo lo hace para el campo de la crianza, con la familia como unidad.

El marco mayor al que esto pertenece es la literatura sobre la no arbitrariedad —iconicidad, sistematicidad y simbolismo fónico (Hinton, Nichols y Ohala 1994; Perniss, Thompson y Vigliocco 2010; Dingemanse *et al.* 2015). El repertorio de la crianza es un caso raro dentro de ese marco, porque la motivación no es un parecido perceptivo entre sonido y referente sino *la producción del propio referente*: el sonido no retrata a la criatura, la cita. Si eso se clasifica mejor como iconicidad, como indexicalidad o como ninguna de las dos, es una pregunta que estos datos pueden informar pero no cerrar.

10.2 Qué falsaría la lectura de repertorio

Como «repertorio» e «inevitabilidad articulatoria» hacen muchas de las mismas predicciones, conviene nombrar dónde se separan, para que la afirmación siga siendo falsable.

Si el repertorio es un objeto *léxico*, debería comportarse como tal: debería sobrevivir en el registro vocativo con más fuerza que en el referencial (el apartado 9 lo predice y CHILDES puede comprobarlo); debería mostrar el sistema de contrastes dentro de la lengua también en lenguas cuya fonología corriente no se lo pone fácil; y debería recuperarse en lenguas cuya ruta del balbuceo al léxico es culturalmente distinta —lenguas de signos con articulaciones bucales, o lenguas con sistemas de parentesco fuertemente supletivos.

Si, en cambio, es convergencia articulatoria y nada más, entonces el vocativo no debería mostrar *ningún* enriquecimiento sobre el término referencial, el contraste madre/padre debería disolverse al controlar la frecuencia, y la evitación de K de 10.3 debería resultar un artefacto de algo que no hemos medido. Cada una de estas es un experimento real, y aquí no se ha corrido ninguno.

10.3 Por qué K no

El cabo suelto más persistente del trabajo se formulaba antes como «por qué T y no K». El apartado 4.2 muestra que la pregunta estaba mal puesta: T no está preferida en absoluto ($h = 0,02$). La pregunta real es por qué se evita K, siendo como es el código monoclase más frecuente del vocabulario corriente (18,1%) y cayendo aquí al 6,7% —una h de $-0,36$, la misma magnitud que S y R.

La explicación candidata más barata es que lo haga la forma de la palabra. Las palabras de crianza son cortas y a menudo reduplicadas, y si la reduplicación misma desfavoreciera las dorsales, todo el patrón se seguiría de la forma y no del campo semántico. Esa explicación falla, y falla limpio. Medido sobre el léxico entero de Lexibank, sin la crianza a la vista:

clase	de los códigos de una consonante	de los reduplicados XX	razón
N	11,3%	16,7%	$\times 1,48$
K	20,4%	25,7%	$\times 1,26$
T	14,9%	18,7%	$\times 1,25$
R	9,7%	9,4%	$\times 0,97$
P	12,3%	9,8%	$\times 0,80$
H	5,0%	3,9%	$\times 0,78$
M	7,4%	5,5%	$\times 0,74$
S	10,8%	6,6%	$\times 0,61$
W	4,7%	2,3%	$\times 0,50$
J	3,6%	1,3%	$\times 0,36$

En el vocabulario corriente una palabra reduplicada es *más* probablemente velar, no menos, y K se reduplica exactamente igual de bien que T. *Kaka* es una forma perfectamente ordinaria para una palabra ordinaria. La evitación de K es específica del campo semántico, no una consecuencia de la forma fonológica que ese campo prefiere.

Quedan cuatro hipótesis vivas, y se separan por experimento más que por argumento.

Orden de adquisición. Los cierres velares exigen elevar el cuerpo de la lengua de forma independiente, y aparecen más tarde y con menos fiabilidad en el balbuceo canónico que los cierres labiales y corales, que caen solos del marco mandibular (MacNeilage y Davis 2000; Vihman 2014). Si el repertorio es balbuceo capturado, lo que el balbuceo no produce de forma fiable no se puede capturar. Esto predice que la fuerza de la evitación de K debería seguir la variación translingüística en la producción temprana de velares.

Saliencia perceptiva y visual. El cierre labial se ve; el coronal se ve a medias; el dorsal no se ve. En la interacción cara a cara del habla dirigida al niño (Fernald y Simon 1984; Kuhl 2004; Saint-Georges *et al.* 2013), un nombre elegido bajo mirada mutua se elige bajo una restricción de visibilidad que el resto del léxico no tiene. Esto predice un orden con lo labial primero, que es lo que los datos muestran (P y M son las dos mayores subidas).

Contraste de registro. Si el registro de crianza funciona en parte por *no* sonar como el vocabulario corriente, entonces la clase más corriente es la peor candidata, y el 18,1% que K ocupa entre las monoclase corrientes es justo lo que la descalificaría. Esto predice que las clases evitadas sean las frecuentes —encaja con K y con S, y encaja peor con R.

Asociación simbólica. Las consonantes dorsales y sibilantes reaparecen en el extremo áspero de las escalas de simbolismo fónico (Ohala 1994; Hinton *et al.* 1994), lo que las desfavorecería en un registro afectivo con independencia de la articulación. Esto predice que las mismas clases deberían evitarse en otros vocabularios afectivos, y es una afirmación directamente comprobable sobre un campo que no hemos tocado.

Nada en estos datos elige entre las cuatro, y el trabajo no pretende otra cosa. Lo que los datos sí establecen es que el hecho a explicar es una evitación y que no se reduce a la forma de la palabra.

10.4 La asimetría de la abuela

Queda una pregunta más que es de este trabajo y no heredada. El abuelo hereda el perfil del padre (54% P, solapamiento 0,91) pero la abuela no hereda el de la madre: se reparte entre tres núcleos y su vecino más próximo es BABY (0,92), no MOTHER (0,71). Nada de lo disponible predice esa asimetría, y queda abierta.

11. Limitaciones

El código se calcula sobre la palabra entera. No se separan morfemas, de modo que una raíz del repertorio con un sufijo encima se sale de la cuenta: el ruso *babuška* da PPSK en lugar del PP de *baba*. Medida por igual en los tres bloques, esta pérdida deja fuera un +6,5% en los cinco seres frente a un +4,5% en el control, lo que confirma que la deuda no es simétrica y que los clusters están subestimados.

El corpus no trae prosodia ni acento. Nada de lo dicho aquí toca la entonación, que en el registro de crianza es un componente central según Ferguson y la dimensión mejor documentada del habla dirigida al niño (Fernald y Simon 1984). Un repertorio de códigos consonánticos es, como mucho, una capa del registro.

La cobertura por familia es desigual. El austronesio aporta 956 lenguas y hay familias con menos de veinte. Por eso las afirmaciones fuertes se dan por familia, y el apartado 7.2 muestra que el resultado sobrevive tanto a quitar las familias grandes como a un remuestreo de una lengua por familia. La potencia de los tests *estratificados* sigue siendo limitada, en cambio —doce familias útiles en el test del eje labial/dental—, y es ahí donde la limitación muere.

El registro afectivo no está. Es la limitación mayor y se ha desarrollado en el apartado 9.

El nivel del fenómeno. Todo lo medido pertenece a la adquisición y la articulación. No es evidencia de una firma genealógica profunda ni debe leerse como tal.

12. Conclusión

Existe un repertorio de cuatro códigos consonánticos —N, M, P, T— que un número grande de lenguas sin relación entre sí emplea para los seres del núcleo de crianza. Se lleva el 80,9% de las palabras monoclasa de ese campo frente al 48,0% del vocabulario corriente, una diferencia de 32,9 puntos con una *h* de Cohen de 0,71, y la ventaja se sostiene dentro de 40 de 45 familias comparadas consigo mismas. Sobrevive a seis codificaciones alternativas, incluida el IPA crudo sin fusionar nada, y sobrevive a quitar las dos familias mayores y a un remuestreo que le da a cada familia el mismo peso.

Dentro de ese resultado hay una costura que conviene arrastrar hacia adelante: el cuarteto no sube en bloque. M, N y P están reclutadas; T solo tolerada; y K, S y R expulsadas con tamaños de efecto tan grandes como los del reclutamiento. Lo que hay que explicar es una evitación.

El repertorio es compartido entre los conceptos, no repartido, y se ordena por un gradiente de especialización que va del 94% de N al 58% de MM, con el eje nasal/oclusiva separando el ser femenino del masculino.

Dos lecturas atractivas no sobreviven a los controles y se consignan como negativas: los cuatro códigos no forman dos ejes cruzados —el eje labial/dental es un efecto de composición, con una razón de momios estratificada de 1,01 [0,65; 1,56] y heterogeneidad real entre familias— y la criatura no tiene cluster propio, sino un código prestado de sus padres.

Y una ausencia condiciona todo el conjunto: no existe corpus comparativo del registro afectivo para la inmensa mayoría de las lenguas tratadas. Nueve registros de 1.740.092. Lo que aquí se ha medido

es el término de referencia; el término de trato —lo que la madre efectivamente dice— está fuera del alcance de los datos disponibles, y los pocos casos documentados indican que caería más dentro del repertorio, no menos.

Referencias

- Blasi, D. E., Wichmann, S., Hammarström, H., Stadler, P. F., y Christiansen, M. H. (2016). Sound-meaning association biases evidenced across thousands of languages. *Proceedings of the National Academy of Sciences*, 113(39), 10818–10823.
- Brown, C. H., Holman, E. W., Wichmann, S., y Velupillai, V. (2008). Automated classification of the world's languages: A description of the method and preliminary results. *STUF – Language Typology and Universals*, 61(4), 285–308.
- Cliff, N. (1993). Dominance statistics: Ordinal analyses to answer ordinal questions. *Psychological Bulletin*, 114(3), 494–509.
- Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences* (2.^a ed.). Hillsdale, NJ: Lawrence Erlbaum.
- Cramér, H. (1946). *Mathematical Methods of Statistics*. Princeton: Princeton University Press.
- Dingemanse, M., Blasi, D. E., Lupyan, G., Christiansen, M. H., y Monaghan, P. (2015). Arbitrariness, iconicity, and systematicity in language. *Trends in Cognitive Sciences*, 19(10), 603–615.
- Dolgopolsky, A. B. (1964). Gipoteza drevnejšego rodstva jazykovyx semej Severnoj Evrazii s verojatnostnoj točki zrenija [Hipótesis probabilística sobre el parentesco más antiguo de las familias lingüísticas de Eurasia septentrional]. *Voprosy Jazykoznanija*, 2, 53–63.
- Ferguson, C. A. (1964). Baby talk in six languages. *American Anthropologist*, 66(6, parte 2), 103–114.
- Ferguson, C. A. (1978). Talking to children: A search for universals. En J. H. Greenberg (ed.), *Universals of Human Language*, vol. 1. Stanford: Stanford University Press, 203–224.
- Fernald, A., y Simon, T. (1984). Expanded intonation contours in mothers' speech to newborns. *Developmental Psychology*, 20(1), 104–113.
- Hammarström, H., Forkel, R., Haspelmath, M., y Bank, S. *Glottolog*. Max Planck Institute for Evolutionary Anthropology. <https://glottolog.org>
- Hinton, L., Nichols, J., y Ohala, J. J. (eds.) (1994). *Sound Symbolism*. Cambridge: Cambridge University Press.
- Jakobson, R. (1960). Why 'mama' and 'papa'? En B. Kaplan y S. Wapner (eds.), *Perspectives in Psychological Theory: Essays in Honor of Heinz Werner*. Nueva York: International Universities Press.
- Kuhl, P. K. (2004). Early language acquisition: cracking the speech code. *Nature Reviews Neuroscience*, 5, 831–843.
- List, J.-M. (2012). SCA: Phonetic alignment based on sound classes. En M. Slavkovik y D. Lassiter (eds.), *New Directions in Logic, Language and Computation*. Berlín: Springer, 32–51.
- List, J.-M., Cysouw, M., y Forkel, R. (2016). Concepticon: A resource for the linking of concept lists. *Proceedings of LREC 2016*, 2393–2400.
- List, J.-M., Forkel, R., Greenhill, S. J., Rzymiski, C., Englisch, J., y Gray, R. D. (2022). Lexibank, a public repository of standardized wordlists with computed phonological and lexical features. *Scientific Data*, 9, 316. <https://doi.org/10.1038/s41597-022-01432-0>
- Locke, J. L. (1983). *Phonological Acquisition and Change*. Nueva York: Academic Press.

- MacNeilage, P. F., y Davis, B. L. (2000). On the origin of internal structure of word forms. *Science*, 288(5465), 527–531.
- MacWhinney, B. (2000). *The CHILDES Project: Tools for Analyzing Talk* (3.^a ed.). Mahwah, NJ: Lawrence Erlbaum.
- Mantel, N., y Haenszel, W. (1959). Statistical aspects of the analysis of data from retrospective studies of disease. *Journal of the National Cancer Institute*, 22(4), 719–748.
- Mortensen, D. R., Littell, P., Bharadwaj, A., Goyal, K., Dyer, C., y Levin, L. (2016). PanPhon: A resource for mapping IPA segments to articulatory feature vectors. *Proceedings of COLING 2016*, 3475–3484.
- Murdock, G. P. (1959). Cross-language parallels in parental kin terms. *Anthropological Linguistics*, 1(9), 1–5.
- Newcombe, R. G. (1998). Interval estimation for the difference between independent proportions: comparison of eleven methods. *Statistics in Medicine*, 17(8), 873–890.
- Ohala, J. J. (1994). The frequency code underlies the sound-symbolic use of voice pitch. En L. Hinton, J. Nichols y J. J. Ohala (eds.), *Sound Symbolism*. Cambridge: Cambridge University Press, 325–347.
- Perniss, P., Thompson, R. L., y Vigliocco, G. (2010). Iconicity as a general property of language: evidence from spoken and signed languages. *Frontiers in Psychology*, 1, 227.
- Robins, J., Greenland, S., y Breslow, N. E. (1986). A general estimator for the variance of the Mantel–Haenszel odds ratio. *American Journal of Epidemiology*, 124(5), 719–723.
- Saint-Georges, C., Chetouani, M., Cassel, R., Apicella, F., Mahdhaoui, A., Muratori, F., Laznik, M.-C., y Cohen, D. (2013). Motherese in interaction: at the cross-road of emotion and cognition? *PLoS ONE*, 8(10), e78103.
- Tarone, R. E. (1985). On heterogeneity tests based on efficient scores. *Biometrika*, 72(1), 91–95.
- Vihman, M. M. (2014). *Phonological Development: The First Two Years* (2.^a ed.). Chichester: Wiley-Blackwell.
- Wichmann, S., Holman, E. W., y Brown, C. H. (2010). Sound symbolism in basic vocabulary. *Entropy*, 12(4), 844–858.
- Wilcoxon, F. (1945). Individual comparisons by ranking methods. *Biometrics Bulletin*, 1(6), 80–83.
- Wilson, E. B. (1927). Probable inference, the law of succession, and statistical inference. *Journal of the American Statistical Association*, 22(158), 209–212.
- Ylönen, T. (2022). Wiktextextract: Wiktionary as machine-readable structured data. *Proceedings of the 13th Conference on Language Resources and Evaluation (LREC 2022)*, 1317–1325. Datos en <https://kaikki.org>

Reproducibilidad

Todo el material está en `endolanguage/mother/`. Los datos primarios provienen de `data/lexicon/lexibank/forms.csv`.

paso	script	salida
extracción de los 33 conceptos, sin vocales	<code>build.py</code>	<code>data/forms.csv</code>
el nulo sobre todo Lexibank apartado 3	<code>baseline.py</code> <code>analyze.py, plots.py</code>	<code>data/baseline_codes.csv</code> <code>out/findings.md, figuras 1–5</code>

paso	script	salida
apartados 4–5	nursery.py, plots_ nursery.py	out/nursery.md, figuras 6–9
apartado 6	groups.py, plots_groups. py	out/groups.md, figuras 10–12
apartado 7	robustness.py, plots_ robustness.py	out/robustness.md, figuras 14–15
apartado 8	child.py	out/child.md, figura 13
apartado 9	endearment_scan.py	data/endearment_raw.csv
intervalos y tamaños de efecto (todos los apartados)	effects.py	out/effects.md, out/ effects.csv

`robustness.py` requiere LingPy (List *et al.* 2022) y PanPhon (Mortensen *et al.* 2016); `effects.py` requiere SciPy y statsmodels. Los cálculos que usan aleatorización llevan semilla fija. El único paso que depende de recursos externos al repositorio es el del apartado 9, que lee los volcados de Kaikki en `data/lexicon/kaikki/dict/`.