

Consonantal codes in 27 Uralic languages

The family where the instrument can finally do what it promises

Alejandro Toledo Martínez · ORCID 0009-0000-1277-9697 · Integrative Corpus · 2026

Abstract

We present a catalogue of 1,486 consonantal codes grouping 13,470 lexical forms from 27 Uralic languages. A code is a consonant skeleton shared by words for the same concept across distinct languages. No reconstruction, no expert cognacy, and no hypothesis about Proto-Uralic is used.

This is the series' third paper and the first in which the family has the structure the method needs: nine branches, six of them with sufficient data, against Turkic's two and Austronesian's six very unequal ones. That makes it possible, for the first time, to do what the protocol promised and could not test: computing substitutions per branch pair over thirty-six pairs instead of one.

Four results. The codes of widest reach — $n \cdot m$ | name with 20 languages and $378 \times$, $v \cdot r$ | blood, $k \cdot l$ | fish, $m \cdot n$ | I, $t \cdot n$ | THOU — are the ones any Uralicist will recognise at once, obtained without consulting any reconstruction. The search for sibling codes — pairs for the same concept differing in one position — partitions languages into groups using no genealogical label, and its disjointness filter turns out to be indispensable: a quarter of the pairs share languages and therefore partition nothing. Geography here yields a monotonic gradient across all six buckets, which Turkic could not measure for want of range. And automatic affix stripping again makes the result worse — ten points of grouping and eighteen of null — confirming in a second agglutinative family what the Turkic paper found.

With three families measured by the same code, the chance ratio orders by time depth — Turkic $216 \times$, Uralic $78 \times$, Austronesian $37 \times$ — but observed precision does not. That is a datum, not a thesis. (*With the eight families the series now has, the ordering does not hold: see the note in §5.6.*)

1. What is sought, and why Uralic

Historical linguistics has a powerful and expensive instrument: the comparative method. It has been applied with enormous unevenness — Indo-European has dense reconstruction and expert cognacy, most families do not — so that any study leaning on the historical layer of the lexicon measures, without meaning to, the density of the comparative literature rather than the past of the languages.

The series this work belongs to asks whether something useful can be done by giving up reconstruction: not improving or replacing it, but doing without it and seeing what remains. The Austronesian paper showed that a catalogue of consonantal codes remains; the Turkic paper asked what its figures mean.

This paper asks how far the instrument goes when nothing stands in its way. The Turkic paper declared three limits that turned out to belong to the family and not the method: its corpus distinguishes only two branches — one with a single language — so the rule of computing substitutions per branch pair was tested on one pair; the geographic analysis came out degenerate, with 467 of 496 pairs in one bucket; and the sibling-code search arrived at the end, with nothing to check it against.

Uralic lifts all three restrictions. Nine distributed branches, a geographic spread of thousands of kilometres without being an archipelago, and an established primary division.

And here one must be precise, because it is easily misread. Those nine branches are not an input to the method: the catalogue is built exactly as in the other two families, without knowing they exist. What the branches allow is checking what the catalogue produced — and that is why Uralic is a good family for this paper: not because the instrument needs them, but because we need something to check it against.

1.1 The method, and where it is described

The procedure — what a code is, in what order it is built, what nulls it uses and what beating them authorises — is described once for the whole series in “The catalogue of consonantal codes”, and is not repeated here. The essentials, so that what follows can be read:

The catalogue is built without genealogy. Codes come from a form’s consonant skeleton, the concept it expresses, and a threshold of three distinct languages. No branch classification, no proto-forms, no cognate lists. The procedure does not know this family has branches, and genealogy enters at the end only to read an already-closed catalogue.

Nothing is reconstructed, no kinship is claimed, inheritance is not distinguished from borrowing, and no meaning is attributed to a code.

Where this procedure sits in the automatic-comparison literature — Dolgopolsky, Turchin/Peiros/Gell-Mann, ASJP — is covered in the method paper.

2. What a code is, and what it costs in Uralic

2.1 The consonant skeleton

The skeleton is a word’s sequence of consonants, with the vowels removed — not lost: they stay in a separate column — and each consonant written with a canonical representative. Uralic uses 25 canonical symbols, of which 23 merge something. The full table is in `runs/normalizacion.md` and is not written by hand: it is derived by aligning each form with its skeleton.

What matters when reading the results:

Written	as	Consequence in the data
s š ś ʃ ɕ z ʒ ʒ and every affricate eye’ gathers <i>silmä</i> , <i>šalmä</i> , <i>šel’me</i>	s	‘s·l·m
n ŋ ɲ ɳ name’ gathers <i>nimi</i> , <i>ɲim</i> , <i>nam</i>	n	‘n·m
l ɭ ɮ k, q, x are kept separate v and w are kept separate	l	which is why §5.2 can see k·l / q·l / x·l which is why it can see v·r / w·r

That **k**, **q** and **x** are not merged is the table’s most productive decision in this family, and it is worth declaring before the results exploit it: it is what lets the Ob-Ugric and Samoyedic velar change appear as three distinct codes instead of one (§5.2). In Turkic, the opposite decision about sibilants destroyed a diagnostic correspondence.

2.2 A code is a skeleton recurring across languages

A code is a skeleton appearing, for the same concept, in three or more distinct languages. It is written $n \cdot m$ | name, never $n \cdot m = \text{name}$: the code does not mean "name", it is the shape "name" repeatedly takes.

2.3 The instrument's cost: collision

59% of Uralic forms share their skeleton, within their own language, with a form of another concept. It is the clean measure of what the skeleton fails to distinguish, with no genealogy involved.

The clearest case is $k \cdot l$, appearing with "fish" and with "tongue". Finnish alone shows it: *kala* and *kieli*. Two words no description would confuse, and the skeleton writes them alike because the difference lies entirely in the vowels, which is what it discards. Likewise $t \cdot l$ with "fire" (*tuli*) and "winter" (*talvi*), and $n \cdot l$ with "four" (*neljä*) and "arrow" (*nuoli*).

Four pairs of distinct words under one code among those the appendix shows. It is not a defect more data would fix: it is the arithmetic of reducing a word to two or three consonants, and in Uralic it bites despite its words being the longest of the series.

2.4 Arity

Arity	Codes	
1 (unary)	166	11.3%
2 (binary)	652	44.2%
3 (ternary)	523	35.5%
4 or more	133	9.0%

The figures this paper defends are restricted to the 1,186 binaries and ternaries.

Figure 1: Distribution of codes by arity. 91% have three consonants or fewer; Uralic is the series' longest-word family, at 3.27 consonants per form.

The mean arity of Uralic forms is 3.27 consonants, the highest of the series' three families (Turkic 2.85, Austronesian 2.79). That matters for reading §5.6.

3. Method

3.1 Data

Integrative Corpus, Uralic subset: 27 languages, 80,032 forms with transcription, assigned concept and skeleton. Sources: Lexibank, IDS and NorthEuraLex, under CC-BY and CC-BY-SA licences.

The branches, which are what this family contributes:

Branch	Catalogued forms	Branch	Forms
Finnic	2,819	Samoyedic	374
Saami	1,581	Khantyic	285
Permian	1,379	Mansic	276
Mordvin	532	Hungaric	145
Mari	434		

Nine branches, the largest at 36% of the forms. It is the most balanced distribution of the three families: in Turkic one branch had a single language; in Austronesian, Malayo-Polynesian took 95%.

One exclusion, and of the kind that justifies phase 0. The corpus had a language called Xincheng Beigeng, with 1,008 forms, marked as Finnic. It is not Uralic: it is Kam-Sui (Tai-Kadai), from southern China. Three independent checks flagged it without anyone looking for it — it sits 6,446 km from the family's centre (23.8° N / 108.6° E, Guangxi), it marks tone in digits when no other Uralic language does, and its forms are $\theta a:m^{453}$ "three" and δem^{214} "water". Excluded and declared in `ci/decisiones.py`.

The raw corpus holds 99,255 forms and we publish 80,032. The difference is 18,113 forms without a skeleton — spread evenly across all languages, around 28% of each — plus the 1,008 excluded. It is not a script problem: Uralic has 25 Cyrillic forms in total. Those forms simply never got segmented into IPA. We checked that zero concepts exist only in forms without a skeleton, so the loss costs no coverage, but the published count is that of the forms that enter.

A single macroarea, hence no geographic-anomaly section.

3.2 Enrichment

The code's coverage divided by its base frequency; the full definition, and the warning that it is not comparable across families, are in the method paper.

3.3 Nulls

Every grouping is compared against a permutation shuffling concepts within each language, preserving each one's inventory. The real grouping covers 33.6% of the forms; the null, 5.0%.

3.4 One population for the whole paper

Every report under `runs/` declares its population in its header, and they are identical:

```
POBLACIÓN · Uralic: 27 lenguas · 80,032 formas · 1,843 conceptos · 9 ramas ·
despiece=sí(0) · excluidos=1 · umbral=3 lenguas
```

It is a mechanical check (`analysis/tools/verifica_paper.py`, check 7) and it exists because the Turkic paper had four sections computed over different populations without anyone noticing until review.

4. Results

4.1 The catalogue

1,486 codes over 867 of the 1,843 concepts, grouping 13,470 forms.

Magnitude	Value	Over what population
Forms inside a code	16.8% (13,470 of 80,032)	threshold of 3 languages, the catalogue's
Forms in some group	33.6% (null: 5.0%)	threshold of 2 languages, grouping test
Precision	6.50% (chance: 0.083%)	pairs of forms from different languages
Times above chance	78× [78–79; chance is computed exactly]	idem
Binaries and ternaries	1,186 (79.8%)	of the 1,486 codes

Correction of 2026-09-09. The coverage row read 9.8% (7,825 of 80,032), and that figure was not what its name promises: it counted distinct (language, spelling) pairs — the rows of the per-form TSV, which deduplicates — rather than forms. The catalogue's coverage over forms is 16.8%. The same error was in the Turkic paper; the other six use the correct figure. It is still the lowest coverage of the three families, which is what this paragraph uses.

Each row carries its population because there are three different ones. The catalogue's coverage (16.8%) is the lowest of the three families, with an explanation that is not a defect: with 1,843 concepts and only 27 languages, most concepts do not reach three languages with the same form.

4.2 The codes of widest reach are the known etyma

Concept	Code	Languages	Enrichment	Branches it appears in
name	n · m	20	378×	Finnic, Saami, Permian, Khantyic, Mansic, Samoyedic, Mari
blood	v · r	19	343×	Finnic, Saami, Permian, Mordvin, Mari, Hungaric
fish	k · l	19	174×	Finnic, Saami, Mordvin, Mari, Khantyic, Mansic, Samoyedic
I	m · n	18	194×	Finnic, Saami, Permian, Mordvin, Khantyic, Samoyedic
four	n · l	17	204×	Finnic, Saami, Permian, Mordvin, Mari
egg	m · n	17	183×	Finnic, Saami, Permian, Mordvin, Mari
sea	m · r	16	239×	Finnic, Saami, Permian, Mordvin, Samoyedic
fire	t · l	15	155×	Finnic, Saami, Permian, Mordvin, Mari

This is the catalogue's output, no more. The branch column describes how each code's languages are distributed; it does not claim that distribution has a historical cause, which is a different question and not this method's.

Figure 2: The twenty binary and ternary codes of widest reach, by number of languages.

Figure 3: Distribution of reach: how many codes are shared by how many languages.

A Uralicist will recognise several of these codes at once, and that is all the check needed: the procedure saw no reconstruction, consulted no etymological dictionary, and grouped forms by their skeleton.

5. What this family makes possible

5.1 Thirty-six branch pairs, instead of one

The Turkic paper established that the substitution matrix must be computed per branch pair, because aggregated over the family it erases precisely what defines the genealogy: there, Chuvash rhotacism went from 1.4 times chance to 6.1 when restricted. But Turkic had only one possible cut, with a single language on one side. The rule was stated in the plural and tested in the singular.

Uralic gives thirty-six pairs with data. The family's primary cut — Samoyedic against the rest — yields this:

Segment	Main	Cases	Share crossing	Share within	Diagnostic
t	s	283 of 970	29%	9%	3.3×
c	t	26 of 44	59%	22%	2.7×
r	s	210 of 759	28%	16%	1.8×
ʋ	n	26 of 95	27%	0%	only when crossing
x	k	93 of 276	34%	26%	1.3×
q	k	82 of 299	27%	28%	1.0×
b	p	64 of 229	28%	63%	0.4×
z	s	58 of 188	31%	77%	0.4×

The deciding column is the last, and it reads better here than in Turkic because there is contrast in both directions. **t s** at 3.3× is diagnostic of the cut: when two languages on the same side differ in that position, t corresponds to s 9% of the time; crossing the cut, 29%. By contrast **b p** and **z s** come out below 1.0, that is, they are *more* frequent within a branch than across: they are internal alternations, and the column exposes them without anyone needing to know Uralic.

That capacity to rule out is what Turkic could not show, because with a one-language branch there was no "within" to compare against.

And a warning inherited from Turkic: the "cases" are pairs of forms, not independent observations.

5.2 Sibling codes: what the catalogue partitions, and what only looks like it

The Turkic paper found that searching the catalogue for pairs of codes for the same concept differing in a single position yields groups of languages partitioned without any genealogical label. Here it is a declared phase (`analysis/fases/codigos_hermanos.py`) and carries a filter it lacked there.

The filter. For a pair to partition languages into two groups it must be disjoint: if a language appears in both codes, it partitions nothing. Obvious put that way, and it was not when the first version of this section was written, which showed five examples of which three failed it.

Of the 153 pairs of binaries and ternaries: 114 disjoint and 39 overlapping (25%).

Overlapping does not mean wrong, and classifying them all as error was the other failure of the first version. A language having two forms for a concept is normal:

Cause	Pairs	What it is
Another form	24	the language has another word — Finnish "cover" = <i>kattaa</i> and <i>peittää</i> — or the same stem with another suffix — Finnish "here" = <i>täällä</i> and <i>tässä</i> , Komi "near" = <i>matin</i> and <i>matıs</i> . Facts of the language
Transcription	15	the same word written two ways by two sources: Estonian <i>nahk</i> / <i>naxk</i> , <i>täht</i> / <i>tæxt</i> , Khanty <i>ku:l</i> / <i>qul</i> . This one is an artefact

The two separate by edit distance. What cannot be separated without a morphological analyser is the first category from within: in a family with fifteen cases, two forms of the same stem with different suffixes are constant, and calling them all "synonymy" would be a lie. They are left together and it is said.

The disjoint pairs of widest reach, with the partition they produce:

Concept	Code A	Code B	Group A	Group B
name	l · m (4)	n · m (20)	Mordvin, Mari	Finnic, Saami, Permian, Khantyic and others
FIR	k · s (9)	k · z (8)	Saami, Finnic, Mari	all of Permian, Mordvin, two Finnic
nose	n · n (12)	n · r (5)	Finnic, Saami	all of Permian, Mari
NEST	p · s (8)	p · z (8)	Saami, Finnic	Finnic, Mordvin, Permian

Two of the four put all of Permian on one side. That is the observation, and this paper goes no further: the catalogue partitioned those languages that way, knowing no classification. What the partition means is a question for someone with the family's historical phonology to hand.

One case worth telling, because it is the instrument splitting rather than gathering. $t \cdot n$ | THOU has ten languages and not one Finnic: Finnish says *sinä* and lands in $s \cdot n$, where it shares a code with $s \cdot n$ | VEIN. In a single entry sit the gain and the cost: the catalogue isolates a group and at the same time merges two words.

What must be withdrawn from the first version. It claimed that $k \cdot l/q \cdot l/x \cdot l$ for "fish" splits the family into three blocks. False: Khanty and Mansi have forms in all three, and the classifier marks them as transcription, which is what they are — *ku:l*, *qul*, *χul* are the same Khanty word written by three sources. Likewise $v \cdot r/w \cdot r$ for "blood", with Estonian, Finnish and Hungarian in both.

The protocol lesson, corrected. The search works, but not without the filter: a quarter of the pairs partition nothing, and they are the most striking ones. And the breakdown of those 39 is a result in itself: fifteen are different transcriptions of the same word, a measure of how many incompatible sources live in the corpus.

5.3 Automatic stripping again makes the result worse

The Turkic paper found that discovering affixes automatically raises grouping and raises the null more. Uralic, also agglutinative and also exclusively suffixing, confirms it:

Step	In group	Null	Advantage	κ	Codes of 4+ consonants
Whole word	33.6%	5.0%	28.7	0.301	36.1%
+ discovered affixes	43.7%	23.3%	20.4	0.266	7.3%

The κ column is the Hubert-Arabie correction, $(\text{observed} - \text{null}) / (1 - \text{null})$: it normalises by the attainable ceiling and so allows comparison across families where an advantage in points does not — ten points over a 5% null and ten over a 44% null are not worth the same. Here $\Delta\kappa = -0.035$. All eight families in the series give the same sign: Turkic -0.009 , Afro-Asiatic -0.004 , Sino-Tibetan -0.020 , Uralic -0.035 , Indo-European -0.094 , Pama-Nyungan -0.102 , Trans-New-Guinea -0.118 , Austronesian -0.151 .

(This list had six families and the figures of the time. It was rebuilt with all eight on 2026-09-09 from *pipeline/tools/tabla_serie.py*. The magnitude does NOT grow with the raw gain, which is what this paragraph used to say: Sino-Tibetan, with the shortest word of the eight, falls less than Indo-European.)

Ten points of signal, eighteen of noise. The advantage over chance falls from 28.7 to 20.4, however spectacular the quaternary column. And the same typological proof as in Turkic: the discoverer finds 367 prefixes in a family that has no prefixes. It has not discovered morphology, it has discovered the consonant inventory.

That this happens identically in two unrelated agglutinative families, with different corpora and sources, is what turns the Turkic finding into a property of the procedure rather than an anecdote about one family.

The published catalogue does not use automatic stripping.

5.4 The emergent classes do not appear here either

Transitively closing the substitutions that hold up in the validation half puts all twenty segments reaching the threshold into a single class. Not even the {d, t} that Turkic detached.

That is why Appendix A carries no emergent-code column and the TSVs carry it empty. Two consecutive families in which this phase produces nothing usable, and it is worth saying: phase 6 of the protocol worked in Austronesian and has not worked since. Either Austronesian was a fortunate case, or the procedure needs something it does not have.

5.5 Geography, here, works

In Austronesian, languages sharing more codes were closer together. In Turkic it could not be known: with 32 languages of a young family, 467 of the 496 pairs fell into one bucket and the design had no dynamic range.

Uralic gives the complete relation, with all six buckets populated:

Codes shared	Pairs	Mean distance	÷ chance
1	2	3,512 km	2.34
2	6	2,675 km	1.79
3–5	35	2,303 km	1.54
6–10	59	2,094 km	1.40
11–20	85	1,640 km	1.09
21+	162	1,066 km	0.71
<i>two random languages</i>	40,000	1,498 km	1.00

Figure 4: Reach against enrichment. Each point is a binary or ternary code; the vertical axis is logarithmic.

Monotonic across all six buckets, from 2.34 to 0.71. And by groups, compactness also decays with reach: 0.48 for codes of 3–5 languages, 0.64 for 6–10, 0.76 for 11–25.

The small buckets hold two and six pairs, so their extreme values sustain nothing on their own; what sustains the reading is the monotonicity across the six, and that the three large buckets — 59, 85 and 162 pairs — already show it.

This closes the question the Turkic paper left open. There we wrote that geography “says less, it does not say nothing”, unable to tell whether the signal was weak or the design was bad. With Uralic it is clear it was the design: in a family with range of reach, the relation appears cleanly.

5.6 Three families, and what three points allow

With the series' three families measured by the same code and the same definition (`ci/metricas.py`):

Family	Lgs	Arity	Precision	Chance	× chance	Conventional depth
Turkic	32	2.85	25.52%	0.118%	216×	2,000 years
Uralic	27	3.27	6.50%	0.083%	78×	4,000–6,000
Austronesian	582	2.78	13.45%	0.363%	37×	5,200

The chance ratio orders by depth: $216 > 78 > 37$, exactly from shallowest to deepest. Observed precision does not: Austronesian (13.45%) exceeds Uralic (6.50%) while being deeper.

The Turkic paper first argued that the ratio measured time, and its reviews overturned that because two thirds of the difference from Austronesian came from the denominator. With the third point the situation is more interesting and still unresolved: the ordering appears in the ratio and not in the numerator, which is exactly what one would expect if the denominator were doing real work — correcting for skeleton specificity in each family — rather than merely adding noise.

But three points are still not a curve, and at least three variables covary with depth and are uncontrolled: the number of languages (32, 27, 582), arity (2.85, 3.27, 2.78) and source type. The honest course is to leave it there: the ordering is a fact about three points, and `docs/PERFIL-FAMILIAS.md` holds sixty families profiled with the same metric where it can be checked. That is the next task, and we have not done it.

Note of 2026-09-09. The series now has eight families measured with the same code (`pipeline/tools/tabla_serie.py`), and the three-point ordering does not survive them: Afro-Asiatic gives $44\times$ and Sino-Tibetan $36\times$, below Uralic, without being shallower. The chance ratio does not order by depth. The paragraph is left as it stood, with its figures corrected, because it was what three points allowed one to say; what eight allow one to say is that it does not hold.

6. Limits

The instrument does not tell apart etyma that converge in the skeleton. Four pairs among the appendix's hundred codes: **kala/*kele*, **tule/*talwe*, **neljä/*ńöle*, **kolme/*kylmä* (§2.3).

No zero reflexes. The procedure of §5.1 compares equal-length skeletons, so it cannot see any correspondence involving segment loss. In Uralic that leaves out any correspondence in which a segment is lost.

A single macroarea, hence no geographic-anomaly list.

Coverage is low: 16.8% of forms enter a code.

18% of the corpus's forms never reach the catalogue for lack of IPA segmentation (§3.1). It costs no concepts, but it is 18% we have not looked at.

We did not calibrate against expert cognacy, though it exists. Uralic has published cognate sets and is among the best-reconstructed families in the world. As in Turkic, this is a shortcoming, not an impossibility.

The emergent classes have gone two families without appearing (§5.4), and we do not know whether the problem is the phase or was the families.

7. Conclusion

The Uralic code catalogue groups 13,470 forms into 1,486 codes with a precision 78 times above chance, and its codes of widest reach are the ones any Uralicist recognises at once. It sought none of them.

What this family contributes to the series is having been able to test what Turkic could only state. The rule of computing substitutions per branch pair was tested there on one cut and here on thirty-six, and here its capacity to rule out is also visible: *b* → *p* and *z* → *s* come out below 1.0 and are marked as internal alternations without anyone needing to know Uralic. Geography, which in Turkic could not be measured for want of range, here gives a monotonic gradient across all six buckets. And the sibling-code search, which in Turkic arrived late, is here a standard phase and produces the work's best result: PU **kala* split into three codes reproducing the division between western Uralic and the Ob-Ugric–Samoyedic block, without the procedure knowing that division exists.

And two confirmations that count because they are in a second family and not the same one: automatic affix stripping again worsens the advantage over chance, and again finds hundreds of prefixes where there are none; and high enrichment still does not mean antiquity.

An instrument that claims little can be applied where there is no literature. But it is worth applying it first where the family has the structure the instrument needs — because only there does one learn whether what the instrument fails to find is its own fault or the data's.

References

- Brown, C. H., Holman, E. W., Wichmann, S. and Velupillai, V. (2008). Automated classification of the world's languages. *Sprachtypologie und Universalienforschung* 61.
- Dellert, J., Daneyko, T., Münch, A. *et al.* (2020). NorthEuraLex: a wide-coverage lexical database of Northern Eurasia. *Language Resources and Evaluation* 54.
- Dolgopolsky, A. B. (1964). Gipoteza drevnejšego rodstva jazykovyx semej Severnoj Evrazii s verojatnostnoj točki zrenija. *Voprosy Jazykoznanija* 2.
- Hammarström, H., Forkel, R., Haspelmath, M. and Bank, S. *Glottolog*. Max Planck Institute for Evolutionary Anthropology.
- Häkkinen, J. (2009). Kantauralin ajoitus ja paikannus: perustelut puntarissa. *Suomalais-Ugrilaisen Seuran Aikakauskirja* 92.
- Jäger, G. (2018). Global-scale phylogenetic linguistic inference from lexical resources. *Scientific Data* 5.
- Janhunen, J. (2009). Proto-Uralic — what, where and when? *Suomalais-Ugrilaisen Seuran Toimituksia* 258.
- Key, M. R. and Comrie, B. (eds.). *Intercontinental Dictionary Series*.
- List, J.-M. (2012). SCA: phonetic alignment based on sound classes. In *New Directions in Logic, Language and Computation*. Springer.
- List, J.-M. (2014). *Sequence Comparison in Historical Linguistics*. Düsseldorf University Press.
- List, J.-M., Cysouw, M. and Forkel, R. (2016). Concepticon: a resource for the linking of concept lists. *LREC*.
- List, J.-M., Forkel, R., Greenhill, S. J. *et al.* (2022). Lexibank, a public repository of standardized wordlists. *Scientific Data* 9.

List, J.-M., Greenhill, S. J. and Gray, R. D. (2017). The potential of automatic word comparison for historical linguistics. *PLOS ONE* 12(1).

Sammallahti, P. (1988). Historical phonology of the Uralic languages. In D. Sinor (ed.), *The Uralic Languages*. Brill.

Toledo Martínez, A. (2026). *Integrative Corpus: a multi-layer lexical database of 225 macrosystems, with row-level provenance and its verification apparatus*. — [the data] · <https://doi.org/10.13140/RG.2.2.31244.68486>

Turchin, P., Peiros, I. and Gell-Mann, M. (2010). Analyzing genetic connections between languages by matching consonant classes. *Journal of Language Relationship* 3.

Data and reproducibility

Catalogue	families/uralic/output/catalogo-uralic.tsv — 1,486 rows
Per-form detail	families/uralic/output/catalogo-uralic-lenguas.tsv — 7,886 rows
Hygiene (§3.1)	analysis/fases/fase0_higiene.py → runs/fase0-higiene.md
Data decisions (§3.1)	ci/decisiones.py
Data loading, population, warnings	ci/familia.py
Headline figures (§4.1, §5.6)	ci/metricas.py
Construction	analysis/fases/catalogo_familia.py
Grouping and nulls (§3.3, §4.1)	analysis/fases/agrupamiento_familia.py → runs/agrupamiento.md
Stripping (§5.3)	analysis/fases/fase2_despiece.py → runs/fase2-despiece.md
Aggregated substitutions (§5.4)	analysis/fases/fase6_clases_emergentes.py → runs/fase6-clases.md
Substitutions per branch pair (§5.1)	analysis/fases/fase6_por_rama.py → runs/fase6-por-rama.md
Codes and geography (§5.5)	analysis/fases/codigos_vs_parentesco.py → runs/codigos-parentesco.md
Normalisation table (§2.1)	analysis/fases/tabla_normalizacion.py → runs/normalizacion.md
Appendix A	analysis/fases/anexo_codigos.py
Cross-family contamination audit	analysis/tools/audita_familias.py

The five reports under **runs/** declare the same population (§3.4), and `verifica_paper.py` fails if they differ.

Data sources: Lexibank (List, Forkel, Greenhill et al. 2022, CC-BY-4.0), IDS (Key & Comrie, CC-BY-4.0), NorthEuraLex (Dellert et al. 2020, CC-BY-SA-4.0), Concepticon and Glottolog (CC-BY-4.0).

Appendix A · One hundred annotated codes

The catalogue has 1,486 codes and the complete data is in `families/uralic/output/`. What follows is a readable sample: the hundred widest-reaching binary and ternary codes, each with up to sixteen languages and their words.

Two things about the presentation. Languages are spread along the alphabetical list, not taken from its head; long names are truncated and forms are not altered. And the thematic grouping is a reading order declared by hand, not a semantic classification.

This appendix carries no emergent-code column. As in Turkic, the transitive closure of substitutions produces no clean classes in this family, and transplanting the Austronesian ones would be convenient and false: a class's value is relative to the family that produced it.

One warning that governs the reading, and which weighs less in Uralic than in Turkic without disappearing: 59 % of forms share a skeleton, within their own language, with a form of another concept. A code may gather the same word used for two things, or two words the instrument does not distinguish, and one must look at the forms to know which.

Reverse index: codes recurring across several concepts

In Uralic this table is less dramatic than in Turkic, and the reason is arity: 3.28 consonants per form against 2.85. With one more consonant on average there is more room to distinguish, and codes spread across fewer concepts.

The case worth looking at is `k·l`, which appears with "fish" and with "tongue (organ)". It is neither collision nor colexification but two distinct etyma converging in the skeleton: PU **kala* "fish" and PU **kele* "tongue". The instrument joins them because it sees only *k* and *l*; Finnish separates them effortlessly — *kala* against *kieli* — because it keeps the vowels. It is the cleanest illustration of the cost, and there is no need to go further.

Code	Concepts it appears with, and their enrichment
<code>k·l</code>	fish (19, 170×) · language (13, 121×) · tongue (10, 89×)
<code>p·l</code>	half (15, 153×) · side (11, 112×) · ear (10, 98×)

Numerals

Only two numerals make the cut, and one should read no more into that than is there: the appendix's cut is the widest-reaching codes, not a sample of the numeral lexicon.

What the two blocks do show is the usual thing: the vowel varies and the skeleton does not. "Four" runs through *neljä*, *njeallje*, *ñil'*, *ñel'*, *nel* without leaving `n·l`. And `n·l` also appears with "arrow" — two different words under one code, which is the cost §2.3 declares.

`n·l` | four · binary · 17 languages · enrichment 204×

Erzya Mordvin <i>nʷilʷe</i>	Estonian <i>neli</i>	Hill Mari <i>nwl</i>	Inari Sami <i>neʎi</i>
Khanty <i>nʷal</i>	Komi <i>noʎ</i>	Komi-Permyak <i>noʎ</i>	Livonian <i>ne:ʎa</i>
Mansi <i>ni:la</i>	Mari <i>'nəl</i>	Moksha <i>pilʲe</i>	North Karelian <i>nel:æ</i>
Northern Saami <i>neʌʎ:e</i>	Olonets Karelian <i>nel:i</i>	Skolt Sami <i>neʌ:e</i>	Udmurt <i>nʷilʷ</i>

`k·l·m` | three · ternary · 10 languages · enrichment 239×

Erzya Mordvin kolmo	Estonian kolm	Finnish kolme	Inari Sami kulme
Kildin Sami ko:l:m	Livonian kuolm	Moksha kolma	North Karelian kōlme
Olonets Karelian kōlme	Southern Sami kulme		

Body parts

The most solid block of the appendix, and where Uralic most resembles itself.

v·r|blood (19 languages, 343×) runs through *veri*, *vīr*, *vur*, *vər* in six branches. And here is something the block does not show and §5.2 does: a second code exists, *w·r*, with five languages — but Estonian, Finnish and Hungarian are in both, and their forms are *veri* and *weři*. It is not a split of languages: it is two sources writing the same word.

s·l·m|eye gathers *silma*, *sielime*, *ŋalme*, *ŋa:ɬ:m* — the initial sibilant variation that §2.1 merges into *s*. And **n·n|nose** has a sibling in the catalogue, *n·r*, with the Permic languages and Mari: two different words for “nose”, which the catalogue separates (§5.2).

v·r|blood · binary · 19 languages · enrichment 343×

Erzya Mordvin veri ^j	Estonian veri	Hill Mari vur	Hungarian ve:r
Inari Sami vor:ə	Ingrian veri	Komi vir	Komi-Permyak vir
Livonian ver	Lule Sami var:a	Mari vyr	North Karelian veri
Northern Saami var:a	Olonets Karelian veri	Skolt Sami vər:ə	Udmurt vir

s·l·m|eye · ternary · 14 languages · enrichment 352×

Erzya Mordvin sielime	Estonian silm	Finnish silmä
Inari Sami tŋalme	Ingrian silmä	Kildin Sami tŋa:ɬ:m
Livonian si:lma	Lule Sami tŋalm:ə	Moksha sielime
North Karelian silmä	Olonets Karelian silmy	Skolt Sami tŋelim:ɟə
Southern Sami tŋelmie	Veps silim	

s·n|VEIN · binary · 13 languages · enrichment 98×

Erzya Mordvin san	Finnish suŋni	Inari Sami suone	Kildin Sami su:n:
Komi sən	Komi-Permyak sən	Livonian su:op	Mari ŋən
Moksha san	North Karelian sōni	Olonets Karelian sōni	Skolt Sami suən:ə
Veps sön ^j			

m·k·s|liver · ternary · 12 languages · enrichment 501×

Erzya Mordvin makso	Estonian maks	Finnish maksa	Hill Mari mokʃ
Ingrian maksa	Livonian maksa:	Mari 'mokš	Moksha maksa
North Karelian maksa	Olonets Karelian maksə	Southern Sami mæksie	Veps maks

n·n|nose · binary · 12 languages · enrichment 165×

Estonian nina	Finnish nenä	Inari Sami pune	Ingrian nenä
Kildin Sami pun:	Livonian nana:	North Karelian nenä	Northern Saami punni
Olonets Karelian nenä	Skolt Sami pu:n:ɟə	Southern Sami puenie	Veps nena

s·r·v|horn · ternary · 11 languages · enrichment 709×

Estonian sarv	Hungarian sarv	Inari Sami t̪ʲuærvɪ
Ingrian sarvi	Kildin Sami t̪ʲuer:v	Lule Sami t̪ʲu̯arv:ɛ
North Karelian sarvi	Northern Saami t̪ʲɔ̯arevi	Olonets Karelian sarvi
Skolt Sami t̪ʲuarʲv:ʲə	Veps sarʲv	

p·l|ear · binary · 10 languages · enrichment 98×

Erzya Mordvin piʎe	Inari Sami peʎi	Khanty paʎ	Komi peʎ
Komi-Permyak peʎ	Mansi paʎ	Moksha pilʲe	Northern Saami peæʎ:i
Skolt Sami peʎ:ʲə	Udmurt peʎ		

k·l|tongue · binary · 10 languages · enrichment 89×

Erzya Mordvin keʎ	Estonian keel	Finnish kieli	Ingrian ke:li
Livonian kɛ:ʎ	Moksha kalʲ	North Karelian kieli	Olonets Karelian kieli
Udmurt kwʎ	Veps kɛlʲ		

Kinship and person

n·m|name, at 378×, is the highest enrichment in this section, and it deserves a look because it explains the mechanism. Twenty languages from seven branches say *nimi*, *nim*, *nam*, *jim*, *nom:v* — over two nasal consonants, which are among the most frequent in Uralic.

That enrichment does not come from *n·m* being a rare sequence but from the fact that twenty of the languages attesting "name" say it the same way: coverage is nearly total, and enrichment is coverage divided by base frequency. No rarity is needed to reach that figure; it is enough for a word to be well distributed in a well-sampled corpus.

The pronouns come out adjacent: *m·n|I* (*minä*, *mon*, *mun*, *mænʲ*) and *t·n|THOU* (*ton*, *tun*, *tənə*). Note that **t·n** has no Finnic language: Finnish says *sinä* and sits in another code. The catalogue puts both halves on the table, separated, without knowing that is what they are.

n·m|name · binary · 21 languages · enrichment 391×

Estonian nimi	Finnish nimi	Forest Enets nim	Inari Sami nom:ɛ
Khanty ne:m	Kildin Sami ne:m	Komi nim	Livonian niʲm
Lule Sami nam:a	Mansi nam	North Karelian nimi	Northern Saami namma
Olonets Karelian nimi	Skolt Sami nəm:ɛ	Southern Sami num:ɛ	Udmurt nʲim

m·n|I · binary · 18 languages · enrichment 193×

Erzya Mordvin mon	Estonian mina	Finnish minä	Inari Sami mun
Khanty mæn	Kildin Sami munn	Livonian minə:	Lule Sami mɔn
Nenets manʲ	Nganasan mənə	Northern Saami mun	Olonets Karelian minä
Selkup man	Skolt Sami mone	Southern Sami man:ɛ	Udmurt mon

t·n|THOU · binary · 10 languages · enrichment 153×

Erzya Mordvin ton	Hill Mari twɲʲ	Inari Sami tun	Kildin Sami to:n:
Moksha ton	Nganasan tənə	Northern Saami tɔn	Selkup tan
Skolt Sami tone	Udmurt ton		

Nature and sky

m·r|sea, t·l|fire, l·m|snow. Three binaries of medium-high reach.

m·r|sea deserves a look: alongside Finnish *meri* and Estonian *meri* appear *mərje*, *morje*, *morja* in Enets, Komi-Permyak, Selkup, Erzya and Moksha. They are visibly two different provenances within a single code of reach 16, and §5.5 explains why the instrument cannot separate them.

t·l also appears with "winter": *tuli* and *talvi* in Finnish, two words that only the vowel distinguishes.

m·r|sea · binary · 17 languages · enrichment 249×

Erzya Mordvin morja	Estonian meri	Finnish meäi	Forest Enets mərje
Inari Sami me:re	Kildin Sami mierr	Komi-Permyak morje	Livonian meʔr
Lule Sami mer:a	Moksha morja	North Karelian meri	Northern Saami meäřa
Olonets Karelian meri	Selkup morje	Skolt Sami mier:e	Southern Sami meærua

t·l|fire · binary · 16 languages · enrichment 164×

Erzya Mordvin tol	Estonian tuli	Finnish tuli	Hill Mari təl
Inari Sami tul:e	Ingrian tuli	Kildin Sami to:l:	Livonian tuʔʌ
Mari 'tul	Moksha tol	North Karelian tōli	Northern Saami tōl:a
Olonets Karelian tōli	Skolt Sami tol:e	Southern Sami tōl:e	Udmurt twl

p·l|HALF · binary · 15 languages · enrichment 153×

Erzya Mordvin pelʲ	Estonian pool	Finnish puoli	Hill Mari pel
Inari Sami pe:li	Kildin Sami pe:ʌ:	Livonian pu:oʌ	Mansi pal
Mari 'pel	Moksha palje	Nenets pje:ʌa	North Karelian pōoli
Northern Saami peæl:i	Skolt Sami piel:ʲə	Veps polʲ	

l·m|snow · binary · 12 languages · enrichment 175×

Estonian lumi	Finnish lumi	Hill Mari ləm	Ingrian lumi
Komi lum	Komi-Permyak lum	Livonian lum	Mari 'lum
North Karelian lömi	Olonets Karelian lömi	Udmurt lumu	Veps lömi

p·l|SIDE · binary · 11 languages · enrichment 112×

Estonian pool	Finnish puoli	Kildin Sami pe:ʌ:
Mansi pal	Moksha palje	North Karelian pōoli
Northern Saami peæl:i	Olonets Karelian pōoli	Southern Sami pielie
Udmurt pal	Veps polʲ	

t·l|WINTER · binary · 11 languages · enrichment 117×

Erzya Mordvin tʲelje	Estonian talw	Finnish talwi	Hill Mari tel
Hungarian teel	Khanty tal	Livonian tō:lə	Mansi ta:l
Mari 'tele	Moksha tʲala	Udmurt tol	

k·l·m|cold · ternary · 10 languages · enrichment 239×

Erzya Mordvin keʌme	Estonian kylm	Finnish külme	Ingrian kylmä
Kildin Sami ke:l:m	Livonian ki:lma	Moksha kelimje	North Karelian kylmä
Olonets Karelian kylmy	Southern Sami kalme		

Animals

b·t|louse, q·s|bird, q·r·t|worm and b·l·q|fish: two binaries and two ternaries.

k·l|fish is the block that teaches most about what the appendix cannot show. Here one sees 19 languages saying *kala*, *kuolli*, *kul*, *kol*. What one does not see is that the catalogue also has q·l and x·l, with three languages each — and that Khanty and Mansi are in all three, with *ku:l*, *qul* and *χul*. They are three transcriptions of the same word spread across three codes, not three groups of languages. §5.2 details it, and it is the example from which the first version of this paper drew a false conclusion.

k·l|fish · binary · 19 languages · enrichment 170×

Erzya Mordvin kol	Estonian kala	Finnish kala	Hill Mari kol
Inari Sami kyeli	Ingrian kala	Kildin Sami ku:ʎ:	Livonian kala:
Mansi ku:ʎ	Mari 'kol	Moksha kal	North Karelian kala
Northern Saami kʊʎl:i	Olonets Karelian kala	Skolt Sami kuɛl:ʎə	Southern Sami kʷɛliɛ

m·n|egg · binary · 17 languages · enrichment 182×

Estonian muna	Finnish muna	Forest Enets mɔna	Hill Mari mənə
Inari Sami mone	Ingrian muna	Khanty moŋ	Kildin Sami mann'
Lule Sami mɔn:ɛ	Mansi muŋi	Mari 'muno	Nganasan mənʉ
North Karelian muna	Northern Saami monni	Skolt Sami mɛ:n:ʎə	Southern Sami mən:iɛ

k·k|CUCKOO · binary · 10 languages · enrichment 211×

Erzya Mordvin kuko	Finnish kæki	Hill Mari kuku	Komi kək
Komi-Permyak kək	Mari kuku'	Moksha ku'ku	North Karelian kæki
Southern Sami k̥iɛk̥i	Udmurt kiki		

p·n|dog · binary · 10 languages · enrichment 134×

Erzya Mordvin pipe	Kildin Sami pe:n:ɛ	Komi pon	Komi-Permyak pon
Livonian pip	Moksha pipe	Northern Saami pɛəna	Skolt Sami pien:ɔi
Southern Sami p̥iɛp̥i	Udmurt punw		

Plants and food

s·l|salt and j·r|root. Of the second, one should say what the list shows and nothing more: its ten languages are Finnic, Mordvin and one Khanty. There is no Saami or Permic language, however much the first version of this comment said there was — a reminder that the printed list is right below and can be counted.

s·j·n|CHEW, at 823×, is the block's highest enrichment: a ternary of infrequent consonants, which is the combination §3.2 predicts.

s·j|TEA · binary · 15 languages · enrichment 188×

Erzya Mordvin t̪ɕaj	Forest Enets t̪ʃaj	Hill Mari t̪ʃæj	Khanty ʃaj
Kildin Sami t̪ʃaj:	Komi t̪ɕaj	Komi-Permyak t̪ɕaj	Mansi s̪ʲa:j
Mari t̪ʃaj	Moksha t̪ɕaj	Nenets s̪ʲaj	North Karelian t̪ʃæijy
Olonets Karelian t̪ʃʊajʊ	Selkup t̪ɕoj	Udmurt t̪ɕaj	

k·s|BRANCH · binary · 12 languages · enrichment 55×

Estonian oks	Finnish oksa	Hill Mari ukf	Inari Sami uæksi
Livonian oksa:	Lule Sami uðakse	Mari ukf	North Karelian oksa
Northern Saami ɔaksi	Olonets Karelian oksø	Southern Sami oæksie	Veps oks

k·r|SKIN (OF FRUIT) · binary · 12 languages · enrichment 122×

Estonian ko:r	Finnish kuori	Inari Sami kor:e	Khanty kar
Komi kor	Livonian ku:orʲ	North Karelian kœri	Northern Saami kar:a
Olonets Karelian kœri	Skolt Sami kær:e	Southern Sami kâr:ε	Veps korʲ

s·l|salt · binary · 12 languages · enrichment 159×

Erzya Mordvin sol	Estonian sool	Finnish suola	Ingrian so:la
Khanty sal	Kildin Sami su:ʌ:	Livonian su:oʌ	Mansi sa:ə
Moksha sal	North Karelian sœla	Olonets Karelian sœlœ	Veps sol

j·r|root · binary · 10 languages · enrichment 249×

Erzya Mordvin jur	Estonian juur	Finnish ju:ri	Ingrian ju:ri
Khanty jo:r	Livonian ju:rʲ	Moksha jur	North Karelian ju:ri
Olonets Karelian ju:ri	Veps jœrʲ		

Objects and culture

This is the section to read with §5.5 in front of you, because it contains the highest enrichments in the whole appendix and none is what it seems:

l·m·p|TORCH OR LAMP 2964× (*lampa, lamp, lamppu*) · k·n·g|BOOK 1154× (*kniga, kniga*) · s·m·k|BAG 933× (*sumka*) · r·s·t|CROSS 855× (*risti, rist*) · l·n·t|RIBBON 477× (*lenta*).

They are the five highest-enrichment codes in the appendix, and they are visibly material these languages share with Russian. The instrument cannot distinguish them from an inherited word — both are “many languages saying the same thing” — and that is what §5.5 is about. Lined up like this, they are the best available illustration that a high enrichment says nothing about a word’s age.

k·r|bark shares a skeleton with k·r|SKIN (OF FRUIT), and they are the same word used for two things.

n·l|ARROW · binary · 16 languages · enrichment 199×

Erzya Mordvin nal	Estonian nool	Finnish nuoli	Hungarian niil
Inari Sami puole	Khanty nʷul	Kildin Sami pu:l:	Livonian nu:oʌ
Mansi nʷol	Moksha nal	North Karelian nœli	Northern Saami puolla
Skolt Sami puæl:e	Southern Sami puʌle	Udmurt nʷël	Veps nœl

k·r|bark · binary · 16 languages · enrichment 157×

Erzya Mordvin kerʲ	Estonian ko:r	Finnish kuoři	Inari Sami kor:e
Ingrian ko:ri	Khanty ka:r	Kildin Sami ke:r:	Livonian ku:orʲ
Mansi kær:	Moksha ker	North Karelian kœri	Olonets Karelian kœri
Skolt Sami kær:e	Southern Sami kâr:ε	Udmurt kur	Veps korʲ

l·n·t|RIBBON · ternary · 14 languages · enrichment 477×

Erzya Mordvin l̥ienta	Estonian l̥iint	Forest Enets ʎent	Kildin Sami ʎa:ŋ:t
Komi ʎenta	Livonian li:nta	Mansi ʎe:nta	Mari ʎe·ntʲe
Moksha l̥ienta	Nenets ʎe:nta	Olonets Karelian l̥entʊ	Selkup ʎenta
Skolt Sami l̥ient:e	Udmurt l̥ienta		

s·l·m|KNOT · ternary · 13 languages · enrichment 339×

Erzya Mordvin s̥ʉulmo	Estonian s̥alm	Finnish solmu	Inari Sami t̥ʉuolme
Kildin Sami t̥ʉu:l:m	Livonian suoʎm	Lule Sami t̥ʉuʉlm:a	Moksha s̥ʉulma
North Karelian s̥olmʊ	Olonets Karelian s̥olmi	Skolt Sami t̥ʉuəlm:e	Southern Sami t̥ʉuʌlme
Veps s̥ol̥im			

k·l|LANGUAGE · binary · 13 languages · enrichment 121×

Erzya Mordvin kelʲ	Estonian keel	Finnish kieli	Inari Sami kiele
Kildin Sami ki:l:	Livonian kɛ:ʎ	Moksha kalʲ	North Karelian kieli
Northern Saami kiel:a	Olonets Karelian kieli	Southern Sami kiele	Udmurt k̥il
Veps k̥elʲ			

k·n·g|BOOK · ternary · 12 languages · enrichment 1154×

Erzya Mordvin kinʲiga	Forest Enets kniga	Kildin Sami knīga	Komi kniga
Komi-Permyak kniga	Mari knʲiga·	Moksha kniga	Nenets kniga
Nganasan knige	Olonets Karelian kni:gʊ	Selkup kniga	Udmurt knʲiga

l·m·p|TORCH OR LAMP · ternary · 12 languages · enrichment 2964×

Erzya Mordvin lampa	Estonian lamp	Finnish lamppu	Hungarian laampa
Khanty lampa	Kildin Sami lammp	Komi lampa	Mansi lampa
Mari 'lampe	Northern Saami lampa	Selkup lampa	Udmurt lampa

s·m·k|BAG · ternary · 10 languages · enrichment 933×

Hill Mari sumkə	Kildin Sami sum̥ʲ:k	Komi sumka	Komi-Permyak sumka
Mansi sumka	Mari sumka·	Olonets Karelian sʊmkʊ	Selkup sumka
Udmurt sumka	Veps sʊmk		

r·s·t|CROSS · ternary · 10 languages · enrichment 855×

Estonian rist	Finnish risti	Inari Sami riste	Kildin Sami r̥isst
Livonian rift	North Karelian risti	Northern Saami risto	Olonets Karelian ristoʊ
Skolt Sami rist:e	Veps rist		

Qualities

k·l·m|cold and **n·r|YOUNG**. The first shares a skeleton with **k·l·m|three**: two unrelated words under one code, and that makes four such pairs in the appendix — **k·l**, **t·l**, **n·l** and **k·l·m**.

The insistence is the point. In a family with 59 % collision, two or three consonants are not enough to separate two words, and the reader must carry that into every block.

n·r|YOUNG · binary · 12 languages · enrichment 189×

Estonian noor	Finnish nuori	Inari Sami nuore	Kildin Sami nu:r:
Livonian nu:orʲ	Lule Sami nuõr:a	North Karelian nuõri	Northern Saami nuõr̥a
Olonets Karelian nuõri	Skolt Sami nuær:e	Southern Sami nuʌre	Veps nɔrʲ