

The periodic table of codes

Mathematical patterns of change across the eight Indo-European systems: a quantitative-linguistics approach

Alejandro Toledo Martínez

Comparative volume of the *coderivative-codes* programme. The eight branch volumes (Indo-Iranian, Balto-Slavic, Germanic, Italic, Celtic, Hellenic, Armenian, Albanian) modelled the consonant skeleton of each word —its binary or ternary *code*— read backwards, from each language to its proto-subsystem. This is not one more branch: it compares the systems to each other in search of mathematical patterns of change. The thesis is that the classic regularities (the satem/centum split, the PIE root constraints, lexical drift) admit a *quantitative* reading —a number, a matrix, a table with gaps— and that those gaps, as in Mendeleev’s table, predict structures not yet observed.

1 What is compared, and how

For each of the eight Indo-European branches the programme produced a reproducible database with: (i) the *code* of each cognate root —the sequence of canonical consonants, binary or ternary—; (ii) the emergent proto → language correspondences; (iii) the full code inventory; (iv) a decay fit of the code over time. This volume takes those eight databases and computes five families of patterns across systems. We are not after obvious “conversion rules” between sister languages (Spanish ~ Portuguese), but *quantitative invariants* that hold across the family.

The method belongs to quantitative and computational historical linguistics —the Neogrammarian regularity of sound change (Osthoff & Brugmann, 1878), the formalisation of sound laws as regular functions (finite-state transducers) (Kaplan & Kay, 1994), the measurement of phonetic correspondences and distances (Kondrak, 2000; List, 2012), and quantitative language phylogeny (Swadesh, 1952; Gray & Atkinson, 2003; Bouckaert et al., 2012)—. The novelty is not the apparatus but the *object*: we apply these tools not to words but to the consonant codes the programme extracted.

2 Pattern 1 · The satem/centum isogloss as a number

The *satem* / *centum* division —whether the PIE palatovelars **k* **ǵ* become sibilants or stay velar (Brugmann, 1904; Meillet, 1908)— is Indo-European’s major split, and always stated qualitatively. We turn it into a continuous index. For each branch we measure, over roots whose reconstruction contains a palatovelar, the probability that the code realises a sibilant, minus that same probability over plain-velar roots:

$$\text{Satem}(\text{branch}) = P(\text{sib} \mid \text{palatovelar}) - P(\text{sib} \mid \text{plain velar}).$$

The subtraction controls for each branch’s background sibilant rate: a positive value means the *palatovelars in particular* sibilate —exactly the definition of satem.

The result reproduces the classic partition without assuming it: it emerges from counting sibilants in the codes. Two honest caveats: Albanian yields a weak index (+0.05) —consistent with its different

branch	group	satem index	P(sib palv.)	P(sib velar)
Indo-Iranio	satem	+0.171	0.43	0.26
Balto-Eslavo	satem	+0.184	0.51	0.33
Armenio	satem	+0.354	0.44	0.08
Albanés	satem	+0.050	0.42	0.37
Germánico	centum	−0.052	0.27	0.32
Itálico	centum	−0.100	0.37	0.47
Céltico	centum	+0.092	0.33	0.23
Helénico	centum	−0.150	0.38	0.53

Sattem mean = +0.190, *centum* = −0.052; *Cohen's d* = 2.1.

Table 1: Satem index by branch. The four traditional satem branches fall on the positive side; the four centum branches near zero or negative. The group means are widely separated (*Cohen's d* $\approx$ 2.1).

merger of the series and its late attestation— and Celtic a slightly spurious positive (+0.09, thin sample). Even so, the group contrast is sharp (satem mean +0.190 vs. centum −0.052). The value of the operation: a textbook qualitative isogloss becomes a magnitude that is comparable, orderable, and—in principle— projectable onto fragmentary languages.

3 Pattern 2 · The periodic table of codes

This is the heart of the book. We pool all *binary* codes of the eight branches ($N = 490$ distinct branch-codes) and build the occupancy matrix $M[c_1][c_2]$: how many branches realise each first-consonant $\times$ second-consonant combination. It is, literally, a periodic table: rows and columns are the consonant “valences” and each cell a possible element.

	r	n	s	l	t	w	d	j	k	g
p	4	2	2	6	5	3	4	6	3	2
k	2	2	5	3	2	6	1	5	∅	∅
s	1	5	∅	4	4	4	4	1	5	1
n	2	1	5	1	2	6	2	1	1	6
g	4	7	6	6	∅	1	4	4	1	∅
w	4	2	4	2	2	∅	7	2	∅	3
m	6	4	5	3	4	1	2	1	∅	3
d	4	3	2	3	1	7	∅	4	∅	3
t	4	4	∅	2	∅	∅	∅	∅	3	2
b	6	4	1	3	2	1	1	2	2	0

Table 2: Occupancy matrix of the binary codes (10 most frequent consonants in each position; number of branches realising each code). In **pink**, the gaps: cells whose expected frequency under independence is ≥ 1 yet no branch realises.

Mendeleev’s question: are the gaps noise, or does a law predict them? We define a *gap* as any cell with zero observed occupancy but expected frequency ≥ 1 under independence ($\hat{e} = N P(c_1)P(c_2)$). There are 33 such gaps; and 24 of them (73%) are predicted by a known phonotactic constraint.

The gaps are not chance: they are laws. Every *geminate* code (s·s, w·w, d·d, k·k, l·l, g·g) is empty—and so is its mirror—: this is anti-gemination, the Obligatory Contour Principle barring two identical adjacent segments (Leben, 1973; McCarthy, 1986), which as a root constraint goes back to the Indo-European root theory of Benveniste (1935) and Meillet (1937). The two-stop gaps (g·t, k·g) reproduce the PIE root *stop co-occurrence law* forbidding certain voicing combinations (Iverson & Salmons, 1992). And the sonorant+sonorant gaps (l·r, l·n) reflect the OCP extended to sonorants (Dolgopolsky, 1999). The periodic table of codes has the same virtue as the table of elements: its blanks are *structured*.

code	expected	mirror obs.	constraint predicting it
s·s	3.66	0	geminate (anti-gemination / OCP)
g·t	2.57	2	two stops (PIE co-occurrence)
w·w	2.42	0	geminate (anti-gemination / OCP)
l·r	2.40	0	two sonorants (sonorant OCP)
t·s	2.35	4	sibilant+coronal
l·n	2.30	1	two sonorants (sonorant OCP)
d·d	2.29	0	geminate (anti-gemination / OCP)
k·k	2.23	0	geminate (anti-gemination / OCP)
k·g	2.23	1	two stops (PIE co-occurrence)
l·l	2.16	0	geminate (anti-gemination / OCP)
g·g	1.94	0	geminate (anti-gemination / OCP)
w·k	1.89	6	—
m·k	1.89	2	—
t·t	1.89	0	geminate (anti-gemination / OCP)

Table 3: The highest-expected gaps, the occupancy of their *mirror* $c_2 \cdot c_1$, and the constraint that explains them.

Further, the mirror asymmetries encode concrete historical facts. The gap w·k is empty but its mirror k·w is realised by six branches —because the labiovelar *k^w is a real *unit* of PIE, and *wk is not—. The gap t·s is empty but s·t appears four times —the directionality of *s-mobile* clusters and the old “thorn clusters”—. The matrix says not only what is missing but *in which direction* the system is allowed to build.

The predictive use (Mendeleev analogy)

Just as the gap between silicon and tin anticipated germanium, a gap *not* explained by a constraint marks a cell a lost or fragmentary Indo-European language *could* have occupied. We thus distinguish two kinds of blank: law-gaps (forbidden: never to be filled) and vacancy-gaps (permitted but unattested: reconstruction candidates). Formally, for a language known only from fragments, the family’s occupancy matrix gives a *prior* over which codes to expect —a lever for reconstructing the undocumented (Anatolian, Tocharian, Messapic, Phrygian) above chance.

4 Pattern 3 · A phylogeny from the codes

If two branches share many codes, they are close. We measure the Jaccard distance between code inventories, $d(A, B) = 1 - |A \cap B| / |A \cup B|$, and take each branch’s nearest neighbour. This is the same principle —lexical distance $\rightarrow$ tree— as glottochronology (Swadesh, 1952) and Bayesian Indo-European phylogeny (Gray & Atkinson, 2003; Bouckaert et al., 2012), but computed over *consonant codes*, not word lists.

The result is striking: from the *consonant structure* of roots alone there re-emerge (i) the Indo-Iranian – Balto-Slavic pair, backbone of satem; (ii) the “north-west centum core” joining Germanic, Celtic and Italic (Ringe et al., 2002); and (iii) the Graeco-Armenian link, one of the most debated Indo-European subgroupings (Clackson, 1994). That three independent phylogenetic hypotheses emerge from the bare consonant skeleton is strong evidence that the code carries historical signal.

branch	nearest neighbour	same group?
Indo-Iranio	Balto-Eslavo	yes
Balto-Eslavo	Céltico	no
Armenio	Helénico	no
Albanés	Céltico	no
Germánico	Céltico	yes
Itálico	Céltico	yes
Céltico	Germánico	yes
Helénico	Balto-Eslavo	no

Table 4: Nearest neighbour by code inventory. It recovers three real Indo-European groupings: Indo-Iranian – Balto-Slavic (satem), the north-western centum core (Germanic – Celtic – Italic), and Graeco-Armenian.

5 Pattern 4 · The code’s decay law

Each volume fitted code conservation over time to a decay curve $d(t) = d_{\infty} (1 - e^{-\lambda t})$, with its half-life $t_{1/2} = \ln 2 / \lambda$ —carrying over to codes the exponential model glottochronology (Swadesh, 1952; Lees, 1953) took from radioactive decay—. Across the eight branches, binary codes are systematically longer-lived than ternary ones: less consonant “mass”, less erosion surface.

branch	binary $t_{1/2}$ (ka)	ternary $t_{1/2}$ (ka)
Indo-Iranio	2.67	1.82
Balto-Eslavo	1.82	0.23
Armenio	5.78	3.47
Albanés	6.93	4.95
Germánico	0.29	0.23
Itálico	0.23	0.23
Céltico	0.23	0.23
Helénico	1.02	3.47

Table 5: Half-lives by branch and arity (thousands of years). Caveat: in Armenian and Albanian the fit is not evaluable (no ③ non-cognate synonym pairs); their figures are given for reference only.

Among the evaluable branches the binary > ternary gradient is constant; and the absolute figures suggest —with caution owing to stage sampling— more conservative branches (Indo-Iranian, Balto-Slavic) versus fast-turnover ones (Italic, Celtic, Germanic). Offered as a quantified hypothesis, not a closed result.

6 Pattern 5 · Recurrent specularity

The programme recorded *specular pairs*: codes $c_1 \cdot c_2$ whose mirror $c_2 \cdot c_1$ also exists. Pooling the eight branches, some mirrors recur across almost the whole family.

The recurrence of $d \cdot w / w \cdot d$ in six branches (the numeral ‘two’, **duo-*, and its reflex) shows that specularity is not a branch accident but a pan-Indo-European feature of the code space. With the programme’s methodological caution: specularity is a fact *of form*; its semantic reading (“qualic inversion”) belongs to the hermeneutic register, not the comparative one.

par espectral / mirror pair	ramas/branches
d·w / w·d	6
g·n / n·g	5
m·s / s·m	4
n·s / s·n	4
d·g / g·d	3
l·s / s·l	3
k·s / s·k	3
g·r / r·g	3
g·j / j·g	2
k·l / l·k	2
g·l / l·g	2
j·k / k·j	2

Table 6: Spectral pairs present in two or more branches (number of branches). Topped by d·w / w·d—the code of *duo- ‘two’ and its mirror—present in six branches.

7 Synthesis and derived hypotheses

Five magnitudes of one family. (1) The satem/centum isogloss is an index with $d \approx 2.1$ separation. (2) The periodic table of codes has *structured* gaps: 73% are predicted by a PIE root constraint, and their asymmetries encode the *k^w unit and s-mobile. (3) The distance between code inventories recovers three real IE subgroupings, Graeco-Armenian included. (4) Code decay is slower the lower the arity. (5) The d·w/w·d specularity is pan-Indo-European.

From these follow testable hypotheses, in line with the brief —“generate a kind of hypothesis”, predict the lost—:

H1 (predictive). For a fragmentary IE language (Phrygian, Thracian, Messapic, Illyrian), the family’s occupancy matrix gives a *prior* over its likely codes; an attested code falling in a *law-gap* signals a misreading or a loan, not inheritance.

H2 (quantitative). The satem index should also order the debated-position languages (Albanian, and formerly Anatolian) on a continuum, rather than forcing a binary.

H3 (structural). The law-gaps (geminate, stop+stop of forbidden voicing) must stay empty when any new branch is added; filling one would falsify the universality of the root constraint within Indo-European.

H4 (evolutionary). If binary > ternary decay is real, temporally deeper languages should preserve longer codes than fast-turnover ones; testable by adding intermediate stages.

This is the comparative *first step* with what is already available. The planned refinements—a proto → proto transition matrix between branches as an algebraic operator, a ternary periodic table, and the addition of the ultra-thin branches (Anatolian, Tocharian) as a predictive test—are left for the second iteration.

Reproducibility

python3 src/cc/compare.py reads the eight coderiv.(branch).db and writes docs/paper-cc-compare/data.gen_compare.py generates this document’s tables. Satem index by palatovelar detection in the reconstruction; periodic table over pooled binary codes; Jaccard distance between inventories; decay from per-branch fits. Primary data: IE-CoR.

References

- Benveniste, É. (1935). *Origines de la formation des noms en indo-européen*. Maisonneuve.
- Bouckaert, R. et al. (2012). Mapping the origins and expansion of the Indo-European language family. *Science* 337, 957–960.
- Brugmann, K. (1904). *Kurze vergleichende Grammatik der indogermanischen Sprachen*. Trübner.
- Clackson, J. (1994). *The Linguistic Relationship between Armenian and Greek*. Blackwell.
- Dolgopolsky, A. (1999). *The Nostratic Macrofamily and Linguistic Palaeontology*. McDonald Institute.
- Gray, R.D. & Atkinson, Q.D. (2003). Language-tree divergence times support the Anatolian theory. *Nature* 426, 435–439.
- Iverson, G.K. & Salmons, J.C. (1992). The phonology of the Proto-Indo-European root structure constraints. *Lingua* 87, 293–320.
- Kaplan, R.M. & Kay, M. (1994). Regular models of phonological rule systems. *Computational Linguistics* 20(3), 331–378.
- Kondrak, G. (2000). A new algorithm for the alignment of phonetic sequences. *Proc. NAACL*, 288–295.
- Leben, W. (1973). *Suprasegmental Phonology*. PhD thesis, MIT.
- Lees, R.B. (1953). The basis of glottochronology. *Language* 29, 113–127.
- List, J.-M. (2012). LexStat: automatic detection of cognates in multilingual wordlists. *Proc. LINGVIS/UNCLH*, 117–125.
- McCarthy, J.J. (1986). OCP effects: gemination and antigemination. *Linguistic Inquiry* 17, 207–263.
- Meillet, A. (1908). *Les dialectes indo-européens*. Champion.
- Meillet, A. (1937). *Introduction à l'étude comparative des langues indo-européennes*. 8th ed. Hachette.
- Osthoff, H. & Brugmann, K. (1878). *Morphologische Untersuchungen*, vol. I. Hirzel.
- Ringe, D., Warnow, T. & Taylor, A. (2002). Indo-European and computational cladistics. *Transactions of the Philological Society* 100, 59–129.
- Swadesh, M. (1952). Lexico-statistic dating of prehistoric ethnic contacts. *Proc. American Philosophical Society* 96, 452–463.