

The Cipher of Sound Change

A cryptographic framing of reconstruction-free comparison, and a theory that unifies substitution, context, and transposition

Alejandro Toledo Martínez — Independent researcher (ORCID: 0009-0000-1277-9697)

Abstract

We argue that a pair of related languages is a pair of ciphertexts of the same messages: a cognate set is one proto-meaning-form enciphered twice, once per language, and the system of regular correspondences is the key. On this reading the comparative method is, structurally, cryptanalysis of a substitution cipher from two messages in depth — recovering the inter-ciphertext key without the plaintext. The framing is not decorative. It shows *why* reconstruction-free comparison is principled rather than merely fashionable (keys are recoverable from depth without plaintext, as in Al-Kindi, Friedman, and the decipherment of Linear B); it supplies a taxonomy that unifies the programme’s papers — monoalphabetic substitution $\leftrightarrow$ context-free (Neogrammarian) change (Papers 1–2), polyalphabetic / keystream ciphers $\leftrightarrow$ context-sensitive laws such as Verner’s (Paper 2), transposition ciphers $\leftrightarrow$ metathesis (Paper 3), non-invertible substitution $\leftrightarrow$ mergers; and it turns qualitative claims into information-theoretic quantities — unicity distance (how many cognates fix a correspondence), key entropy / minimum description length (when a conditioned rule earns its name), and equivocation (the provable residue that bounds reconstruction). Read through this lens, the programme’s empirical finding is a sharp, falsifiable statement: *the cipher of sound change is, to first order, a monoalphabetic, context-free substitution per language pair; the celebrated context-sensitive laws are the polyalphabetic exceptions — rare, weak, or, where their key (e.g. accent, for Verner’s Law) has been overwritten in the ciphertext, uncrackable in principle*. We develop the framing rigorously, connect it to the modern lineage of computational decipherment, and mark its honest limits (sound change is natural, not adversarial; the channel is lossy; contact mixes messages).

0. Thesis in one paragraph

A pair of related languages is a pair of ciphertexts of the same messages: a cognate set is one message (a proto-meaning-form) enciphered two ways, once per language. The key is the system of regular correspondences. The comparative method — recovering the regular correspondences from cognate sets, without access to the plaintext — is, structurally, cryptanalysis of a substitution cipher from two messages in depth. This is not a loose metaphor: it explains *why* reconstruction-free comparison is principled (you can recover the key without the plaintext), it supplies a taxonomy (monoalphabetic $\leftrightarrow$ context-free change, polyalphabetic $\leftrightarrow$ conditioned laws, transposition $\leftrightarrow$ metathesis, non-invertible $\leftrightarrow$ mergers), and it turns qualitative claims into information-theoretic quantities (unicity distance, key entropy / description length, equivocation). Our empirical finding then reads as a sharp, falsifiable statement: *the cipher of sound change is, to first order, a monoalphabetic (context-free) substitution per language pair; the celebrated context-sensitive laws are the polyalphabetic exceptions, and they are rare, weak, or — where their key was erased — uncrackable in principle*.

1. The setup: languages as ciphertexts, cognates as messages in depth

Fix two related languages A and B . Historical linguistics posits, for a cognate set, a proto-form π (the plaintext) and two histories h_A, h_B that turned π into the attested reflexes w_A, w_B (the ciphertexts):

$$w_A = E_{h_A}(\pi), \quad w_B = E_{h_B}(\pi).$$

Reconstruction is the attempt to invert one leg: recover π from w_A (and $w_B, w_C, \dots$). We do not attempt that. We work with the *composed* map between the two ciphertexts,

$$K_{AB} = E_{h_B} \circ E_{h_A}^{-1},$$

the inter-ciphertext key: how A 's sounds correspond to B 's. A single aligned cognate pair (w_A, w_B) is a fragment of ciphertext-in-depth; a corpus of cognate sets is many such fragments. Recovering K_{AB} from them, regularity by regularity, is exactly what the comparative method does — and it is exactly the classical cryptanalytic situation of two messages enciphered in depth (the same plaintext under two keys), where the analyst recovers the *relationship between the keys* without ever seeing the plaintext.

Three consequences follow immediately, and they are the backbone of everything below.

1. The protoform is the plaintext we deliberately do not need. K_{AB} is defined and estimable without π .
2. Alignment is depth-registration. Lining up cognates so corresponding segments share a column (our MSA step) is the cryptanalyst's act of *bringing two in-depth messages into register*.
3. Regularity is the whole game. A cipher is breakable because it is *regular* (systematic); a sound correspondence is recoverable because sound change is *regular* (the Neogrammarian principle). These are the same statement.

1.1 Three refinements that keep the setup honest (and make it stronger)

The elegant $K_{AB} = E_{h_B} \circ E_{h_A}^{-1}$ is a first cut; three precisifications make it correct.

(i) The key is a *stochastic channel*, not a composed bijection. Writing $E_{h_A}^{-1}$ presumes the inverse exists — but §3.3 says it often does not (mergers, losses, reanalysis make E_{h_A} non-invertible). The correct object is a Markov kernel / stochastic channel, a conditional distribution

$$K_{A \rightarrow B}(b | a) = \sum_z P_B(b | z) P(z | a),$$

where z is a latent correspondence class, *not* a reconstructed protoform (it carries no committed pronunciation). This single move absorbs mergers, splits, probabilistic and context-hidden correspondences, alignment uncertainty, and irreversible loss — and it still composes cleanly, $K_{A \rightarrow C}(c | a) = \sum_b K_{B \rightarrow C}(c | b) K_{A \rightarrow B}(b | a)$, the composition law of channels (the setting of *Markov categories*). We keep "key" as the intuitive name but define it throughout as a stochastic inter-language transduction, so reconstruction-free comparison never has to pretend every change is reversible.

(ii) "Depth" here means a shared *source*, not a shared *key*. In classical cryptanalysis two messages are "in depth" when enciphered under the *same key*. Our situation is the dual — *one* source π under *two different* channels, $w_A = E_A(\pi)$, $w_B = E_B(\pi)$ — i.e. parallel ciphertexts from a common latent source (equivalently: multi-view / cross-channel cryptanalysis of correlated sources). This is a more precise label than "messages in depth," and it links the programme to multi-view learning and correlated-source coding, not only to classical depth. We keep the vivid phrase where it aids intuition, but the exact claim is *paired encipherments of a shared source*.

(iii) Form and meaning are two *coupled* channels, not one plaintext. A cognate is a continuity of *form*, which does not guarantee continuity of *meaning* (homologous forms drift apart semantically; unrelated forms converge on a concept). So the single plaintext π should split into a phonological source $\pi^{\text{form}} \xrightarrow{E^{\text{phon}}} w_L$ and a semantic source $\pi^{\text{sem}} \xrightarrow{E^{\text{sem}}} m_L$, *coupled* rather than identical, $P(w_L, m_L \mid z)$. This paper’s cipher is the transformation of π^{form} ; a full endolinguistics must also model E^{sem} (metaphor, specialisation, taboo replacement, form-borrowing without concept-borrowing) — §10. “The cognate is the same message” is the first-order approximation this paper works in, stated as an approximation.

2. Why reconstruction-free is cryptographically principled (not merely an attitude)

The reconstruction-free stance is often heard as a refusal or a fashion. Cryptography shows it is a theorem-shaped method. Kerckhoffs’s and Friedman’s cryptanalysis, and the entire history of *decipherment*, recover keys and even languages without the plaintext, from structure alone:

- Frequency and distribution break a monoalphabetic substitution with no crib at all (Al-Kindi, 9th c.).
- Messages in depth — two texts under related keys — are broken *by exploiting their relation*, not by knowing what they say (the classic “in-depth” attack; Friedman’s *Index of Coincidence*).
- Structural decipherment without the language: Ventris cracked Linear B’s *script* by grid/frequency structure (building on Kober’s regularities) before anyone “knew” the plaintext meanings; the Rosetta Stone worked as a crib (known-plaintext), the exception that proves the rule.

Reconstruction-free comparison is the linguistic instance of this: recover K_{AB} from cognate sets in depth. We lean on the plaintext (π , the protoform) nowhere in the estimation; it can enter only afterwards, as an external check. This is the precise sense in which the programme is reconstruction-free — and it is a cryptographic sense, not a stylistic one.

3. A taxonomy of ciphers $\leftrightarrow$ a taxonomy of sound change

The payoff of the framing is that the classical cipher types map cleanly onto the classical types of sound change, and each mapping is quantitative and testable.

3.1 Monoalphabetic substitution $\leftrightarrow$ context-free (Neogrammarian) change

Each plaintext symbol maps to one ciphertext symbol *regardless of position*. This is a regular, unconditioned sound change: PIE $p \rightarrow$ Germanic f everywhere. *Grimm’s Law is a structured substitution alphabet — indeed a chain shift, a systematic rotation of the manner dimension of the stop system (voiceless stop $\rightarrow$ voiceless fricative $\rightarrow \dots$), the phonological analogue of a Caesar-like offset in feature space rather than a random permutation. Paper 1’s repertoire and Paper 2’s $P(b \mid \text{pair}, a)$ both measure this* layer: the near-regular, context-free substitution that dominates. Empirical result: it dominates — $P(b \mid \text{pair}, a)$ is nearly deterministic; the cipher is, to first order, monoalphabetic.*

3.2 Polyalphabetic / keystream ciphers $\leftrightarrow$ context-sensitive (conditioned) laws

A polyalphabetic cipher switches its substitution alphabet according to a key stream. A conditioned sound law does the same: the correspondence for a segment depends on its environment. Verner's Law is the paradigm case, and the parallel is exact and instructive. Verner's "exceptions to Grimm" are governed by the *position of the Proto-Indo-European accent* — a key stream. Crucially, later Germanic *moved the accent to the initial syllable*, erasing the conditioning environment from the daughter forms. So the law looks *irregular* in the reflexes alone; it becomes regular only when the hidden key (accent, still visible in Sanskrit/Greek) is restored. In cipher terms: the keystream was overwritten in the ciphertext, so the polyalphabetic structure is uncrackable from the ciphertext alone and requires an external witness to the key. This is not a colourful aside — it is *precisely* our empirical wall (§5): on a corpus with no stress, the accent keystream is absent, so Verner-class laws are information-theoretically unrecoverable, not merely hard.

3.3 Non-invertible (lossy) substitution $\leftrightarrow$ mergers, and the limits of reconstruction

Not every cipher is a bijection. Mergers are many-to-one: if B collapses o and a to a , then B 's a does not determine which proto-vowel it came from. The encipherment E_{h_B} is not invertible; the residual uncertainty about π — the equivocation $H(\pi \mid w_B)$ — is strictly positive and *cannot be driven to zero by any amount of ciphertext*. This is the crisp, information-theoretic statement of a fact historical linguists know in their bones: reconstruction is under-determined, and irreversibly so wherever information was destroyed. The programme's reconstruction-free stance is, in this light, not a limitation we accept but a *recognition of an information-theoretic bound*: we characterise the recoverable inter-ciphertext key K_{AB} rather than chase a plaintext that the channel has, in places, provably erased.

3.4 Transposition ciphers $\leftrightarrow$ metathesis (Paper 3)

A transposition cipher reorders symbols without substituting them. Metathesis is a transposition cipher on the skeleton (Spanish *murciélagos* $\leftarrow$ *murciégalo*; *cocodrilo* $\leftarrow$ *crocodilo*). Paper 3's object — the operator gaining a permutation component $\pi \in S_k$ — is exactly the transposition layer of the cipher, orthogonal to the substitution layer of Papers 1–2. The cryptographic taxonomy is what *unifies the three papers under one theory*: substitution (Papers 1–2) and transposition (Paper 3) are the two classical cipher families, and real sound change is their product.

3.5 Indels and channel noise $\leftrightarrow$ epenthesis/loss, borrowing, analogy, chance

Insertion and deletion of segments (epenthesis, loss) are null-cipher / padding operations — handled at the alignment (gap) layer, and, honestly, *not yet scored* in our dissimilarity (Paper 1 §2's declared blind spot). Borrowing, analogy, and chance resemblance are channel noise: symbols from a different message injected into the stream. LexStat's cognate detection is the noise filter that decides which fragments are genuinely in depth — the cryptanalyst's judgement of which intercepts belong to the same cipher. Its fallibility is our noise floor, and our nulls exist to measure it.

3.6 Two channels, two attacks: the consonant skeleton and the vowel frequency signature

The taxonomy above (substitution, transposition, mergers) describes *operations*. There is a second, orthogonal cut — the classical division of any cryptanalysis into its two attacks — and in language it falls cleanly along the consonant / vowel axis. The two attacks are complementary: they need different evidence, they are corrupted by different confounds, and together they are two witnesses that do not share a bias.

Attack A — key recovery from depth (the consonant channel). Everything in §§3.1–3.4 is this attack: bring cognates into register and read the correspondence, i.e. recover K_{AB} from *messages in depth*. It is carried mostly by consonants, because consonants carry lexical/root identity — the extreme is Semitic templatic morphology, where the root is a consonant skeleton and vowels are a separate, grammatical overlay. Consonants undergo both substitution (Papers 1–2) and transposition (metathesis, Paper 3); they are the message-bearing skeleton, the part of the ciphertext where the key is worth recovering. This attack *requires depth*: without aligned cognates there is no correspondence to read.

Attack B — frequency analysis (the vowel channel). The very first move of classical cryptanalysis against a monoalphabetic cipher is not registration but counting: tabulate symbol frequencies and match them to the known fingerprint of the plaintext language (Al-Kindi, §2). This attack needs no depth at all — no cognates, no alignment — only a mass of monolingual text. And it is naturally the vowel attack, because vowels are where the frequency signal lives and moves. Each language has a characteristic vowel-frequency profile (how often /a e i o u .../ and their length/quality variants are used), and — this is the point — *a shift in that profile between related languages is itself the signature of a sound change*, readable without ever aligning a single cognate. The two attacks are therefore independent instruments: Attack A needs depth and is confounded by cognate-detection noise; Attack B needs only monolingual mass and is confounded by morphology and corpus composition. Agreement between two channels that fail differently is strong evidence; that is the whole reason to add the vowel channel.

Why the split is principled, not just convenient. A famous decipherment result makes it rigorous: Sukhotin’s algorithm separates vowels from consonants in an *unknown* script using only distributional / adjacency statistics — you can discover *which symbols are vowels without knowing the language*. This is Attack B in its purest, reconstruction-free form, and it strengthens the programme’s core stance: not only can we recover the inter-ciphertext key without the plaintext (Attack A), we can even partition the alphabet into its functional classes from distribution alone (Attack B). Consonant = message; vowel = a more mutable, frequency-bearing overlay, closer to a modulation / keystream than to the message itself. This is a hypothesis to be measured, not an axiom — the Semitic templatic case is the extreme, not the universal (tone-, length-, or phonation-heavy families may distribute the load differently). Its falsifiable form compares the information each channel carries about lexical identity, $I(C; \text{lexeme} \mid V)$ versus $I(V; \text{lexeme} \mid C)$, decomposed into unique / redundant / synergistic parts (*partial information decomposition*); the typological claim is then that in *some* families $I(C; \text{lexeme} \mid V)$ systematically dominates, and this is checked per family, not assumed.

What the vowel channel could reveal (frequencies as sound-change signatures). Reading vowel-frequency profiles across related languages turns several classic phenomena into *distributional* signatures, each an independent, cognate-free check on the consonant-channel story:

- Reduction / centralisation (unstressed vowels → schwa) shows up as frequency mass concentrating on central vowels — a rise in schwa/central-vowel rate down a lineage.
- Chain shifts (e.g. the Great Vowel Shift) are, in the vowel space, a permutation of positions: they surface as a systematic re-mapping of frequency-ranked vowels between stages, the vowel analogue of a transposition on the consonant skeleton.
- Vowel harmony is a within-word co-occurrence constraint on vowels — cryptographically a *polyalphabetic constraint keyed on the word’s other vowels*, and detectable as adjacency/correlation structure in the vowel stream rather than in the marginal frequencies alone.
- Quantity / length collapse (loss of a long–short contrast) is a merger of frequency peaks, and connects directly to §3.3: a merged vowel system has positive equivocation, so the pre-merger distinction is information-theoretically unrecoverable from the daughter’s vowels alone.

A scalar summarises the channel — but the right one is the entropy rate $h(V) = \lim_n \frac{1}{n} H(V_1, \dots, V_n)$

and its conditioned versions $h(V \mid \text{stress, morphology, position})$, *not* the marginal $H(V)$: marginal entropy is not channel capacity, and vowel harmony lives in sequential dependency invisible to $H(V)$ — harmony surfaces as $I(V_t; V_{t-k} \mid \text{morphology}) > 0$. Its change over time tracks expansion, reduction, and merger of the vowel space.

Honest limits of the vowel channel. Vowel frequency is *not* natural adversarial text: it is confounded by morphology (inflectional endings inflate certain vowels), by corpus composition and lemma selection, and by word-frequency effects (Zipf). So a vowel-frequency profile is a signature, not a proof — its evidential force comes precisely from being *independent* of the consonant-correspondence channel, not from standing alone. As with every claim in the programme, a shift in the profile earns interpretation only against a null that holds the confounds fixed (matched morphology / register), exactly as the consonant channel is judged against its permutation nulls (§4).

4. The information-theoretic layer (which we already use, now named)

Shannon’s *Communication Theory of Secrecy Systems* (1949) is the bridge: cryptography and information theory are one subject, and our “how many bits does knowing X buy?” metric is already living in it. Three classical quantities become the programme’s core measurements.

- Redundancy and unicity distance — and why it is *not* a permutation test. A cipher becomes uniquely solvable once there is enough ciphertext for one key-class to concentrate the posterior — the *unicity distance*. This is strictly stronger than a permutation null. A permutation test answers “*is there more regularity than chance?*”; unicity distance answers “*is there now enough data that a single equivalence-class of keys carries almost all the posterior mass?*” Operationally it is a posterior-entropy threshold, $n_\epsilon^* = \min\{n : H(K \mid D_n) \leq \epsilon\}$ (or a MAP-concentration form $P([K_{\text{MAP}}] \mid D_n) \geq 1 - \delta$), taken over observational-equivalence classes $[K]$ — two channels that induce the same observable distribution need never be told apart. The permutation nulls (Paper 1 label, Paper 2 environment) establish only the *first, weaker* fact; the posterior-entropy curve $n \mapsto H(K \mid D_n)$, per pair / phoneme / context class, is what answers *how many cognates fix /p/ $\rightarrow$ /f/*, how many a conditioned rule needs, how much the distance grows when the accent key is missing, and which correspondences are never identifiable.
- Key entropy / description length. The correspondence system K_{AB} has a size — the number of bits to specify it. This is the MDL object of Paper 2: a context-sensitive rule *earns its name* only if adding it to the key shortens the total description (key + data-given-key). “Is this a law?” becomes “does it reduce the key’s description length?” — a crisp, non-arbitrary criterion.
- Equivocation. $H(\pi \mid w_A, w_B, \dots)$ measures how much about the plaintext (protoform) remains undecided after all ciphertext — the formal residue of §3.3’s non-invertibility. It quantifies *how much of reconstruction is in principle impossible*, and it is a first-class object of the theory rather than an embarrassment to be hidden.

Our results in these terms. Paper 2 found: the within-pair substitution is near-regular (low conditional entropy, monoalphabetic); adding environment does not shorten the description on average (the polyalphabetic layer does not pay its key-cost) though it is detectable above a shuffled-key null ($p \approx 0.002$); and the strongest apparent conditioning was morphological noise at word edges, not a phonological keystream. Stated cryptographically: *the cipher is dominated by a context-free substitution key; the polyalphabetic component is small, and part of it (accent-keyed, Verner-class) is erased from this ciphertext.*

5. Historical linguistics as support — precedent, ground truth, and hidden keys

Now, and only now, we bring historical linguistics in — as support, never as an input we depend on. It provides three things.

1. Precedent and vindication. The comparative method *is* cryptanalysis avant la lettre. The Neogrammarian *Ausnahmslosigkeit* (“sound laws admit no exceptions”) is nothing but the assumption that *the cipher is regular* — the very assumption that makes any substitution cipher breakable. Historical linguistics discovered, on 19th-century evidence, the regularity principle that cryptography would later formalise; the two fields corroborate each other.
2. A worked key-recovery (Verner). Verner (1876) is a documented cryptanalytic triumph: a residue of “irregular” correspondences was made regular by locating a hidden key (the PIE accent) that had been overwritten in the ciphertext. We use it as the canonical illustration of §3.2 and as an honest boundary marker (§3.2, §4): when the key is not in our data, we cannot and do not claim to recover it.
3. Ground truth and the automatic-decipherment lineage. The known laws (Grimm, Verner, Grassmann, rhotacism, lenition) are our post-hoc test set — a *frozen inventory* (Paper 2 guardrail G4) we contrast against *after* discovery, never supply as input. And the programme sits squarely in the modern lineage of computational decipherment — Kober and Ventris on Linear B; Shannon on secrecy systems; the EM-based cipher-breaking of Knight and colleagues (the Copiale and Zodiac ciphers); the statistical decipherment of Ugaritic (Snyder, Barzilay & Knight, 2010) and the neural decipherment of Ugaritic and Linear B (Luo, Cao & Barzilay, 2019). We inherit their toolkit (frequency structure, in-depth registration, EM/MDL, held-out evaluation) and their discipline (a decipherment is trusted only when it *generalises* and is *checked against an external witness*).

The dependence runs one way. Historical linguistics supports us with precedent, a hidden-key case study, and a ground-truth inventory; we do not depend on its reconstructions — no protoform enters the estimation of K_{AB} . This is exactly the “apoyo, no dependencia” stance the programme has held from the start.

5b. Why we recover no named sound laws — and what languages’ regularity really is

A fair challenge. Historical linguistics has a long catalogue of named laws — Grimm, Verner, Grassmann, RUKI, Winter, Osthoff, the palatalisations, rhotacism, and dozens more. Our analyses, meanwhile, keep returning nulls: Paper 2 finds no extractable conditioned rule; Paper 3 finds metathesis rare and its relaxation at chance. Are languages, then, not mathematically regular after all — only cultural convention?

The honest answer is neither defeat nor paradox. It has four parts.

1. The laws live on the branches; we look, on purpose, at the composed key between two cousins. A named law is a *directed* change from a *proto*-stage to a daughter: PIE $p \rightarrow$ Germanic f . Our *reconstruction-free* key $K_{A \rightarrow B}$ between two attested cousins is $E_{h_B} \circ E_{h_A}^{-1}$ — the composition* of two whole histories, run forwards on one leg and backwards on the other.

Grimm is a clean law on the Germanic branch. But the German $\leftrightarrow$ Hindi correspondence is Grimm composed with every other Germanic change and with the inverse of every Indo-Aryan change. That

tangle is not itself any single law. We discarded the proto and the arrow of time on principle — and with them the frame in which “Grimm” is even a statement.

2. The unconditioned laws we *do* recover — as the monoalphabetic key, so well they look like the default. Grimm, in our terms, is a regular context-free substitution: $p \rightarrow f^*$ everywhere. That is exactly the monoalphabetic key Paper 2 measured and found near-deterministic.

So we are not missing the Grimm-type laws. We found their regularity so completely that it became the *background* — the substitution cipher itself — rather than a special “rule” standing out against it. An unconditioned consonant shift is the thing our method is *best* at.

3. The laws we cannot recover are the conditioned ones — and our instrument dropped exactly their conditioning. Most named laws are *conditioned*: they fire only in a context. In our vocabulary they are polyalphabetic, keyed on something. And our instrument is blind to almost every one of those keys — by design or by corpus:

law type	keyed on	do we have it?	verdict
Grimm (chain shift)	nothing (unconditioned)	—	found — the monoalphabetic key
Verner, Dybo, Hirt, Sausure	accent (stress / pitch)	no prosody in the corpus	unrecoverable in principle
Grassmann, RUKI, palatalisations, rhotacism	segmental context, often a vowel (front V, u/i, intervocalic)	vowels dropped (consonant skeleton)	mostly blind
Osthoff, Winter, Brugmann, ablaut, umlaut, Szemerényi	vowel length / quality	vowels dropped	out of reach
metathesis	segment <i>order</i>	yes (Paper 3)	targeted — rare, local
Wackernagel, Tobler–Mussafia	clitic <i>syntax</i>	n/a	out of scope

Read the third column on its own: no prosody, no vowels, coarse context, no direction. Between them, those blindnesses remove the substrate of nearly the whole catalogue.

So “we find no laws” is, to first order, “we deleted the ingredients the laws are made of, then correctly reported that we cannot cook them.” Each deletion is principled; the honesty is in naming what it costs. (Add Paper 2’s scale limit: at $n = 25$ on core vocabulary, even a conditioned rule that *was* present would sit below the resolution.)

4. So — mathematically regular, or only convention? Both, at different levels. This is the real payoff, and the cryptographic framing states it cleanly. A sound law has two parts.

Its content — that p goes to f rather than to θ — is a frozen historical accident, a specific key. Nothing mathematical forces that* substitution; in that sense it is a cultural-historical convention.

Its form — that the change is *regular* (a function, not a lottery), *composable*, *low-entropy*, and *transitive* across the family — is the mathematics. And that form is exactly what we measure, and keep confirming: the per-pair map is near-deterministic (monoalphabetic), correspondences concentrate after a handful of cognates (unicity), and composing pairwise keys reconstructs the direct key (transitivity — Ω held in 60/60 triples).

We are therefore not in conflict with the catalogue; we work one level above it. Historical linguistics catalogues the keys — which sound went where, on which branch, under which condition — and to do that it *needs* the proto, the direction, the vowels and the prosody. We measure the mathematics of the key-space — that there *is* a regular, composable, transitive key at all — which is measurable

without any of those. The Neogrammarian slogan that “sound laws admit no exceptions” is itself the meta-law we test, and it holds.

What it would take to re-discover the named laws inductively. The path is the operator model (Paper 4): treat each law as a transformation operator and let recurrent operators emerge from data — exactly the inductive, data-first stance the challenge proposes. But to make Grimm or Verner *fall out*, we must feed the operator the substrate we currently withhold: direction (dated stages, or a proto taken as one gauge choice), vowels and prosody (the richer corpus), and finer context.

Our nulls, then, are not an absence of regularity. Read the whole exercise as an ablation study on the comparative method: each ingredient we deliberately withhold — the proto, the direction, the vowels, the prosody — is a controlled cut, and the law-classes that go dark under each cut tell us, precisely and per class, which ingredient that law *requires*.

So a null here does double duty. It rules out where a law is not — a subtractive result, cheap and certain — and in the same stroke it specifies what would have to be restored for the law to appear. Cutting prosody moves us away from the prosodic laws; dropping vowels, from the vocalic ones; refusing the proto and the branches, from everything that is defined relative to them. The map of where the laws are *not*, and of what each absent class needs, is therefore the useful result — and it says exactly what to build next.

6. What the framing buys the programme

- One theory over three papers. Substitution space and its geometry (Paper 1); monoalphabetic vs polyalphabetic substitution (Paper 2); transposition (Paper 3, metathesis). The cipher taxonomy is the through-line; the atlas is the recovered key in full.
- A precise vocabulary that replaces vague talk of “structure”: context-free vs conditioned becomes monoalphabetic vs polyalphabetic; “is it a law?” becomes “does it pay its key-cost (MDL)?”; “is the signal real?” becomes “are we past unicity distance (permutation null)?”.
- Falsifiable, quantified claims about the *nature* of sound change (e.g., “the cipher is first-order monoalphabetic”), not just classifications of languages.
- A principled reconstruction-free method (in-depth cryptanalysis) and a principled account of *why full reconstruction is bounded* (equivocation under non-invertible mergers).
- A bridge to an active computational field (automatic decipherment / neural cipher-breaking), which supplies methods, benchmarks, and reviewers who already think in these terms.
- Two independent channels of evidence (§3.6): key recovery from depth on the consonant skeleton, and frequency analysis on the vowels — attacks that need different data and fail under different confounds, so their agreement is a genuine cross-check rather than one measurement restated.

7. Honest limits and dis-analogies

The framing must not be over-sold; its boundaries are part of its honesty.

- Sound change is natural, not adversarial. No designer chose the key to resist analysis; the “cipher” is the by-product of articulation, perception, and transmission. Consequences: keys are *not* uniform (each language pair has its own), and they *drift continuously* rather than being fixed — closer to a slowly-rekeying stream than to a static cipher.

- The channel is lossy and the plaintext is itself uncertain. Mergers destroy information (§3.3); and π is a reconstruction, not a recovered original — there is no Rosetta crib for most of prehistory. We therefore characterise K_{AB} (recoverable) rather than π (often not).
- Contact mixes messages. Borrowing injects fragments from other ciphers; a "language" is a confluence, not a single clean channel (the programme's coderivative/confluence stance). The cipher is really a *mixture*.
- Morphology and prosody are edits outside the substitution alphabet. Analogy re-shapes forms non-phonologically; affixation adds material at edges (our Paper 2 word-final artefact); stress conditions changes invisibly. Some of these are noise to be filtered, some are missing keys (§5). Naming them keeps the metaphor from doing work it cannot.

8. Programme placement and next steps

This document is the theoretical framing the three empirical papers can cite for their shared vocabulary and claims. Concretely:

- Paper 2 adopts the monoalphabetic/polyalphabetic language for its verdict and cites §3.2/§4 for the Verner recoverability boundary.
- Paper 3 adopts §3.4 (metathesis = transposition) as its founding analogy.
- Paper 1 can, in revision, note that its repertoire is the *substitution alphabet* and its geometry the structure of that alphabet.
- Open theoretical tasks (future): make unicity distance operational (how many cognates to fix a correspondence, empirically); estimate key entropy / MDL of K_{AB} for real families; measure equivocation under attested mergers; and explore whether automatic-decipherment models (EM, neural) recover our correspondences *without* cognate labels — a genuine decipherment test.
- Vowel channel (Attack B, §3.6): build vowel-frequency profiles per language from the monolingual corpus (no cognates required), measure profile drift and vowel-channel entropy down attested lineages, and test whether the drift signatures (centralisation, rank re-mapping, peak merger) *agree* with the consonant-channel correspondences recovered from depth — the two-witness cross-check. Confounds (morphology, register, Zipf) fixed by matched nulls.

9. From pairwise keys to a global correspondence field

The paper so far is *local*: it recovers one channel $K_{A \rightarrow B}$ per pair. The programme's larger claim is that pairwise correspondences are local charts of one multilingual atlas, and that the interesting history lives in how those charts do — and do not — fit together. Four moves make this precise; they are the seeds of Paper 5.

9.1 Synchronisation and gauge (why reconstruction-free is a theorem, not a preference). Given languages $L_1, \dots, L_n$ and estimated relative channels K_{ij} , in a roughly-reversible region one may posit latent states g_i with $K_{ij} \approx g_j g_i^{-1}$. Recovering all g_i from noisy relative K_{ij} is a group-synchronisation problem, and it is identifiable *only up to a common right-multiplication* $g_i \mapsto g_i h$ (which leaves every $g_j g_i^{-1}$ unchanged). That invariance is the reconstruction-free stance, now formal: there is no absolute origin to recover — only the network of relative relations, defined up to a choice of coordinates. A protoform is one gauge choice, not the ontologically privileged object; reconstruction is *gauge-fixing*, and the physics is in the gauge-invariant relations.

9.2 Cycle consistency and sheaves (a map of systemic defects). For any triangle A, B, C , compare the direct channel with the composed one and define the cycle defect

$$\Omega_{ABC} = D(K_{AC}, K_{BC} \circ K_{AB}).$$

Small everywhere $\Rightarrow$ a coherent, purely-genealogical field; a large Ω_{ABC} marks where *the compositional explanation fails to close* — a candidate for contact, a hidden condition, a false correspondence, semantic divergence, misalignment, unmodelled loss, or too small an operator set. It does not by itself mean "contact"; it means the local charts are locally incompatible there. Assembling these compatibilities is exactly a cellular sheaf: languages are nodes, each phonological space a stalk, each K_{ij} a restriction map; a global section is a coherent multilingual reconstruction; the sheaf Laplacian measures disagreement and its cohomology localises the obstructions. In one line: *pairwise sound correspondences are local data; a language family is the global-section problem they generate* — and where the cocycle conditions ($K_{ii} = I$, $K_{ji} \circ K_{ij} \approx I$, $K_{ik} \approx K_{jk} \circ K_{ij}$) fail is precisely where the history is.

A first empirical probe (`transformations/global-field/cycle_defect.py`). We estimated the pairwise keys reconstruction-free from IE-CoR gold correspondences — 25 languages, a coarse consonant alphabet — and measured Ω over every triple with data.

Transitivity holds. In all 60 of 60 triples, composing through the *best real* intermediary reconstructs the direct key better than through a random one (median Ω 0.034 vs 0.061). This is the first empirical sign that the pairwise keys are local charts of *one coherent, transitive atlas*.

The ranking of defects is linguistically sensible. The tight Indo-Aryan trio (Hindi–Urdu–Nepali) sits at the floor, where composition is near-perfect ($\Omega \approx 0.02$). The highest residuals are triples whose intermediary sits *off the genealogical path* — routing two Welsh varieties through Breton, or West-Slavic Old Czech and Old Polish through South-Slavic Old Church Slavonic — exactly where a compositional account *should* fail.

Caveat. The channels are sparse at $n = 25$ and coarsened, so Ω is a ranking, not an absolute. Telling genuine contact from an off-path intermediary or from estimation noise needs more data and anchor languages — the honest next step for Paper 5.

9.3 Cognate sets are hyperobjects, not bags of pairs. A cognate set over ten languages is *one* joint observation of ten correlated reflexes, not $\binom{10}{2} = 45$ independent pairs; the pairwise reduction discards joint constraints, inflates evidence, and manufactures pseudo-replication. Represent it as a tensor $X_{\ell,c,p,f}$ (language, coderivative set, aligned position, feature) or a hypergraph (hyperedges = one multilingual position), and factor it into universal + family + areal + language-specific + noise components — answering what a distance matrix cannot: *how much of a correspondence is a general phonetic tendency, how much genealogical tradition, how much a contact area, how much idiosyncrasy?*

9.4 A directed informational order and the geometry of systems. Symmetric distance $d(A, B)$ ranks similarity but not *who preserves what*. Introduce a directed *garbling* order — $A \preceq_Z B$ when some channel G gives $E_A = G \circ E_B$ (all of A 's information about the latent Z is a degradation of B 's). Languages then form a partial order of complementary preservations (A keeps a consonant contrast B lost, B keeps a vowel contrast A lost; the two are *incomparable*), which is exactly why several languages jointly identify more than any one: $I(Z; A, B, C) \geq \max\{I(Z; A), I(Z; B), I(Z; C)\}$, with genuinely *synergistic* distinctions recoverable only in combination. And because a vowel is an acoustic *distribution* (F1, F2, duration, phonation), inventories should be compared by optimal transport — Wasserstein when they share an acoustic space, Gromov–Wasserstein (relational, no shared coordinates) when they do not — so a chain shift is a coordinated flow, a merger a mass collapse, a split a bifurcation: transport of mass, not reordering of letter frequencies.

10. Toward a mathematical endolinguistics (this paper as the first instrument)

The cipher is a *local* theory of transition between representations. The larger programme — a mathematical endolinguistics — is a theory of how each language *builds* the representation it then transforms, and how such systems transmit and coordinate across a population. Three reframings drive it: pairs $\rightarrow$ networks, keys $\rightarrow$ stochastic channels, and "same message" $\rightarrow$ coupled form/meaning channels over a shared experiential field quantised per language. A language becomes a local coordinate system over several partially-shared fields (phonological, conceptual, rhythmic, bodily, social); correspondences are the transition maps; global regularities are the invariants; mergers are rank losses; contact is mixing and cycle-defect; and *systems of thought* are different geometries and inference rules over the conceptual field — each language a lossy rate–distortion quantiser $q_L(w \mid x)$ with its own saliency $p_L(x)$ and distortion d_L (the frame already validated for colour and emotion colexification, imported here as support, not dependence).

Six mathematical pillars carry it:

1. Stochastic channels / Markov categories (the object of §1.1).
2. An algebra of transformation operators — generators (substitution, insertion, deletion, permutation, fusion, fission, copy, spread, resegmentation, mixture), non-commutative composition ($T_2 \circ T_1 \neq T_1 \circ T_2 \Rightarrow$ relative chronology without protoforms, via the commutator $[T_i, T_j] = D(T_i \circ T_j, T_j \circ T_i)$), weighted-transposition metathesis on the Cayley graph of S_k , all realised as weighted finite-state transducers and selected by MDL, $T^* = \arg \min_T [L(T) + L(D \mid T)]$.
3. Networks, gauges, and sheaves (§9): synchronisation, cycle consistency, cohomological obstruction.
4. Information and transport geometry: Wasserstein / Gromov–Wasserstein on phonological *and* semantic spaces; the geometry of keys as points on a product of simplices (Jensen–Shannon, Fisher–Rao).
5. Population and rhythmic dynamics: continuous-time Markov generators $P(t + \Delta) = P(t)e^{Q\Delta}$ and replicator–mutator flow with a non-biological fitness $F_v = -\alpha A_v - \beta C_v + \gamma S_v + \delta R_v$ (articulatory cost, confusability, social value, rhythmic fit), so a "sound law" is the *macroscopic fixation* of variants, not an operator that appeared over the whole system at once; contact is structured coupling between layers of a multiplex network, not merely noise; and endorhythm enters as coupled oscillators $\dot{\phi}_i = \omega_i + \sum_j K_{ij} \sin(\phi_j - \phi_i) + \eta_i$ plus a rhythmic-energy term that predicts *where* metathesis, syncope, and epenthesis occur.
6. Semantic quantisation: rate–distortion / information-bottleneck models of how each language partitions experience (colexification, polysemy, bipolar conceptual axes).

Its honesty backbone is a set of identifiability results — what provably *cannot* be recovered: hidden keys with identical observables; pre-merger states; global states only up to gauge; contact indistinguishable from parallel change without anchors; cognitive representations underdetermined by expression — plus a hierarchical MDL that *defines* universal / areal / genealogical / idiosyncratic by what compresses across how many families, and a standing leave-one-family-out test against genealogy-, area-, or sampling-driven artefacts (the anti-apophenia guardrails, at programme scale).

On this map, *The Cipher of Sound Change* is not the whole theory but its first instrument: the mathematical theory of the phonological transformative layer. The rest is a programme of companion papers.

- Paper 4 — A Calculus of Linguistic Transformations. Generators, composition, non-commutativity, finite-state transducers, relative chronology, MDL.

- Paper 5 — The Global Correspondence Field. Synchronisation, gauges, cycle consistency, sheaves, hypergraphs, joint multilingual correspondence inference (§9 in full).
- Paper 6 — Languages as Lossy Semantic Quantisers. Rate–distortion, information bottleneck, colexification, polysemy as perceptual compression.
- Paper 7 — Endorhythm as a Dynamical Constraint on Language Change. Coupled oscillators, multiscale prosody, metathesis / epenthesis / syncope as energy minimisation.
- Paper 8 — Coding Pressure and the Evolution of Phonological Systems. Functional load, lexical distance, perceptual robustness, mergers and systemic compensation.
- Paper 9 — Endolinguistic Codes as Latent Operators. OAS and binary/ternary codes as bipolar semantic axes (a direction v_z with contextual sign $\sigma \in \{-1, +1\}$, resolving the “opposite meanings” puzzle: purify/rot, join/separate are the two poles of one transformative field), signed-graph semantics, evaluated strictly leave-one-family-out. This is the *reserved*, *blind* contrast — OAS is not used as an input before then, consistent with the standing programme rule.

References (support — to complete)

Shannon (1949) *Communication Theory of Secrecy Systems*; Friedman, *The Index of Coincidence*; Al-Kindi (freq. analysis); Sukhotin (vowel/consonant separation from distribution); Kober & Ventris (Linear B); Verner (1876); the Neogrammarians (Osthoff & Brugmann) on exceptionlessness; Knight et al. (EM decipherment; Copiale); Snyder, Barzilay & Knight (2010, Ugaritic); Luo, Cao & Barzilay (2019, neural decipherment); plus the transformations pilot and Paper 2.

For §§1.1, 9–10 (*mathematical endolinguistics, to complete*): Fritz, Perrone & others on Markov categories / categorical probability; Singer and Bandeira et al. on group synchronization and synchronization over groups; Hansen & Ghrist on cellular sheaves and the sheaf Laplacian; Peyré & Cuturi on optimal / Gromov–Wasserstein transport; Williams & Beer, and Griffith & Koch on partial information decomposition; Zaslavsky et al. on rate–distortion / information-bottleneck models of colour naming; Jackson et al. (2019) on emotion colexification; Wedel et al. on functional load and mergers; Nowak et al. on replicator–mutator language dynamics; and the finite-state-transducer tradition in computational phonology.