

ABAD Adaptive ZD30 v5.1

Controlled Middle English Syllabic Mapping • Frozen G2-S1

SEALED OUTCOME

FAIL — COVERAGE ROSE, COHERENCE DID NOT.

The syllabic model resolved 55.23% of held-out tokens, beating the 41.07% v4.5 word-level coverage. However, the complete search performed no better than equally flexible null searches ($p = 0.756$), structural alignment was negative ($z = -1.36$), and the worst folio remained 70.18% unresolved.

Decision: The increased coverage is attributable to mapping flexibility rather than a stable Middle English language layer. No syllabic mapping is locked and no plaintext claim is permitted.

Primary comparison

Measure	v4.5 word mapping	v5.1 syllabic mapping	Required
Held-out coverage	41.07%	55.23%	$\geq 40\%$ and improve
Uncertainty	58.93%	44.77%	$\leq 60\%$
Structural z	-0.81	-1.36	> 0
Phase z	-3.33	+0.71	≥ 0
Penalized composite	-10.507	-8.819	Beat baseline + nulls
Selection-aware p	0.005	0.756	< 0.05

The apparent contradiction—higher coverage but a worse p -value—is expected when a more flexible mechanism can manufacture dictionary-supported forms in both the real and null searches. Selection-aware nulls measure whether the real structure benefits more than that flexibility alone.

What this test was for

Change the encoding unit — Test whether the v5.0 failures resulted from assuming one Voynich token behaved like one mapped word.

Preserve the contributor block — Use the exact 84-folio G2-S1 assignment created without language evidence.

Demand transport — Learn syllable assignments on four page folds and apply them unchanged to the fifth.

Charge flexibility — Repeat the complete syllabic search in every one of 200 matched null replicates.

Beat the prior mechanism — Require improvement over both v4.5 coverage and its penalized composite, while still clearing every absolute gate.

Frozen mechanism

Element	Frozen rule
Segmentation	Three-glyph frozen prefix plus greedy two-glyph residual chunks; 2–4 components per token
Syllable source	128 orthographic syllables derived from the top 50,000 ranked forms in the frozen 229,691-form CME lexicon
Mapping	One component → one syllable; injective; maximum 72 component assignments
Search	Eight deterministic restarts; top 120 training token types; no held-out optimization
Complexity	0.03 per mapped component plus existing page-null penalties
Nulls	Next-class and illustration labels permuted in training; entire mapping search rerun; held-out pages untouched

Gate results

Gate	Outcome
Selection-aware $p < .05$	FAIL
Coverage $\geq 40\%$	PASS
Structural $z > 0$	FAIL
Phase $z \geq 0$	PASS

Gate	Outcome
No unit >50%	PASS
Every page null $\leq$ 60%	FAIL
Beat best null	FAIL
Beat v4.5 score	PASS
Beat v4.5 coverage	PASS

Fold behavior

Held-out fold	Coverage	Worst page null	Mapped components
0	51.60%	66.67%	55
1	59.59%	55.56%	59
2	59.80%	55.81%	58
3	53.55%	67.80%	57
4	51.59%	70.18%	55

Coverage is repeatable across folds (51.59%–59.80%), but that repeatability is not language evidence because the same search process also achieves comparable or better scores after the relevant structural labels are destroyed.

Why the resolved words are not a reading

The generated forms change with the training fold: for example, *daiin* becomes *synne*, *name*, or *sone*, while *cthy* becomes *make*, *herte*, or *wele*. That is legitimate for testing transport, but it is not a unique manuscript-wide reading. Combined with the failed null comparison, these forms cannot function as translation anchors.

Interpretation

v5.1 answers the immediate question: changing from whole-word substitution to a controlled syllabic composition model can increase dictionary coverage substantially, but the increase does not attach to the manuscript’s held-out structural organization. The model is good at generating historical-English-looking

forms; it is not good at predicting which structural or contextual behavior those forms should have.

This closes the exact injective Middle English syllabic branch on G2-S1. It does not reject every possible syllabic system. Homophonic syllables, nomenclators, null-bearing code groups, and language-independent generative systems were not tested here.

Next defensible move

Do not increase the syllable inventory or relax injectivity; that would add the same kind of flexibility the nulls already exposed. The next distinct test should be a language-independent nomenclator/control-code model: ask whether frequent Voynich forms encode a small inventory of operations, categories, or references rather than pronounceable language. It must reproduce the established final-glyph → next-class mechanism, hand/section strength differences, and illustration modulation before any dictionary is introduced.

Bottom line: We found a stronger collision generator, not a decryption. The structural layer still carries the reliable signal; the Middle English surface layer still does not lock.

Receipt and reproducibility

Specification hash:

cac791f03d0bdade588fa7dbf1c73744ee47cc35d3d3c4959d54d43282650e07

Terminal receipt hash:

f0962c0a4552a3070e52b5c87edb4ed15039370b7b5c5562b3720007a8426f57

Parent G2-S1 atlas hash:

0371430f4dbb9f056f58a465f5ce7e1bb4299d006545f95cf9388816d8e33faf

v4.5 baseline hash:

414424abbf49f73076331d6c6915fd59b04b0ad8b6f6741f6784c1eadc88ea6e

Syllable inventory hash:

f83c407cea0b8c4a23d879b8d61146d8c915e425de0b52fa7dcadb7687aba725

Machine-readable SPEC.json and RECEIPT.json accompany this report. All 200 null composites and the five frozen held-out mappings are retained in the receipt.