

RESEARCH REPORT | 11 SEPTEMBER 2026

Local-Cycle Written-Order Grammar

Final 474-cycle ordering test against two within-cycle null controls

Juan Gabriel Molina
Independent Researcher

Framework by Christine Preedy | Independent testing, reconstruction, and report by Juan Gabriel Molina.

Abstract

This report documents the final 474-cycle written-order grammar test. The analysis starts from the supplied cycle segmentation and retains the 474 cycles marked with explicit ends. Glyph blocks are reparsed in actual written order into a limited functional-role alphabet; eight previously declared first-occurrence pair relations are scored and averaged equally. The observed aggregate score is 0.835487. Two frozen within-cycle controls are generated with seed 20260911 and 999 replicates each. A full role shuffle has mean 0.527867 (SD 0.012182); a stricter control that fixes every GATE, local-end and major-close position has mean 0.692602 (SD 0.009419). Neither ensemble contains a score at least as large as the real value, giving empirical upper-tail $p = 0.001$ under each control. The result is a final PASS against these two conditional nulls. It establishes residual ordering in the supplied explicit-end cycles; it does not independently validate the supplied segmentation or the literal meanings of the role names.

Research question and scope

Does the actual written glyph order within the supplied 474 explicit-end local cycles exhibit stronger first-occurrence role ordering than expected when role positions are randomized within those same cycles? The experiment deliberately conditions on the supplied cycle boundaries so that the nulls test order rather than re-deriving segmentation.

Keywords: Voynich Manuscript; local cycles; written-order grammar; permutation null; procedural ordering; reproducibility

1. Corpus slice and written-order reconstruction

Property	Frozen specification
Block rows	2,352
Supplied cycle profiles	1,593
Primary explicit-end cycles	474
Scored relations	8
Primary score	Unweighted mean of eight pair-concordance rates
Null replicates	999 per family
Seed	20260911

The source profile table identifies the cycle groups and marks the 474 cycles used for the primary endpoint. This release treats those boundaries and inclusion flags as inputs. It does not claim an independent segmentation algorithm.

1.1 Role extraction from written glyph order

The parser keeps the order in which relevant glyph forms occur inside each stored glyph block and then concatenates blocks in cycle order. The limited scoring map is: c = PROJECT; h = RECEIVE; o/d/n/y = GATE; f/k/p/q/t and the atomic horizon forms f+/k+/p+/t+ = PLANET; a = A_INIT; a contiguous run of i = I; s = FLUID; r = R_LOCAL_END; and m/g = MAJOR_CLOSE. Unscoored forms are ignored for this specific relation statistic.

1.2 First-occurrence convention

For each declared pair A before B, a cycle is eligible only when both roles occur. Concordance is determined by the first occurrence of A and the first occurrence of B. Multiple later repetitions do not alter that pair decision. The overall statistic is the unweighted mean of the eight pair rates, so each relation contributes equally regardless of its number of eligible cycles.

2. Pair definitions and observed counts

Relation	Eligible	Concordant	Rate
GATE before PROJECT	264	138	0.522727
PLANET before PROJECT	159	111	0.698113
PROJECT before RECEIVE	279	278	0.996416
RECEIVE before A_INIT	103	54	0.524272
A_INIT before I	102	102	1.000000
I before R_LOCAL_END	102	100	0.980392
PROJECT before R_LOCAL_END	263	258	0.980989
RECEIVE before R_LOCAL_END	263	258	0.980989

Table 1. Final written-order counts for the 474-cycle scoring set. The aggregate 0.835487 score is the arithmetic mean of the eight displayed rates, not a pooled numerator/denominator ratio.

Two relations are near chance in raw written order: GATE before PROJECT ($138/264 = 0.5227$) and RECEIVE before A_INIT ($54/103 = 0.5243$). Four endpoint-related or process-order relations exceed 0.98, and PROJECT before RECEIVE is $278/279$. The aggregate test therefore reflects a mixture of weak and strong relations; the null comparison is applied to the complete eight-relation score.

2.1 Why the raw written order is primary

The purpose of this release is to score the sequence that is actually represented by the glyph order after a fixed role parser, rather than a stored role list that may itself encode interpretive ordering decisions. This makes the test easier to audit and places the ordering claim on a directly reproducible sequence transformation.

3. Null controls and decision rule

3.1 Null A - all roles shuffled within cycle

Every parsed role position in each already-selected cycle is randomly permuted. Cycle length and role multiset are preserved exactly; cycle boundaries remain fixed. This asks how much of the eight-pair score follows merely from the roles present in each cycle rather than their observed order.

3.2 Null B - gate and endpoint positions fixed

Every GATE, R_LOCAL_END and MAJOR_CLOSE position is held fixed. All other parsed roles are permuted only among the remaining positions in the same cycle. This is a harder conditional control because it preserves the placement of the most obvious timing and terminal landmarks.

$p = (1 + \text{number of null scores} \geq \text{real score}) / (999 + 1)$. PASS requires $p \leq 0.01$ under both controls.

Neither null reselects cycles or reconstructs segmentation. The experiment therefore tests within-cycle order conditional on the supplied explicit-end cycle set. That conditioning is a limitation and also a deliberate design choice: it isolates the ordering signal from the separate segmentation question.

4. Null results

Null family	Real	Null mean	Null SD	Null max	Exceeders	p
All roles shuffled	0.835487	0.527867	0.012182	0.574220	0/999	0.001
Gate/endpoints fixed	0.835487	0.692602	0.009419	0.718786	0/999	0.001

Table 2. Both controls produce zero exceeders. The minimum resolvable upper-tail p-value under 999 simulations and the declared plus-one correction is $1/1000 = 0.001$.

The observed-minus-null-mean differences are 0.307620 for the all-role shuffle and 0.142885 for the position-fixed control. Relative to the simulated SDs, these are approximately 25.25 and 15.17 null standard deviations. Those standardized separations are descriptive; the empirical p-values remain the formal reported comparison.

4.1 Distribution ceiling

The largest stored null score is 0.574220 under the all-role shuffle and 0.718786 under the position-fixed control, both below the observed 0.835487. The result therefore does not depend on a single near-tie at the top of either simulated distribution.

5. Interpretation and limits

5.1 Supported statement

Within the supplied 474 explicit-end cycles, the actual written order of the parsed roles produces a substantially higher eight-relation score than either declared within-cycle randomization control. The test therefore supports non-random residual ordering under this role projection and cycle set.

5.2 Not a semantic proof

The labels PROJECT, RECEIVE, PLANET, GATE and related terms are role names inherited from the framework. The nulls scramble positions, not meanings. A passing ordering test cannot by itself prove that a symbol literally denotes a planet, vessel, projected substance or medical instruction.

5.3 Segmentation remains conditioned

The 474-cycle inclusion mask is supplied rather than inferred inside the null. If cycle boundaries were themselves selected using role information, part of the observed ordering could be conditioned by that selection. A future segmentation-aware null could address that broader question. This release makes the narrower conditional claim only.

5.4 Aggregate versus individual relations

The $p=.001$ values apply to the predeclared aggregate statistic. They are not individual p -values for each of the eight rows in Table 1. In particular, a near-perfect raw rate should not be described as independently significant unless its own null distribution is explicitly tested.

5.5 Relation to higher nulls

These two shuffles are not Markov-1, Markov-2 or CSSR surrogate tests. The companion structural-null reports address different sequence statistics and generator families. Passing both types of experiment may be complementary, but their numerical evidence should not be multiplied as if the tests were independent.

5.6 Reproducibility versus independent replication

Rerunning the included code and frozen 474-cycle inputs reproduces this conditional statistic and its two null ensembles. It is not a statistically independent replication, and it does not remove dependence on the supplied segmentation or role projection.

6. Reproducibility, provenance and references

The clean handoff contains the final block and profile tables, standalone written-order parser/test runner, full 1,998-row null distribution, pair-count table, result JSON, protocol, dependency file and SHA-256 manifest.

```
python -m pip install -r requirements.txt
python code/run_local_cycle_test.py
python verify_manifest.py
```

Input fingerprints:

local_cycle_blocks.csv: 37ca502eb5a80ecca8a71398254db0b6ba0a3569324c1a45e8f186b5567bdd75

local_cycle_profiles.csv: f850091bcb9e7d33a1a09625f963ba6aa08b5e86ab456d7d2cfe047535a06b72

References and source records

[1] Preedy, C. and Molina, J. G. (2026). Local-Cycle Written-Order Grammar Handoff, version 2026-09-11. README.md and PROTOCOL.md.

[2] Same package. data/local_cycle_blocks.csv; data/local_cycle_profiles.csv.

[3] Same package. code/run_local_cycle_test.py; results/result.json; results/pair_counts.csv; results/null_distributions.csv.

Contributions: Christine Preedy originated the framework and proposed interpretations. Juan Gabriel Molina contributed the testing architecture, computational analysis, reproducibility packaging and report.