

VOYNICH MANUSCRIPT | COMPUTATIONAL VALIDATION

Broad-Role Codebook Transfer

Cross-folio tournament over all 5,040 role mappings

Juan Gabriel Molina Independent Researcher

Framework by Christine Preedy | Independent testing, computational analysis, reproducibility package and report by Juan Gabriel Molina.

1 / 5,040

IDENTITY RANK

13,461

OUT-OF-FOLD POSITIONS

57.499%

IDENTITY ACCURACY

Abstract

This report documents an exhaustive cross-folio validation of the supplied seven-role mapping. Thirty-nine folios are evaluated by leave-one-folio-out prediction, producing 13,461 out-of-fold target positions. Under the declared order-2 categorical count model with $\alpha=0.2$ smoothing and backoff, every one of the 5,040 test-time role permutations is scored. The supplied identity mapping ranks first with no tied rival, with log loss 1.066470 nats per target and accuracy 57.499%. The next-best mapping has loss 1.298259. The result supports stable cross-folio identity of the supplied broad-role encoding under this predictive architecture.

Research question. Does the supplied seven-role encoding remain the best mapping when all 5,040 role permutations are evaluated on folios held out from model fitting?

PASS - IDENTITY MAPPING UNIQUE RANK 1 / 5,040

1. Frozen data and evaluation design

The executable package uses 39 source folios and evaluates 13,461 noninitial target positions out of fold. Each folio is held out once; the model is trained on the remaining 38 folios and the held-out folio is then scored under the full rival family.

PLANET, GATE, PROJECT, RECEIVE, STRING, DOSE and OTHER are tournamented; HORIZON, SOLAR and SPECIAL remain fixed.

The scoring model uses second-order categorical conditional counts with additive $\alpha=0.2$ smoothing. When a second-order context is unseen, the model backs off to first-order probabilities and then to the unigram distribution.

The declared ranking statistic is aggregated out-of-fold log loss. Lower loss is better. Ties are evaluated at tolerance $1e-12$.

2. Exhaustive rival tournament

The rival space contains exactly 5,040 candidates and is enumerated exhaustively. The supplied identity mapping is compared directly with every alternative under the same trained-fold predictions.

Rank	Identity	Mapping	Log loss	Accuracy
1	Yes	PLANET->PLANET; GATE->GATE; PROJECT->PROJECT; RECEIVE->RECEIVE; STRING->STRING; DOSE->DOSE; OTHER->OTHER	1.066470	57.499%
2	No	PLANET->DOSE; GATE->GATE; PROJECT->PROJECT; RECEIVE->RECEIVE; STRING->STRING; DOSE->PLANET; OTHER->OTHER	1.298259	52.804%
3	No	PLANET->PLANET; GATE->GATE; PROJECT->PROJECT; RECEIVE->RECEIVE; STRING->STRING; DOSE->OTHER; OTHER->DOSE	1.351422	55.575%

Identity result: **unique rank 1/5,040**; loss 1.066469853; accuracy 57.4994%. The weighted first-order baseline loss is 1.175878964; the weighted unigram baseline is 1.940218651.

3. Cross-folio diagnostics

Every prediction contributing to the headline tournament is generated for a folio that was not used to fit that fold's model. The included per-folio table preserves the number of scored positions, identity second-order loss, identity accuracy, first-order loss and unigram loss for each of the 39 folds.

4. Relabeling symmetry control

A label-equivariant count model should be invariant when the same relabeling is applied consistently to both training and testing and the model is refit. Representative consistent relabelings were checked across the folds. The maximum observed loss difference was 0.0, satisfying the declared numerical tolerance.

This symmetry control clarifies the interpretation of the tournament: the evidence concerns stability of the supplied fixed encoding against test-time alternatives. Arbitrary names applied consistently to the entire model cannot be distinguished by a label-equivariant architecture.

5. Interpretation

This is a fixed-label structural transfer test. It identifies the supplied role mapping as uniquely preferred under the declared predictive architecture; it does not independently establish that the English role names are the only possible literal semantic glosses.

The cross-folio result is therefore best read as evidence that the supplied identity is structurally stable across the tested folios under a fixed predictive representation and exhaustive rival enumeration.

6. Reproducibility package

The accompanying handoff contains the frozen source-occurrence table, executable scorer and rival generator, complete candidate scores, per-folio cross-validation diagnostics, result JSON, protocol, citation file, dependency list, verification script and SHA-256 manifest.

From the package root, run the analysis script first and then the manifest verification script. The analysis regenerates the result JSON and CSV outputs from the packaged source data.

7. Current result

Metric	Value
Design	39-fold leave-one-folio-out cross-validation
Out-of-fold positions	13,461
Rival space	5,040
Identity rank	1/5,040 - no tied rival
Identity loss	1.066469853
Identity accuracy	57.4994%
First-order baseline	1.175878964
Unigram baseline	1.940218651
Verdict	PASS

PASS - IDENTITY MAPPING UNIQUE RANK 1 / 5,040

Version 2026-09-12