

HierLex: Forensic Pattern Analysis

Manual

Section 1: Lawyer's Guide to Interpretation

Recommended Interpretation Workflow:

- Start with the **Winning_Cases_Overview** tab (specifically the `plaintiff_prevailed` column) to establish the baseline for successful strategies.
- Consult **Doctrinal_Taxonomy_Definitions** to identify the main extracted Root- and Sub-patterns.
- Focus on the **Case_Pattern_Alignment** tab: Prioritize extracted patterns with higher scores on *Alignment_Score_Final*, *DRatio_LegalBERT*, and *DRatio_Jina*. All court cases are analyzed in the order of Sub-Patterns. For example, if ten cases are analyzed with parameters $K = 2$ and $N = 3$, all ten cases are evaluated sequentially across Root-Pattern 1 / Sub-Pattern 1, and so on.

HierLex is a forensic audit instrument designed to map legal arguments against verified historical success paths. Use the table below to interpret the `HierLex_Final_Pattern_Analysis.xlsx` workbook generated by your analysis. Note that the columns listed below are found within the `Case_Pattern_Alignment` tab.

Column Name (Case_Pattern_Alignment Tab)	Meaning & Interpretation
Alignment_Score_Final	A composite score (0 to 1) representing the argument's overall relevance. <i>Rule of Thumb: 0.7+ indicates a strong tactical alignment.</i>

Column Name (Case_Pattern_Alignment Tab)	Meaning & Interpretation
S_Global	Thematic overlap. Measures alignment with the general legal doctrinal topic.
S_Diff	The "De-noised" score. Measures specificity by penalizing generic legal noise. In Version 1.5, this is technically equivalent to <i>Score_LegalBERT_Winning</i> as the noise penalty is set to 0.
Score_LegalBERT_Winning / Losing	Similarity scores against winning/losing arguments using LegalBERT.
Score_Jina_Winning / Losing	Similarity scores against winning/losing arguments using Jina v2.
DRatio_LegalBERT / Jina	Discriminant Ratios for respective models. High values indicate strong tactical differentiation from losing patterns.
Winning_Distinguishable_Argument	Extracted tactical patterns found uniquely within successful case corpora.
Losing_Distinguishable_Argument	Extracted tactical patterns found within unsuccessful case corpora. Indicates "danger zone" arguments. The <code>Losing_Distinguishable_Argument</code> is not matched against individual court cases directly. Instead, it is compared against semantically extracted argument groups derived from the losing case corpus. As a result, the losing argument associated with a winning case reflects patterns drawn from losing cases that share the same Root Pattern—not from the winning case itself. This is by design: the losing argument

Column Name (Case_Pattern_Alignment Tab)	Meaning & Interpretation
	serves as a forensic contrast anchor, not a description of the case being analyzed.

Note on Case Diagnostics: The `Progress_Report` is provided as a **separate tab** within the workbook. This tab contains the saved run log, including: PROFILE, PREVAILED status, individual WIN/LOSE mechanism lines, and comprehensive pipeline summary lines.

Technical Note 1: Understanding Similarity Scores and Empty Arguments:

Users may observe non-zero similarity and `DRatio` values in the similarity matrix even when the system identifies "empty" or missing losing arguments for a specific case. This is a deliberate design feature:

- **Hierarchical Fallback (Centroid Projection):** When no explicit losing arguments for a specific sub-pattern are identified, the system dynamically falls back to the **Root Centroid**—the geometric mean of all losing arguments for that parent legal category.
- **Interpretation:** This ensures that every case is compared against the core "Losing Profile" of its associated legal root, maintaining statistical continuity in the `DRatio` calculation even when specific sub-patterns are absent.

Outcome-Based Similarity Adjustment: Following computation of the raw cosine similarities, HierLex applies a deterministic outcome-consistency adjustment before normalization. When the analyzed case was historically successful (`plaintiff_preveled = True`), the opposing (losing) similarity score is multiplied by a fixed factor of **0.05**. Conversely, when the analyzed case was historically unsuccessful (`plaintiff_preveled = False`), the winning similarity score receives the same 0.05 penalty. Consequently, the reported similarity scores and all downstream quantities—including `DRatio_LegalBERT`, `DRatio_Jina`, and `Alignment_Score_Final`—are computed from these outcome-adjusted similarity values rather than the raw cosine similarities. This deterministic forensic weighting reinforces tactical distinctions between successful and unsuccessful historical strategies while preserving the underlying semantic relationships identified by the embedding models.

Technical Note 2: Score Normalization and Neutral Similarity:

All cosine similarity scores are computed on the conventional **[-1, +1]** scale and are subsequently normalized to the **[0, 1]** interval using the transformation $(\text{Cosine Similarity} + 1)/2$.

A normalized score of **0.50** therefore corresponds to a raw cosine similarity of approximately **0.0**, indicating **no meaningful semantic preference** rather than moderate similarity.

Important: When HierLex cannot construct a losing (or winning) semantic representation because no opposing cases are available, the corresponding raw cosine similarity is defined as zero and therefore appears as **0.50** after normalization. This value should be interpreted as **neutral (absence of semantic evidence)** rather than evidence of moderate similarity.

Accordingly, similarity scores should be interpreted **relatively** by comparing competing Root-Sub Patterns within the same analysis rather than as absolute measures across unrelated analyses.

Absent Root Centroids: In rare circumstances, HierLex may be unable to construct a Root Centroid because no opposing semantic evidence exists for the corresponding legal category. In these cases, the raw cosine similarity is explicitly defined as **0.0**. After normalization using the transformation $(\text{Cosine Similarity} + 1)/2$, the reported similarity therefore becomes approximately **0.50**. This value should likewise be interpreted as a neutral absence of semantic evidence rather than as evidence of moderate doctrinal similarity.

Important Disclaimer: HierLex is a forensic pattern analysis tool intended for research and tactical support. It does not provide legal advice or guarantees of judicial outcomes. In Version 1.5, we apply a deterministic 0.05 weighting factor to opposing strategy scores based on the LLM-identified "plaintiff_prevalled" status. This judgment over the base models ensures that winning cases prioritize winning tactics, and losing cases isolate procedural failures, providing a commonsensical forensic filter for legal professionals.

Section 2: Technical Specifications & The "Engine Room"

This section provides a formal technical overview of the scoring mechanics for forensic auditors.

Case Vector Construction (V_{case})

The vector representing each legal case is synthesized using a weighted approach to capture both semantic nuance (the classification token) and broad document content (the pooled mean).

$$V_{\text{case}} = (\alpha_1 \times V_{\text{pooled_mean}}) + ((1 - \alpha_1) \times V_{\text{cls}})$$

Similarity Centroids and Computation

Similarity scores are computed by measuring the cosine distance between the case vector (V_{case}) and the established *Distinguishable Argument* vectors.

- Sim_{global} : The engine identifies the centroids of sub-patterns belonging to the same parent *Root-pattern* to calculate thematic alignment.

The Discriminant Ratio (DRatio)

$$\text{DRatio} = \frac{\text{sim_winning}}{\max(\text{sim_losing}, 0.05) + 1\text{e-}4}$$

A high DRatio confirms the argument successfully differentiates itself from failure-prone historical patterns.

The Alignment Score Final

$$\text{Score} = (\alpha_2 \times S_{\text{diff}}) + ((1 - \alpha_2) \times S_{\text{global}})$$

Variable Definitions:

- α_1 : Context Density coefficient. Manages the weight of local legal context versus global classification.
- α_2 : Sub-pattern Validation coefficient. Balances local tactical specificity (S_{diff}) against global doctrinal theme (S_{global}).
- S_{diff} : Calculated as $(Sim_{\text{winning}} - 0.0 \times Sim_{\text{losing}})$.