

SYSTEM ARCHITECTURE DOCUMENT

Production System Design & MLOps Blueprint

Project Title	Automated End-to-End Bibliometric Extraction & Predictive Lifecycle	Document Version	2.0.0-Stable
Author/Engineer	Priyanshu Vijay (Data Engineer & ML Analyst)	Target Domain	Data Engineering & MLOps

1. System Overview & Core Philosophy

This document establishes the official technical blueprint for a production-grade machine learning system designed to ingest unstructured web layouts, execute deterministic quality cleaning, validate non-linear regression boundaries, and serve dynamic real-time predictions.

The analytical baseline focuses on the extensive publishing career of author Stephen King. The architecture transitions completely away from volatile, static data environments by managing the entire lifecycle dynamically: moving programmatically from live ingestion to feature engineering arrays, cross-validated ensemble modeling, and structural terminal interface serving.

Personal System Inspiration: This infrastructure is personally driven by a deep passion for cinema, suspense thrillers, and immersive structural storytelling. By monitoring patterns across a prolific author's career, the engine models publication characteristics over half a century of literary history.

2. Infrastructure Architecture & MLOps Topology

The infrastructure is explicitly decoupled into three independent runtime layers to guarantee system reliability, prevent cross-contamination of execution environments, and ensure modular isolation of concerns:

- **Layer 1 (Ingestion & In-Memory Preprocessing):** Connects to live web resources, strips interface noise via regular expressions, and serializes clean arrays into a persistent local cache storage.
- **Layer 2 (Automated Pipeline Training):** Normalizes cached payloads, applies dimensional matrix feature transformations, coordinates rotating cross-validation passes, and exports serialized mathematical weights.
- **Layer 3 (Dynamic Runtime Server Interface):** Rehydrates frozen production binaries to handle live, continuous terminal console parameter queries with rigorous out-of-vocabulary fallback layers.

3. Data Engineering Lifecycle & Schema Specifications

3.1 Scraping & Content Harvesting Layer (pipeline.py)

The ingestion layer connects to live network sources over standard transport protocols utilizing customized user-agents to prevent access rate clamping or firewall drops. The parser isolates targeted table elements matching the explicit wikitable class name. Crucially, the loop tracks markdown document section nodes (such as header contexts) dynamically to parse and label the format style (e.g., Novel, Collection, Non-Fiction) before processing string matrices.

3.2 Automated Data Quality (DQ) Gateways & Recovery Playbook

Unstructured community layouts contain substantial structural noise. The ingestion layer applies automated data sanitization rules to guarantee array compatibility before persistent storage:

- **Categorical Corruptor Mitigation (ISBN Catching):** Monitors the publisher text string using regular expression verification algorithms (`\d{3}-\d | \d(9)`). If an industrial retail tracking or ISBN sequence is caught within the category column, it is scrubbed and assigned to an Unknown intercept to protect feature array weights from target label pollution.
- **Table Structural Shift Recovery:** If community tables wrap elements unevenly, causing the page count variable to exactly mirror the publication calendar year column, the loop flags the index shift. If it captures historical entries points (such as *The Shining* or *Salem's Lot*), it injects accurate hardcoded indices; otherwise, the row is discarded to save training set purity.
- **Hyperlink Clutter Stripping:** Run regular expression vectors (`\\[\\d+\\]`) across book title fields to remove active bracketed anchor footnote fragments before file writing.

3.3 Data Storage Contract (retrieve_data.json)

Sanitized arrays are written directly into a local data warehouse file cache structured according to the following strict schema:

```
{
  "status": "success",
  "count": 37,
  "data": [
    {
      "Title": "The Stand",
      "Year": 1978,
      "Publisher": "Doubleday",
      "Pages": 823,
      "Book_Type": "Novel"
    }
  ]
}
```

4. Machine Learning Engineering & Feature Blueprints

4.1 Feature Dimensional Matrix Engineering (train_model.py)

Data structures are dynamically converted into structural arrays optimized for matrix math transformations:

- **Continuous Temporal Array:** The Year column is preserved as an unscaled integer vector to cleanly map production velocity over time.
- **High-Cardinality Dimensions:** String categories for Publisher and Book_Type are transformed via One-Hot Matrix Encoding (`pd.get_dummies(drop_first=True)`) into discrete multi-dimensional boolean matrices (1 or 0).
- **Operational Variance Stabilization:** The training compiler tracks frequency distributions. Categories presenting less than two historical occurrences are programmatically merged into a shared Other column vector, protecting matrix operations from sparse coefficient dispersion errors.

4.2 Machine Learning Estimator Specification

The core engine utilizes an ensemble **Random Forest Regressor Architecture** configured with 150 independent decision trees and a fixed tree depth max cap of 7. This decision tree ensemble approach natively handles high-cardinality binary one-hot arrays without requiring strict input feature scaling. It effectively captures sharp non-linear boundaries, such as rapid changes in book length or printing layout styles across distinct career eras.

4.3 Validation Protocol & Generalization Testing

To prevent traditional overfitting and verify real-world generalization performance, the training lifecycle runs a robust **5-Fold Cross-Validation Framework**. The full array matrix is split into five rotating, randomized partitions. Across five distinct iterations, each partition serves as a blind test set against an ensemble optimized on the remaining four folds, capturing real error variance:

- **Mean Absolute Error (MAE):** Discloses the average structural physical page deviation expected from a production forecast query.
- **Coefficient of Determination (R² Score):** Defines the precise percentage of target book variance successfully captured and explained by the engineered feature interactions.

5. Deployment Framework & Inference Lineage

5.1 Production Binary Preservation

Upon meeting cross-validation constraints, the model fits its parameters against the complete dataset matrix and serializes the operational states into binary repository storage files:

- **models/book_length_regressor.pkl:** Stores the frozen collections of multi-decision node tree arrays.
- **models/model_features.pkl:** Logs the exact ordered structural array of engineered hot-encoded columns required to ensure proper array dimension shape alignment during runtime inference queries

5.2 Live Runtime Inference Engine Logic (predict.py)

The terminal application rehydrates the saved binary components to spin up an interactive parameter query terminal dashboard. It transforms user terminal arguments into real-time feature arrays on the fly:

```
PS C:\Users\VP\Py\Hub\web_scraping_pipeline\Stephen's king> python predict.py
@Initializing Production Inference Engine...
@Rehydrating serialized artifacts from disk...
@Model layers rehydrated successfully. Ready for real-time inference.

=====
@STEPHEN KING BOOK-LENGTH PREDICTION DASHBOARD
=====
Options: Type your parameters below, or type 'exit' to quit.
=====
Enter a custom Publication Year (e.g., 2027): 2027

@Available Book Types:
(1) Non-Fiction
(2) Novel
=====
Select a Book Type (Enter number or text): 2

@Known Production Category Publishers:
(1) Grant
(2) Hard Case Crime
(3) Other
(4) Scribner
(5) Signet Books
(6) Viking
(7) Viking Press
(8) Other / New Publisher
=====
Select a Publisher (Enter number or text): 4

@Target constraint matched for category: 'Scribner'
@Target constraint matched for book type: 'Novel'
C:\Users\VP\Anaconda3\Lib\site-packages\sklearn\utils\validation.py:2742: UserWarning: X has feature names, but RandomForestRegressor was fitted without feature names
warnings.warn(

=====
@INFERENCE ML RESULT:
@Predicted Expected Page Count: 571.8 pages
=====
```

5.3 Out-of-Vocabulary (OOV) Error Resilience Playbook

If a user submits parameters through the terminal console that were not captured during the model training lifecycle, explicit defensive programming fallback exception routines catch the input variance:

- Unseen or newly introduced publishing entities are mapped directly to the baseline Publisher_other vector column.
- Unseen formatting configurations fall back immediately to the baseline structural novel intercept (assigned to the omitted Novel baseline reference array), shielding the model from dimensionality mismatch errors.

6. Maintenance, Monitoring & Lifecycle Management

- **Pipeline Monitoring Intercepts:** If the ingestion script logs a count value falling beneath a baseline threshold (<30 records), an upstream layout corruption is assumed, and execution terminates immediately to protect active production cache files.
- **Retraining Triggers:** Complete model re-optimization is triggered whenever a fresh literary work is launched or corporate printing distributions migrate to an alternate publishing ecosystem.
- **Decoupled Engine Upgrades:** The decoupled lifecycle architecture isolates the machine learning training runtime from the scraping runtime, allowing developers to upgrade the estimator core to high-performance gradient boosting models (e.g., XGBoost) without altering data engineering layers.

7. WorkFlow

