

Apocalypse Clock JSON Data Model

Methodology, Data Structure, Calibration Logic, and Scientific Interpretation

Prepared as a public methodology document for website download.

Document version: 1.0

Date: 9 May 2026

Primary dataset reference: Apocalypse Clock JSON parameter dataset, v1.7.1 metadata revision, with later source-refresh and score-recalibration audit workflow.

Core statement

The Apocalypse Clock JSON file is not a raw measurement database and it is not a deterministic prediction of future events. It is a normalized, source-traceable, uncertainty-aware model layer for comparing global systemic risks across a common set of dimensions.

Use note

This document is intended to help readers understand how the JSON file is structured, how values are interpreted, what the ``mu``, ``lo``, and ``hi`` fields mean, why source evidence does not automatically translate into numerical changes, and how model limitations should be communicated transparently.

Contents

1. Executive summary	3
2. What the JSON file contains	3
3. Threats and metrics	3
4. How source evidence is obtained	4
4.1 Official institutional sources	4
4.2 Peer-reviewed scientific sources	4
4.3 Proxy sources	4
4.4 Expert and anchored judgment	4
5. What `mu`, `lo`, and `hi` mean	4
6. The 0-5 ordinal metrics	5
7. The growth_rate metric	5
8. The threshold metric	6
9. Evidence strength labels	6
10. How source evidence becomes a model score	6
11. Source refresh versus score recalibration	6
12. How the dashboard uses the JSON	7
13. Versioning and audit trail	7
14. Scientific limitations	7
15. Recommended public interpretation	8
16. Summary of the model logic	8
Appendix A. Parameter evidence card template	8
Appendix B. Glossary	9
Appendix C. Public disclaimer	9

1. Executive summary

The JSON file functions as the data layer for the Apocalypse Clock / Global Systemic Risk Monitor. It organizes a set of global systemic threats into a common structure so that different risk domains can be compared and processed by the dashboard. The file is designed for traceability: each parameter contains numerical estimates, uncertainty bounds, source metadata, a qualitative evidence-strength label, and explanatory notes.

The model does not claim that its numerical scores are directly measured by the cited sources. Instead, scientific reports, institutional datasets, peer-reviewed papers, and expert judgments are translated into a normalized model format. This translation requires explicit calibration rules, careful distinction between empirical measurements and proxy indicators, and transparency about uncertainty.

The most important methodological principle is separation of functions: the JSON file is a model input layer, not a raw evidence archive. Raw evidence informs the scores; it does not automatically equal the scores.

2. What the JSON file contains

The dataset follows a flat legacy schema. Instead of nested objects such as ``climate: { scale: ... }``, each model parameter is stored as a key of the form ``<threat>.<metric>``, for example ``climate.scale``, ``ai.growth_rate``, or ``nuclear.threshold``.

The core structure is:

23 threats x 8 metrics = 184 non-meta parameters

The ``_meta`` object describes how to interpret the file. It includes the schema version, list of threats, list of metrics, dimensional definitions, growth-rate calibration classes, revision records, and methodological safeguards.

Element	Meaning
Threat	A systemic risk domain, such as climate, AI, nuclear risk, cyber risk, pandemics, biodiversity, debt, water, or geopolitical conflict.
Metric	One of eight dimensions used to characterize each threat.
Parameter	A single <code>`<threat>.<metric>`</code> entry, such as <code>`climate.acceleration`</code> .
mu	Central model estimate.
lo	Lower plausible model estimate under a conservative interpretation of the evidence.
hi	Upper plausible model estimate under a worse-case or higher-risk interpretation of the evidence.
source/url/accessed	Source traceability fields: the cited source, its URL, and the date on which it was reviewed.
strength	A qualitative label describing how directly and strongly the source supports the parameter.
note	The methodological explanation linking the source evidence to the model value.

3. Threats and metrics

Each threat is assessed using the same eight metrics. This creates comparability across otherwise heterogeneous domains. Climate, nuclear conflict, cyber risk, antimicrobial resistance, and financial instability cannot be measured in the same natural unit. The model therefore uses normalized dimensions to make them analytically comparable.

Metric	Scale	Interpretation
scale	0-5	Potential systemic magnitude of harm if the threat intensifies.
urgency	0-5	Near-term priority and degree to which action is required in the present decade.
acceleration	0-5	Observed or plausibly inferred rate of worsening in drivers, exposure, capabilities, or impacts.

Metric	Scale	Interpretation
interdependence	0-5	Degree to which the threat couples to, amplifies, or is amplified by other systemic threats.
irreversibility	0-5	Persistence of damage and difficulty of recovery after harmful states are crossed.
gov_failure	0-5	Risk that institutions, governance systems, or collective action fail to contain the threat.
growth_rate	domain-specific	Annualized effective systemic risk-growth proxy; not a 0-5 score.
threshold	0-10	Normalized model threshold for systemic destabilization; not an empirical tipping point.

4. How source evidence is obtained

The file uses several categories of evidence. The preferred sources are official institutional reports and peer-reviewed scientific work. Where direct measurement is unavailable, the model uses carefully marked proxy evidence or anchored expert judgment.

4.1 Official institutional sources

These are often the strongest sources when they directly match a parameter domain. Examples include IPCC and WMO for climate, WHO for antimicrobial resistance and health security, IEA for energy-system scenarios and critical minerals, SIPRI for military expenditure, IMF and World Bank for financial and debt risk, IPBES for biodiversity, ENISA for cyber threat landscapes, and FAO or UN water reports for food, land, water, and soil systems.

4.2 Peer-reviewed scientific sources

Peer-reviewed articles are especially important for conceptual anchors such as planetary boundaries, nuclear winter, ecosystem regime shifts, climate tipping processes, or systemic-risk frameworks. A peer-reviewed source can support a score strongly, but only if it measures or directly informs the same phenomenon as the parameter.

4.3 Proxy sources

Some model parameters cannot be measured directly. In those cases, the file uses proxy indicators: related measurements that plausibly track the underlying systemic pressure. Examples include threatened-species counts as a proxy for biodiversity pressure, cyber incident reporting as a proxy for cyber-risk growth, or satellite inventory growth as a proxy for orbital-system pressure. Proxy use is acceptable only when it is explicitly identified in the note and reflected in the evidence-strength label.

4.4 Expert and anchored judgment

For some dimensions, no single empirical variable can directly support a numerical value. This is especially common for threshold values, advanced AI irreversibility, governance failure, and interdependence. In such cases the model uses expert judgment or anchored judgment. Anchored judgment means the number is inferential, but it is tied to one or more named evidence sources.

5. What `mu`, `lo`, and `hi` mean

Every parameter normally contains three numerical fields: `lo`, `mu`, and `hi`. These are not always statistical confidence intervals. They are structured uncertainty intervals used by the model.

```
mu = central model estimate
lo = lower plausible model estimate
hi = upper plausible model estimate
```

A value such as ``climate.scale.mu = 4.7`` does not mean that a source measured climate scale as exactly 4.7. It means that the evidence has been translated into the model's normalized scale and that 4.7 is the central calibrated estimate. The interval around that central value expresses uncertainty, scenario spread, proxy limitations, and methodological caution.

The wider the gap between ``lo`` and ``hi``, the more uncertainty the parameter carries. Wide intervals are appropriate where the source is indirect, proxy-based, scenario-dependent, or expert-judgment based. Narrower intervals are more defensible when the source is direct, quantitative, methodologically stable, and closely matched to the metric.

6. The 0-5 ordinal metrics

Six metrics use a normalized 0-5 ordinal scale: ``scale``, ``urgency``, ``acceleration``, ``interdependence``, ``irreversibility``, and ``gov_failure``. These values are model scores, not natural measurements. Decimal values are used because the model is more granular than a simple five-point survey, but the underlying logic remains ordinal rather than physically measured.

Range	Interpretation
0.0-0.9	Negligible systemic relevance or mostly local impact.
1.0-1.9	Limited or localized risk with weak systemic coupling.
2.0-2.9	Relevant sectoral or regional risk with partial global relevance.
3.0-3.9	Globally relevant systemic risk, but still substantially governable.
4.0-4.6	Very high systemic risk with plausible cascading effects.
4.7-5.0	Extreme systemic risk with potential civilizational or irreversible global consequences.

The following interpretations guide the six 0-5 metrics:

- `scale` asks how large the potential harm could become if the threat intensifies.
- `urgency` asks how much action is required in the present decade.
- `acceleration` asks whether drivers, exposure, capabilities, or impacts are worsening faster over time.
- `interdependence` asks how strongly the threat couples to and amplifies other systemic risks.
- `irreversibility` asks how persistent or unrecoverable the damage would be after severe states are reached.
- `gov_failure` asks how likely existing institutions and collective-action mechanisms are to fail in preventing or containing the threat.

7. The growth_rate metric

``growth_rate`` is a special metric. It is not a 0-5 score. It is an annualized effective systemic risk-growth proxy. A value of 0.02 is read as approximately 2 percent per year, but the underlying empirical object may differ by threat domain.

For example, AI may use frontier-compute or capability-growth proxies, cyber may use incident-reporting growth, space may use orbital-inventory growth, and pandemics may use event-frequency proxies. Each proxy must be interpreted through the note and the ``growth_kind`` field.

The model's raw-mode conversion logic is:

```
effective_growth = min(log1p(raw_indicator_growth) x risk_conversion, threat_specific_cap)
```

The logarithmic transformation reduces the tendency of very high raw growth rates to dominate the model. The ``risk_conversion`` factor reflects the fact that growth in an indicator is not identical to growth in systemic risk. The ``threat_specific_cap`` prevents any single proxy from generating an implausibly high effective risk-growth value.

The dataset uses safe-calibrated effective growth mode. This means that ``growth_rate.mu``, ``growth_rate.lo``, and ``growth_rate.hi`` already store calibrated effective values. They must not be passed through the formula again. Reapplying the conversion to values that are already calibrated would create a double-conversion error.

8. The threshold metric

``threshold`` is also a special metric. It uses a 0-10 scale and represents a normalized destabilization threshold inside the model. It is not a 0-5 severity score, not a physical boundary, not a directly observed empirical tipping point, and not a deterministic prediction.

A threshold value indicates how the model positions a threat relative to systemic destabilization. It is a relative model anchor. For example, climate, nuclear conflict, and governance failure can be used as reference anchors for interpreting other threats, but the threshold values should not be read as natural constants.

Because threshold values are model anchors, revising them requires a relative-model review, not a search for a single empirical tipping point. A threshold should be changed only if the relative positioning of the threat becomes inconsistent with new evidence or with the model's internal logic.

9. Evidence strength labels

The ``strength`` field is essential because it tells readers how directly the cited evidence supports the parameter. It also prevents a false impression that all numbers have the same empirical status.

Label	Meaning
strong	Direct support from a high-quality source or closely matched quantitative evidence.
moderate	Reasonable support from credible sources, but with non-trivial uncertainty or indirect mapping.
weak	Limited, proxy, small-sample, commercial, or otherwise lower-confidence evidence.
expert_judgment	Expert judgment without a direct URL anchor for the exact parameter.
anchored_judgment	Expert or model judgment calibrated to one or more named source anchors; used when numeric mapping is inferential.

A strong source for one metric may be only moderate evidence for another. For example, an annual climate report may directly support climate acceleration, but it may only indirectly support a governance-failure score. The evidence label should therefore be assigned at the parameter level, not at the source level alone.

10. How source evidence becomes a model score

The transformation from evidence to score is not mechanical for most 0-5 metrics. It follows an evidence-to-anchor calibration procedure.

source evidence -> metric interpretation -> normalized anchor scale -> `mu/lo/hi` -> source note

For 0-5 metrics, the evidence is compared against predefined anchor meanings. For growth rates, a raw indicator is converted into an effective proxy using the growth-rate formula. For thresholds, the value is reviewed as a relative model anchor.

The correct calibration question is not: 'What number did the source state?' The correct question is: 'What does the source justify on the model's scale, given the parameter definition, source fit, metric fit, uncertainty, and cross-threat comparability?'

11. Source refresh versus score recalibration

A newer source does not automatically justify a numerical change. A source can update the evidence base without altering ``mu``, ``lo``, or ``hi``. This distinction is central to responsible model maintenance.

Operation	What changes	What does not necessarily change
Source refresh	source, url, accessed, and sometimes note	mu/lo/hi
Evidence strengthening	strength, note, and possibly interval width	central mu may remain unchanged
Score recalibration	lo/mu/hi and associated methodological notes	schema and metric meaning

A numerical change is justified only when the new source provides a directly comparable empirical indicator, a clearly stronger measurement, or a defensible change in the model anchor. Otherwise, the appropriate action is to refresh the source, update the note, adjust strength, widen or narrow the interval, or mark the evidence as monitoring-only.

12. How the dashboard uses the JSON

The dashboard reads the JSON as a structured input file. The exact dashboard formula belongs to the application layer, but the data layer supplies the variables needed for model computation and explanation.

The JSON has two functions:

- Computational function: ``mu``, ``lo``, ``hi``, ``growth_rate``, ``threshold``, and related calibration fields are used by the model logic.
- Explanatory function: ``source``, ``url``, ``accessed``, ``strength``, ``note``, ``calibration_note``, and ``model_anchor_note`` provide transparency and auditability.

In conceptual terms, the model can distinguish current systemic stress from projected horizon. Current stress reflects present systemic load. Projected horizon reflects how quickly the system may approach modeled destabilization thresholds if current dynamics persist. These are distinct outputs and should not be conflated.

13. Versioning and audit trail

The file records its own revision history. This matters because a scientific model should preserve not only its current values, but also why values, notes, source labels, or calibration safeguards changed over time.

Version	Methodological role
v1.4	Raw growth-rate values were replaced by calibrated effective systemic risk-growth proxies.
v1.5	Growth-rate metadata fields were added, including <code>growth_kind</code> , <code>risk_conversion</code> , <code>raw_indicator_growth</code> , and <code>calibration_note</code> .
v1.6	<code>Growth_kind</code> values were harmonized and threat-specific caps plus do-not-reconvert safeguards were added.
v1.7	Score-transparency revision added model-anchor threshold caveats and selected conservative interval changes.
v1.7.1	Metadata cleanup clarified revision notes and preserved model scores.
v1.7.2	Source-refresh layer: updates sources without changing scores.
v1.8.0	Score-recalibration audit layer: checks whether new evidence justifies numerical changes.

14. Scientific limitations

The model is scientifically disciplined, but it is not free of limitations. The most important limitations should be stated openly.

- Normalization loses some information from the original measurement units.
- Some parameters rely on proxy indicators rather than direct measurements.
- Some values are model anchors or anchored judgments rather than empirical constants.
- Decimal values can create a false impression of precision if shown without explanatory context.
- A high-quality source does not automatically validate the exact numeric score; the mapping from source to score must be documented.
- Threshold values should not be interpreted as physical tipping points or deterministic forecasts.

These limitations do not invalidate the model. They define how the model should be interpreted: as a structured, transparent comparative risk framework rather than as a direct observational instrument.

15. Recommended public interpretation

For public communication, it is important to avoid both exaggeration and false precision. The following language is recommended:

This model compares global systemic threats using a normalized evidence-based framework. The scores are not direct measurements and should not be read as deterministic forecasts. Each value is supported by sources, methodological notes, and uncertainty intervals. The goal is transparent comparative risk assessment, not prophecy.

In user-facing tooltips, each metric should be explained in plain language. For example: `Scale measures the potential systemic magnitude of harm if this threat intensifies. It is a normalized model score from 0 to 5, not a direct measurement from a single source.`

16. Summary of the model logic

The entire JSON system can be summarized as follows:

- It collects evidence for 23 global systemic threats.
- Each threat is described with the same eight metrics.
- Six metrics are normalized 0-5 ordinal scores.
- Growth_rate is an annualized effective systemic risk-growth proxy, not a 0-5 score.
- Threshold is a 0-10 normalized model anchor, not an empirical tipping point.
- Each parameter contains a central estimate and uncertainty interval: `lo`, `mu`, and `hi`.
- Each parameter includes source metadata and a methodological note.
- The dashboard uses the file to compute stress and horizon outputs while preserving source traceability.

A concise definition is: the Apocalypse Clock JSON file is a scientifically anchored, normalized, source-traceable model layer for comparative assessment of global systemic risks.

Appendix A. Parameter evidence card template

A full scientific audit should review each of the 184 parameters using a consistent evidence card. This prevents arbitrary changes and improves reproducibility.

Field	Purpose
key	Parameter being reviewed, such as `climate.scale`.
current lo/mu/hi	Current model estimate and uncertainty interval.
current source	Existing source and URL.
source type	Primary empirical, official institutional, peer-reviewed, proxy, dashboard, commercial, expert judgment, or anchored judgment.
source fit	Whether the source directly supports the parameter, partially supports it, or only works as a proxy.
metric fit	Whether the source measures the same phenomenon as the metric definition.

Field	Purpose
numeric fit	Whether the source permits direct numerical recalibration.
calibration method	0-5 evidence-to-anchor, growth-rate raw-to-effective conversion, or threshold model-anchor review.
recommended action	Keep, source update, note update, strength change, interval change, numerical recalibration, or monitoring-only.
confidence	High, medium, or low.

Appendix B. Glossary

Term	Meaning
Normalized score	A value converted to a common model scale for comparison across domains.
Ordinal score	A score that indicates ordered severity but is not a natural physical unit.
Proxy indicator	A measurable variable used as an indirect indicator of a harder-to-measure risk process.
Model anchor	A reference value used inside the model for relative calibration rather than direct empirical measurement.
Safe-calibrated growth	A mode in which growth_rate values have already been converted into effective model values and must not be reconverted.
Structured uncertainty interval	The `lo` to `hi` range around a central `mu`, representing plausible uncertainty rather than a strict statistical confidence interval.
Source refresh	Updating source metadata without necessarily changing numerical values.
Score recalibration	Changing numerical values after a defensible change in evidence or calibration.

Appendix C. Public disclaimer

This model is designed to support transparent reasoning about global systemic risks. It is not a prediction engine, not a prophecy, and not an official scientific consensus statement. Its values are structured model inputs derived from cited evidence, calibrated assumptions, and documented uncertainty. The correct use of the model is comparative, interpretive, and methodological: it helps organize evidence and reveal systemic relationships, but it does not eliminate uncertainty.