

The Purchase Mechanics™ Standard for Consumer Market Datasets

Publication Draft V0.1

Version: 0.1

Date: August 18, 2026

Status: Internal publication candidate — technical method draft

Maintainer: Measured Choice / Purchase Mechanics

ABSTRACT

The Purchase Mechanics Standard defines how Measured Choice turns changing consumer-market information into versioned, auditable decision data. The standard separates product identity from retailer offers, preserves missing and conflicting information instead of silently filling gaps, distinguishes manufacturer claims from independent evidence and owner signals, and applies decision methods only when the available evidence supports them.

The standard is designed for two audiences at once. The consumer-facing layer should remain readable and direct. The technical layer should make each important claim traceable to a source class, a field definition, a dated dataset release, and a stated method. The objective is not to create the appearance of mathematical precision. The objective is to make it difficult for unsupported precision to enter the system.

This publication draft intentionally differs from an earlier research draft titled “Purchase Mechanics: Data Methodology & Reproducible Research Framework.” That research draft contains useful design ideas, but it also describes future-state capabilities as if they were already implemented, including continuous scraping, modal 14-day pricing, automated imputation, bitemporal storage, fixed error thresholds, and a deterministic Choice formula. Those items are not part of the current production standard unless and until they are separately implemented, tested, versioned, and added to this document.

1. PURPOSE AND SCOPE

Purchase Mechanics is a consumer-decision method. The Measured Choice Decision Dataset, or MCDD, is the data structure that supports that method.

The standard covers five practical questions:

1. What exactly is the product or configuration being compared?
2. What current offer or price observation is being used, and when was it observed?
3. What evidence supports each product fact or ownership claim?
4. What calculations or decision rules were applied to those facts?
5. What would cause a published conclusion to be reopened or changed?

The standard applies to observational consumer-market datasets built from public product pages, manufacturer documentation, retailer offers, independent testing, owner evidence, public safety information, and transparent derived calculations. It does not claim that an observational dataset is equivalent to a randomized experiment, a representative consumer panel, or a complete census of every unit sold.

1.1 Operating principles

The current standard follows these rules:

- Separate identity from commerce. A product configuration is not the same thing as a retailer offer.
- Preserve provenance. Important facts should retain source class, source location, observation date, and any material qualification.
- Preserve missingness. If a current price, lifespan, failure rate, or repair cost is not known, the default action is to leave it unknown rather than impute a decision-grade value.
- Preserve conflicts. When two credible sources disagree, the conflict is stored and carried into the decision instead of being silently averaged away.
- Separate evidence types. Manufacturer claims, independent tests, owner reports, direct observations, and modeled assumptions are different classes of evidence.
- Use hard gates before preference trade-offs. A product should not be promoted by a favorable score if it fails a basic evidence, safety, identity, or ownership-quality requirement.
- Avoid arbitrary weighted scores when the weights cannot be defended.
- Allow zero winners. A category release may conclude that no product currently earns a Choice designation.
- Stop downstream modeling when upstream evidence is not ready. More advanced mathematics does not repair weak input evidence.
- Version the conclusion. Every public recommendation should be tied to a dated release and an explicit reopen condition.

2. CORE DATA MODEL

MCDD currently separates the research system into several related layers.

2.1 Configuration Master

The Configuration Master represents the canonical product configuration being evaluated. A configuration may include brand, canonical product name, model identifiers, form factor, known technical specifications, and other relatively stable attributes.

The configuration identifier is an internal Measured Choice ID. GTIN, UPC, MPN, retailer model aliases, and manufacturer model codes are treated as evidence used to resolve identity when available; they are not assumed to exist for every product.

A configuration is not automatically treated as a unique physical platform merely because a retailer gives it a unique listing name.

2.2 Retailer Offers

Offers are separate observations linked to a configuration. An offer may include:

- retailer and seller,
- observed price,
- regular price when displayed,
- promotion or code,
- shipping terms,
- availability,
- condition,
- observation timestamp,
- source class,
- offer qualification status,
- notes about material limitations.

This separation allows the same physical configuration to have multiple prices at different retailers or at different times without duplicating the product master record.

2.3 Raw Facts

Raw Facts store evidence observations before they are compressed into decision fields. A raw observation may record a specification, warranty term, cleaning instruction, owner-duration signal, independent test result, repair-path finding, or other relevant fact.

Raw Facts are intended to make it possible to trace a decision-layer statement back to the evidence that produced it.

2.4 Decision Dataset

The Decision Dataset is a smaller, decision-focused table produced after screening and targeted enrichment. It should contain only fields that materially affect a purchase decision or the confidence of that decision.

For the countertop nugget-ice release, the decision fields include qualified price, claimed output, first-ice time, storage, water capacity, dimensions where resolved, warranty, cleaning profile, parts support, smart controls, independent-test evidence, reliability evidence, evidence maturity, safety/recall checks, market-frontier status, hard-gate read, cost-to-own status, Choice eligibility, primary unknowns, and flip conditions.

2.5 Release Manifest

The Release Manifest records the dataset version, method version, frozen population, observation counts, collection dates, known limitations, QA status, and the release that was superseded.

The release manifest is the present mechanism for reproducibility. The current system does not yet claim a fully event-sourced or bitemporal database. That architecture may be valuable later, but publication language must describe the system that actually exists.

3. MARKET POPULATION AND FREEZE RULES

3.1 Define the category before collecting products

A category dataset begins with a written inclusion and exclusion rule. The rule should specify geography, product function, form factor, configuration treatment, condition, and the time window of the release.

The purpose of the rule is to prevent the category definition from changing after attractive or unattractive products are discovered.

3.2 Frozen population

A release may freeze a set of configurations for analysis. A frozen set is a versioned research frame, not a claim that no other product exists.

For MCDD-ICE-001, the frozen U.S. countertop nugget/pebble/chewable configuration frame contains 84 configurations. Post-freeze discoveries are recorded separately and may be considered in a later release rather than silently changing the denominator mid-analysis.

3.3 Coverage language

The terms “census,” “complete market,” “total addressable market,” and “saturation” should not be used unless the evidence supports them.

The current standard prefers language such as:

“We mapped 84 configurations in the frozen U.S. market frame used for this release.”

It does not authorize:

“We identified every nugget-ice maker sold in the United States.”

A future dataset may support stronger coverage claims if manufacturer registries, retailer strata, GTIN coverage, shipment data, or other external sources make the denominator defensible.

4. ENTITY RESOLUTION AND SHARED-PLATFORM LOGIC

Entity resolution is the process of deciding whether multiple listings represent the same sellable configuration, closely related variants, or genuinely different products.

4.1 Strong identity evidence

Strong evidence can include:

- exact manufacturer model number,
- exact GTIN or UPC,
- exact MPN,
- manufacturer documentation,
- exact retailer/manufacturer crosswalk,
- matching certification or regulatory identifiers where relevant.

4.2 Supporting identity evidence

Supporting evidence can include:

- identical dimensions,
- identical wattage or electrical specifications,
- matching tank and basket capacities,
- matching control layout,
- matching manual diagrams,
- highly specific shared component or chassis details.

Supporting evidence alone should not automatically convert two products into a confirmed shared-OEM node.

4.3 Confidence labels

The public system should distinguish at least:

CONFIRMED SHARED PLATFORM — strong identifiers or direct manufacturer/retailer crosswalk.

LIKELY SHARED PLATFORM — several specific technical matches but no decisive identifier.

POSSIBLE RELATIONSHIP — visual or partial technical similarity requiring more evidence.

NO CROSSWALK ESTABLISHED — insufficient evidence.

The proposed TRUE_OEM_NODE field remains partially implemented. Until its rules are fully formalized, public wording must state the confidence level and the evidence used.

5. OFFER QUALIFICATION AND PRICE STATISTICS

5.1 Current price means observed offer, not transaction price

A qualified current price is an observed, decision-usable offer. It is not necessarily the price a consumer ultimately paid.

A qualified offer should normally have:

- a live product page,
- an identifiable retailer and seller,
- a new-condition product unless the category says otherwise,
- a price that applies to the relevant configuration,
- sufficient availability information to treat the offer as actionable,
- no obvious bundle or identity mismatch,
- an observation timestamp.

5.2 Promotions

Displayed base price and promotions should be stored separately when possible.

Example: if a product page displays \$499.99 and separately exposes an \$80 public code, the qualified base remains \$499.99 while the code is stored as a promotion. A derived \$419.99 checkout sensitivity may be shown only with wording such as “if the code applies at checkout.”

5.3 Snapshot price statistics

A release may calculate non-parametric descriptive statistics over the qualified-price population.

For MCDD-ICE-001 Screening V0.1, the qualified model-price snapshot has:

Q1 = \$199.00

Median = \$239.99

Q3 = \$299.99

These are descriptive statistics for the frozen qualified-price snapshot. They are not permanent category truths and are not consumer transaction-price statistics.

5.4 Longitudinal price statistics

The fields MODAL_14D_PRICE, PRICE_30D_RANGE, and PRICE_VOLATILITY_30D remain in development.

The current standard does not authorize calling a single observed sale the “normal” price. A modal 14-day price or 30-day volatility measure may be added after sufficient repeated observations exist and the collection window, retailer coverage, missing days, and promotion treatment are documented.

5.5 Outliers

An unusual price should be flagged and reviewed rather than automatically deleted solely because it falls outside an IQR rule.

The current MCDD uses explicit offer qualification and price-outlier classes. An IQR-based alert may become one input to future automation, but it is not presently a universal automatic deletion rule.

6. MISSING DATA, CONFLICTS, AND IMPUTATION

6.1 Missing data

Current decision data should preserve missingness. If a model lacks a qualified current offer, a verified dimension axis, an annual consumable cost, or a duration-qualified reliability estimate, that field remains missing or explicitly unresolved.

6.2 Conflicting data

Conflicting credible values should be recorded as a conflict or range until the discrepancy is resolved.

For example, the Euhomy Jet SE evidence contains conflicting output claims across related listings and brand material. The current dataset preserves the conflict instead of choosing the most favorable number.

6.3 Imputation

The current Purchase Mechanics standard does not use carry-forward imputation, peer-class mean imputation, or hedonic replacement adjustment to manufacture current decision facts.

Those methods may be appropriate for some statistical time-series purposes, but they must not be imported into a consumer decision dataset without a separate validation showing that the imputed value is suitable for the specific decision being made.

If an imputed or modeled value is introduced in a future version, it must be labeled as modeled, versioned, and kept distinguishable from an observed value.

7. EVIDENCE PROVENANCE

7.1 Evidence classes

At minimum, Measured Choice distinguishes:

OFFICIAL / MANUFACTURER — specification, warranty, manual, support policy, or manufacturer commerce page.

MAJOR RETAILER — retailer listing or retailer commerce observation.

INDEPENDENT TEST — hands-on testing by a source independent of the manufacturer.

OWNER SIGNAL — owner report or review, preferably with explicit duration when used for reliability context.

PUBLIC SAFETY / REGULATORY — recall, incident, certification, or other public safety record where relevant.

DERIVED — arithmetic or analytic output calculated from observed inputs.

MODELED ASSUMPTION — scenario or assumption not directly observed.

7.2 Evidence hierarchy is claim-specific

No source class is universally strongest for every question.

A manufacturer manual may be the strongest source for a cleaning procedure. An independent test may be stronger for first-hour production. A written warranty is stronger than an owner comment for warranty length. Owner evidence can reveal failure modes but usually cannot produce a population failure rate without a defensible denominator.

7.3 Duration-qualified owner evidence

Reliability evidence should store explicit ownership duration when the source states it.

Examples such as “three days,” “seven months,” or “eleven months” are materially different evidence. If duration is not stated, it should not be invented.

Small review samples may support a basic quality-floor discussion or identify a possible failure mode. They do not justify incidence claims such as “10% of units fail in the first year.”

8. DATA INSIGHTS AND PUBLIC EVIDENCE LABELS

Measured Choice may publish original or derived Data Insights, but the publication state of each field must be explicit.

The current public-label system is:

MEASURED DATA — directly computed from a defined dataset or recorded observation set.

EVIDENCE SIGNAL — supported by real observations but not mature enough to be treated as a population statistic.

OUR READ — a transparent Measured Choice interpretation of measured data and decision constraints.

IN DEVELOPMENT — a proposed field or score that is not yet authorized to influence the verdict.

TOOL BEHAVIOR — a statement about what an interactive tool currently does.

Examples in MCDD-ICE-001:

MEASURED DATA: qualified-price median, quartiles, offer observations, warranty terms where verified.

EVIDENCE SIGNAL: shared-platform crosswalks, cleaning burden, repairability path, duration-qualified failure-mode notes.

OUR READ: the statement that spending above roughly the current \$300 Q3 boundary should buy a distinct benefit the buyer values.

IN DEVELOPMENT: biofilm vulnerability tier and normalized CPSC incident ratio.

9. DECISION DERIVATIONS

9.1 Diagnostic ratios

Simple ratios such as price per claimed daily output may be used as diagnostics. They should not be treated as complete measures of value because they ignore reliability, storage, cleaning burden, support, size, features, and buyer fit.

9.2 Pareto screening

Pareto analysis may be used to identify non-dominated products under a defined set of dimensions. A product is non-dominated if no other product is at least as good on every included dimension and strictly better on at least one.

Pareto status does not mean a product is an “indisputable best choice.” The result depends on the included dimensions, measurement quality, product class, and buyer constraints.

9.3 Role and sub-frontier analysis

Products with materially different functions should not be collapsed into a single ranking merely to create one winner.

For the nugget-ice category, ordinary basket machines, smart variants, compact dispensers, high-output dispensers, and premium high-output products can occupy different buyer-specific frontiers.

9.4 Hard gates

Hard gates are applied before any preference trade-off that might otherwise make a weak-evidence product appear attractive.

Typical gates include:

- identity sufficiently resolved for the comparison,
- current offer qualified when current price is decision-critical,
- basic utility or performance floor,
- evidence quality sufficient for the strength of the claim,
- safety or recall findings handled appropriately,
- lifecycle or reliability evidence sufficient when the recommendation depends on long-term ownership value.

9.5 No arbitrary weighted score

The current Choice methodology does not authorize an arbitrary weighted total score as the primary selector.

Weights may be appropriate in personalized fit tools when they are transparent, sensitivity-tested, and subordinate to hard constraints. They should not be used to manufacture a universal category winner.

10. CHOICE ELIGIBILITY AND THE STOP RULE

A Choice designation must be earned after upstream evidence requirements are satisfied.

The current decision sequence is:

1. qualification and identity checks,
2. evidence-confidence checks,
3. cost-to-own normalization where decision-critical,
4. market/frontier analysis,
5. uncertainty analysis only when the underlying variables are mature enough,
6. knee-point or tie-break methods only if a valid Choice-eligible set remains.

For MCDD-ICE-001 V0.1, 0 of 9 finalists cleared the full upstream evidence plus lifecycle/cost-to-own requirements.

Therefore the release stopped. Monte Carlo uncertainty propagation, Kneedle, and DEA were not run to force a winner from an ineligible set.

The valid category-level result is:

ZERO CATEGORY-LEVEL CHOICE PRODUCTS — V0.1.

This does not mean every product is poor. It means the current evidence supports buyer-specific conditional recommendations more strongly than it supports a universal Choice designation.

11. RELIABILITY, REPAIRABILITY, AND COST TO OWN

11.1 Reliability

The standard distinguishes product-age evidence, duration-qualified owner evidence, manufacturer durability testing, warranty protection, and actual population reliability.

A long warranty may reduce buyer downside. It is not evidence that the population will survive for the full warranty term.

A manufacturer cycle-test claim may be useful evidence about a test program. It is not automatically a field failure-rate estimate.

11.2 Repairability

Repairability should describe the path that can actually be verified, for example:

OPEN PARTS — model-specific parts openly listed or sold.

SUPPORT-ONLY PARTS — parts may be supplied through support, often during warranty.

AUTHORIZED SERVICE — repair is directed through manufacturer or authorized service.

NO OPEN CATALOG FOUND — a targeted search did not surface an open model-specific catalog.

UNKNOWN — insufficient evidence.

“No open catalog found” must not be rewritten as “parts do not exist.”

11.3 Cleaning and maintenance

Cleaning burden should consider more than the advertised self-clean cycle. Relevant inputs may include cycle length, recommended frequency, descaling requirements, manual disassembly, removable access points, rinse steps, and recurring consumables.

Cleaning convenience should not be converted into a sanitation or health claim unless a separate validated method supports that conclusion.

11.4 Total cost of ownership

A cost-to-own model may include acquisition cost, energy, consumables, repair, replacement, and other ownership costs when those inputs are supportable.

If lifespan probabilities, repair incidence, or consumable dose counts are missing, the model should be partial or scenario-based rather than a fabricated expected-value total.

12. VERSIONING, REPRODUCIBILITY, AND CHANGE CONTROL

12.1 Release-based reproducibility

Each decision release should identify:

- dataset ID,
- dataset version,
- schema version,
- method version,
- frozen configuration count,
- raw observation count,
- offer observation count,
- collection dates,
- freeze date,
- known limitations,
- QA status,
- superseded release.

This allows a reader or auditor to identify the data state used for a conclusion.

12.2 Reopen conditions

A published decision should state what can change it.

Typical triggers include:

- material current-price movement,
- newly resolved model identity,
- new recall or safety information,
- duration-qualified reliability evidence,
- parts or repair-path changes,
- warranty changes,
- cost-to-own evidence that changes the lifecycle comparison,
- a new configuration that materially changes the frontier.

12.3 Corrections

A correction should not silently erase the fact that an earlier release existed. The release manifest should record what release supersedes what, and a public technical page should maintain a concise change log for material corrections.

13. QUALITY CONTROL

The present standard does not publish invented accuracy thresholds such as “master data must be 99.5% correct” unless those thresholds have been operationally defined and measured.

Current QA relies on:

- explicit source-class tracking,
- manual checks on decision-material fields,
- conflict preservation,
- missing-field preservation,
- qualified-offer rules,
- role/sub-frontier comparison,
- evidence maturity notes,
- release-level completeness reporting,
- stop rules when the evidence is not sufficient.

Future automated QA may add sampled error-rate measurement, transformation logs, or stronger temporal storage. Those capabilities should be documented only after implementation.

14. PUBLICATION STANDARD

The technical publication layer should be designed as a primary-source research asset, not only as a PDF attachment.

A finished paper should normally include:

- a stable canonical HTML page,
- title and paper ID,
- version and publication date,
- abstract,
- definitions,
- methods,
- results or worked examples,
- limitations,
- source mapping,
- change log,
- data dictionary where relevant,
- downloadable CSV/JSON where rights and data quality permit,
- PDF mirror for archival/download use,

- a short “How to cite this release” block.

Structured data may be added when it accurately describes the visible page and the underlying dataset. Structured data is not a substitute for visible methods and evidence.

The objective is to make Measured Choice an original source for facts it actually measured or derived. No publication format can guarantee citation by search engines, journalists, or AI systems.

15. PRIVACY AND PERSONALIZED TOOLS

Category-level datasets and technical papers may be public. Sensitive personal inputs do not need to be.

The Living Within Your Means prototype uses take-home pay, essential expenses, required debt payments, savings/reserve target, other commitments, cash set aside, and user-defined cushions to calculate a purchase-fit read.

The method is public; the personal values should remain local to the browser where practical and should not be sent to analytics by default.

The current calculator intentionally does not use 50/30/20 as the controlling affordability rule. A benchmark may be shown as optional context later, but the buyer’s stated obligations and cushions control the result.

16. CURRENT LIMITATIONS

The current Purchase Mechanics / MCDD implementation has important limitations:

- The market frame is versioned and broad, but it is not claimed to be an exhaustive global census.
- Current prices are offer observations, not transaction-price data.
- Longitudinal 14/30-day pricing is not yet mature for the nugget-ice category.
- Shared-OEM/platform resolution is partial and confidence varies by product.
- Long-horizon reliability evidence is sparse across the nugget-ice finalist set.
- Owner-review evidence does not provide a defensible population failure denominator.
- Cost-to-own modeling is incomplete for most finalists.
- Biofilm vulnerability scoring is not validated or published.
- A normalized CPSC incident ratio is not published because a defensible exposure denominator is not yet available.
- The current system uses release snapshots rather than a fully implemented bitemporal datastore.
- The standard does not yet publish measured operational error rates for the data pipeline.

These limitations are part of the method. They are not defects to be hidden. They define what the current release can and cannot support.

APPENDIX A — NUGGET-ICE RELEASE EXAMPLE

Dataset: MCDD-ICE-001

Frozen configuration count: 84

Choice-analysis finalists: 9

Category-level Choice products in V0.1: 0

Qualified-price snapshot median: \$239.99

Qualified-price Q1: \$199.00

Qualified-price Q3: \$299.99

Raw observations in release 0.4/0.5 baseline: 602

Offer observations in release 0.4/0.5 baseline: 89

Primary unresolved category-level issue: long-horizon reliability and complete lifecycle/cost-to-own evidence.

The example demonstrates the purpose of the stop rule. A dataset can be complete enough to make useful conditional recommendations while still being incomplete for a universal Choice designation.

APPENDIX B — PUBLIC CLAIM LABEL DEFINITIONS

MEASURED DATA

A value or fact directly observed or computed from a defined dataset and a stated release.

EVIDENCE SIGNAL

A real observation or pattern that is decision-relevant but not mature enough to support a population statistic.

OUR READ

A Measured Choice interpretation of measured data, buyer constraints, and trade-offs. It must not be presented as a directly measured fact.

IN DEVELOPMENT

A proposed metric, model, or field that is not authorized to influence the published verdict.

TOOL BEHAVIOR

A statement describing the actual behavior of an interactive tool in the audited version.

APPENDIX C — SOURCE BASIS FOR THIS PUBLICATION DRAFT

This publication draft is based on the current MCDD architecture and release records in the OIE V1.0 — Opportunity Intelligence Engine workbook, including MCDD Config Master, MCDD Raw

Facts, MCDD Offers, MCDD Screening V0.1, MCDD Decision Dataset V0.1, MCDD Choice Analysis V0.1, MCDD Publication QA V0.1, and MCDD Release Manifest.

It also audits and substantially narrows the earlier research document “Purchase Mechanics: Data Methodology & Reproducible Research Framework.” Future-state concepts from that document remain research inputs unless this publication standard explicitly adopts them.

END OF PUBLICATION DRAFT V0.1