

Vol. 1 — How Unsold grades

Methodology of record: GRADING.md v2.2.0 · unsoldlife.com · typeset from source, nothing paraphrased

This is the whole method, in public. No secret sauce, no "proprietary algorithm" — that phrase is the disease Unsold exists to cure. Every number below lives in `src/engine/score.ts` and `src/engine/ingredients.ts`, and every grade in the catalog can be reproduced on your machine in seconds.

What the letters mean

Grade	Plain language
A+	Top-of-range clinical doses, best-in-class forms, nothing hidden. Rare on purpose.
A / A-	Dosed like they actually read the research. Buy with confidence.
B+ / B / B-	Mostly honest and mostly effective, with a soft spot or two — a light dose, a cheaper form, one thing undisclosed.
C+ / C / C-	Middling. The right ingredients shy of the doses that make them work, or fine on paper and cheap in the details. Not a scam, not a standout.
D+ / D / D-	Sprinkles of the good stuff priced like the full dose. Mostly filler and fairy dust.
F	Marketing with a Supplement Facts panel attached.

The cutoffs are the familiar academic ones (97/93/90 for A+/A/A-, and so on down). A grade is a compressed summary — the per-ingredient breakdown is the actual argument, and it is always one tap away.

The pipeline

Each product runs through the same six steps, in order, every time:

1. **Dose normalization.** Every labeled dose is converted into the reference ingredient's canonical unit (mg ↔ mcg, IU → mcg for vitamin D). A dose we cannot honestly convert is not guessed at — it is set aside.
2. **Clinical-range banding.** The normalized dose is compared against the clinically studied range in the reference table: effective (interpolated 90–100 by position inside the range), above-range (86), excessive (68), slightly under (62), underdosed (42), trace (24), meaningless trace (12). Concealed doses score 34.
3. **Form-quality modulation.** Ingredients are not interchangeable chemicals. Magnesium glycinate absorbs; magnesium oxide mostly passes through. Each form carries a 0–1 quality factor, and a cheap form can dock up to 40% of the band score.

4. **Blend-concealment penalties.** On top of each hidden ingredient's low score, the whole grade is shaved in proportion to how much of the formula is hidden inside proprietary blends (up to 40% of the grade), plus a flat extra penalty when a hero active — the ingredient you are actually buying the product for — is the thing being hidden.
5. **Weighted mean.** Ingredient scores roll up into one 0–100 number, weighted by how much each ingredient matters (hero actives up to 2.0, incidental fairy dust as low as 0.6).
6. **Letter cutoffs + deterministic flags.** The number maps to a letter, and pure rules over the computed scores emit the flag chips: *proprietary blend hides the doses*, *underdosed hero ingredient*, *cheap ingredient form*, *most of the formula is hidden*, *every dose disclosed*, and so on. The one-line verdict comes from a fixed template table. No step anywhere in this pipeline involves randomness, the clock, the network, or a model generating text.

The stance

Why blends are penalized. A hidden dose is worse than a known-bad dose. A known-bad dose you can at least see; a hidden one asks you to pay for a promise. "Proprietary blend" is not a formulation secret — the FDA already requires every active in the blend to be named, just not quantified. The only thing concealed is the part that lets you judge value. So concealment is graded as what it is: a choice to keep you from checking.

Why forms matter. Two labels can print "400 mg magnesium" and deliver wildly different amounts of usable magnesium. Grading the molecule and ignoring the form would grade the marketing, not the product.

Why unscored rows exist. If an ingredient is not in the clinical reference table, we show it and score *nothing*. We never invent a dose range we cannot cite. An unrecognized ingredient is neutral — it neither helps nor hurts the grade — because "we haven't reviewed it" is not evidence of anything.

What "unscorable" means. If a product contains *nothing* we can score — an all-novel formula — it does not get a fabricated grade. It gets an explicit "nothing recognizable to grade" result. A grade we cannot support is worse than no grade.

A worked example of the philosophy: the concealed-blend fix

For a while, the concealment math had a hole. The blend penalty counted hidden rows only when the hidden ingredient was *also* in the reference table. A row hidden in a blend **and** unrecognized escaped entirely — it didn't count as hidden mass, and in products where it was the *only* blend row, the product didn't even register as having a blend.

The predictable result: a multi-collagen product hid its entire active formula behind one row labeled "Proprietary Blend" and graded **B+** on its two disclosed token vitamins. A pre-workout parked a 13.5-gram stimulant matrix in unrecognized blend rows and graded **B** on the handful of ingredients it chose to disclose.

That is backwards. A label does not get *less* suspicious because the thing it hides is obscure — the less we can identify the hidden mass, the more the concealment is the story. The fix: concealment now counts **every** blend-hidden row, recognized or not. Unrecognized hidden rows still get no per-ingredient score (we never invent a range), but their mass counts toward the "most of the formula is hidden" penalties and flags. Those two products dropped to **C** and **C-**, and 47 other products with the same pattern moved down with them. No product

was special-cased; the rule changed, the math followed.

Reproducibility

The engine is a pure function. Same label data in, same grade out — no randomness, no dates, no network, no LLM.

```
```bash npm run verify # 15 checks: worked examples, determinism, flag rules,
npm test # unit tests + a full-catalog snapshot test ```
```

The snapshot test (`src/engine/snapshot.test.ts`) records every product's grade, numeric score, concealment fraction, and flags. Any tuning change — a constant, a reference range, a data edit — shows up in the pull request as a complete, line-by-line diff of its consequences. CI runs `verify`, tests, and a build on every PR and posts the grade distribution as a comment. If a grade in the app ever disagrees with what `npm run verify` prints on the same commit, that is a bug, and we want to hear about it.

## What grades are NOT

- **Not a safety assessment.** An F does not mean a product is dangerous and an A+ does not mean it is safe *for you*. Interactions, contraindications, allergies, pregnancy, medication — none of that is in scope.
- **Not medical advice.** Unsold grades whether a label delivers the doses and forms the clinical literature studied. Whether you should take anything is a conversation for a clinician who knows your history.
- **Not a review of the brand.** We grade formulas, not companies. A brand can earn an A on one SKU and an F on the next; several do.
- **Not permanent.** Labels reformulate. Each grade is a dated snapshot of a specific panel from the source cited on its result page. If the tub in your hand prints a different panel, the label wins — confirm against the physical product before you buy.