



The 30-Day Terminology Audit

20 assets · 2 reviewers · 30 days — the pilot that turns a language opinion into a findings report

The smallest honest version of a terminology program: audit the twenty assets clients actually see, with two reviewers, inside thirty days. The output is a findings report and a seeded standard — evidence, not opinion. The audit measures the corpus, never the authors: no names in findings, severity attaches to assets.

The shape: 20 / 2 / 30

- **20 assets** — the client-facing set: top knowledge articles by traffic, the sales deck, the onboarding guide, the proposal template, recent social posts. Twenty is enough to establish rates and patterns; it is deliberately not a boil-the-corpus exercise.
- **2 reviewers** — one who owns content, one who doesn't (fresh eyes catch normalized idiom). Both score independently against the same inventory, then reconcile; disagreements are findings too — they mark Tier-2 territory your standard must decide.
- **30 days** — week 1 select and baseline, weeks 2–3 review and reconcile, week 4 write findings and seed the standard. A pilot that runs longer than a month becomes a program without ever having earned one.

Asset selection rubric

Pick by exposure, not convenience. Score each candidate asset 1–3 on **reach** (how many client eyes), **authority** (does it speak *as* the company), and **longevity** (how long it stays in circulation); take the top twenty by total. Tie-break toward diversity of surface — an audit of twenty KB articles finds one pattern; an audit spanning KB, sales, delivery, and social finds the culture.

The scoring sheet

One row per finding. Findings cite the term inventory by entry number so both reviewers flag identically.

Field	Values	Note
Asset	id/name	From the selection list
Location	section / slide / paragraph	Enough to find it in one minute
Term	the exact phrase found	Verbatim, in quotes
Inventory #	1–40, or NEW	NEW = not in the inventory; candidate for your standard
Severity	S1 · S2 · S3	See scale below
Surface	KB · sales · delivery · social · internal	Where it lives
Proposed fix	replacement text	From the inventory, or drafted

Severity scale. S1 — client-visible and charged: a term with exclusionary or ethnic freight in an asset clients see (inventory Tier 1 hits on external surfaces). **S2 — client-visible and imprecise:** idiom or drift that costs clarity or retrieval on an external surface, and any Tier 1 hit on an internal one. **S3 — internal drift:** imprecision on internal surfaces; fix opportunistically. Severity is about *exposure* × *freight*, which is why the same term can score differently on a sales deck and a runbook.

Reading the results — the three rates

1. **Findings per asset** (density): the headline number, and the re-audit's before/after.
2. **S1 count** (exposure): every S1 is one screenshot away from being someone else's social post; the pilot's most persuasive line for leadership is usually "we found *n* of these."
3. **Consistency rate** (retrieval): of concepts appearing in 3+ assets, the share using one canonical name. This is the number that ties language to cost per answer — every sub-100% concept is a split retrieval key.

Outputs

- **Findings report** — the three rates, the top patterns, the S1 list with proposed fixes, and the NEW terms the inventory didn't cover. Severity-ordered, author-blind.
- **Prioritized swap list** — every S1 and S2 with its one-line fix; the S1s get fixed inside the pilot window (finding a problem you then ship for another quarter is worse than not looking).
- **Seeded terminology standard** — the template's §2 and §3 tables, filled from what the audit actually found. The standard's first draft should smell like your corpus, not like this toolkit.
- **The re-audit date** — same twenty assets (or their successors), two quarters out. The delta is the program's proof.

After the pilot

The findings report is the champion kit's evidence base — the deck's "pattern" section quotes your own corpus instead of industry examples. From there the four-phase program (guide §*Running the program*) proceeds: ratify the standard, install the gate line-items, enable. If the pilot found little, say so and stop — a clean audit is a real and creditable result, and the gate addendum alone will keep it that way.

Contextkeeping is a KCS-aligned operating model. KCS® is a service mark of the Consortium for Service Innovation; this material is not affiliated with or endorsed by the Consortium.