



DOCKET / RESEARCH & PRACTICAL GUIDANCE

How to Re-Screen Targets When Investment Criteria Change

Build a more reviewable acquisition pipeline with Docket. Our deal-origination software and managed research help private equity teams screen companies against a mandate and understand the evidence behind each finding.

Published by Docket · 2026-09-27 · acquisition target screening · private equity · investment criteria · mandate versioning · re-screening · deal sourcing · target universe · decision audit trail · evidence standards

Original article: <https://docket.capital/articles/rescreen-targets-when-investment-criteria-change>



In this article

1. Executive Summary
 2. Introduction and Background
 3. Define the Mandate Version and Its Change Classes
 4. Classify Changes by Their Decision Effect
 5. Run Impact Analysis Before Choosing a Rerun
 6. Execute the Re-Screen and Preserve Decision History
 7. Choose an Execution Arrangement and Handle Exceptions
 8. Preserve Comparability, Approval, and Reporting Context
 9. Data Analysis and Evidence
 10. Implications and Future Directions
 11. Frequently Asked Questions (FAQs)
 12. Conclusion
-

Executive Summary

An acquisition target should be re-screened when an approved change could alter its eligibility, priority, evidence status, or next action. The unit of control is a **mandate version**, not a revised spreadsheet name. It contains the questions, field definitions, thresholds, exclusions, required evidence, decision logic, approval, and effective date that governed a screen. A prior decision remains a historical event; a changed mandate produces a new decision event linked to the first. This approach adapts NIST's controlled baseline concept and the World Wide Web Consortium's provenance vocabulary as design patterns, not as private equity requirements. (csrc.nist.gov) (www.w3.org)

The practical sequence is to freeze the old version and evidence cutoff, classify each change, map every changed rule to the fields and records it can affect, then select **no rerun, selective rerun, or full rerun**. A wording clarification with unchanged meaning may need a documented no-rerun decision. A threshold change can often use retained values, provided the data and all downstream dependencies are complete. A new ownership exclusion or sector taxonomy may require new research and a wider rerun. An incomplete dependency map makes a selective rerun unreliable. Published investor criteria show why the distinctions matter: firms specify different size ranges, ownership types, geographies, and add-on exceptions. (www.gppfunds.com) (www.mpepartners.com) (www.southfieldcapital.com)

For a **fictional 500-record universe**, this report works through a change from an inclusive revenue floor of \$5 million to \$8 million, with the upper bound unchanged. It supplies a diff template, a dependency matrix, boundary tests, and a transition ledger. The outcome cells remain blank until actual records are evaluated. The team reports pass to fail, fail to pass, unchanged, and now unknown as counts divided by the appropriate eligible population. It separately reports records awaiting evidence. Treating missing size data as a failed threshold would confuse uncertainty with ineligibility; ordinary null comparisons in PostgreSQL illustrate that distinction. (www.postgresql.org)

Finally, funnel and cohort dashboards must carry a **comparability flag**. Two pass rates are directly comparable only when their mandate versions, universe definitions, evidence cutoffs, and missing-data policies are compatible or have been restated on a common basis. Statistics guidance makes a similar methodological point about explaining material changes and constructing a consistent back series where possible. (osr.statisticsofauthority.gov.uk) (osr.statisticsofauthority.gov.uk) A changed mandate may alter today's shortlist without implying that earlier research was wrong. The retained old decision explains what the team knew and applied at the time. (www.w3.org)

Introduction and Background

Private equity origination teams often refine a live market map. A fund may move its size floor, add a geographic carveout, narrow a sector definition, change the ownership profile it will consider, or introduce a new exclusion. The research operations problem begins after a screening rubric is already in use: which target records need fresh work, which earlier screens are still valid for their original purpose, and how should a team compare the resulting pipeline? Published fund pages demonstrate that these dimensions are operational rather than abstract. **Great Point Partners** publishes a \$10 million to \$100 million revenue range and a \$2 million to \$12 million earnings before interest, taxes, depreciation, and amortization (EBITDA) range (www.gppfunds.com); it also accepts minority and majority positions (www.gppfunds.com). **MPE Partners** distinguishes platform geography from add-on geography (www.mpepartners.com).

A rubric design guide can explain how to define observable criteria and missing-answer treatment. Docket's existing screening guide does that. (docket.capital) This report starts at the next step: controlling a change to an already applied rubric. It treats a screen as a dated decision based on a specific version of rules and a defined body of evidence. It does not assess which investment thesis is financially attractive.

The recommended control is deliberately lightweight. A team needs a version identifier, an approved diff, a dependency map, a re-screen plan, and immutable decision events. It need not imitate a security compliance program. The National Institute of Standards and Technology (NIST) defines a baseline in the context of federal information-system configuration management; the transferable idea is an agreed reference point whose changes are explicit. (csrc.nist.gov) (csrc.nist.gov) The World Wide Web Consortium (W3C) PROV model offers relationships for revisions, derivations, attribution, and time. Those relationships provide a useful vocabulary for linking a new screen to the old one, the changed rules, the evidence, and the reviewer. (www.w3.org) (www.w3.org)

The reader should expect a decision procedure, not a universal numerical threshold for materiality. The cost of re-screening, risk of missed eligible companies, and data coverage differ by mandate. A version number alone cannot make an ambiguous phrase such as "attractive niche" measurable. The team must resolve meaning in the rubric itself, test it on boundary cases, and preserve any approved human exception separately from the rule.

Define the Mandate Version and Its Change Classes

What the version freezes

A **mandate version** is the complete, executable specification used for a screening decision. At minimum, it freezes the population definition; questions and accepted answers; units and period definitions; lower and upper thresholds with inclusive or exclusive endpoints; sector and geographic taxonomies; ownership rules; disqualifiers; evidence requirements; scoring weights; treatment of unknowns; and rules that turn answers into pass, fail, priority, or review states. The version also names its approver and effective date. PostgreSQL's explicit treatment of inclusive BETWEEN endpoints is a useful reminder that even a seemingly obvious threshold can contain an unstated decision. (www.postgresql.org)

Version 2.0.0 could represent a changed eligibility rule, **2.1.0** a new evidence question that leaves eligibility logic intact, and **2.1.1** a wording correction proven not to change interpretation. This is a local convention inspired by Semantic Versioning, whose major, minor, and patch numbers describe compatibility in software interfaces. A PE team should publish its own mapping and never assume the numbers explain a change without a diff. (semver.org) DataCite likewise leaves the judgment of major versus minor changes to data stewards while retaining version relationships. (support.datacite.org)

Use this **mandate diff template** for every proposed release:

- **Identity:** mandate ID, old version, proposed version, requester, and reason.
- **Timing:** request date, approval date, effective date, and whether any restatement is authorized.
- **Rule diff:** old and new question, definition, unit, threshold, exclusion, precedence, and score effect.
- **Evidence diff:** accepted source types, freshness rule, unknown treatment, and retained source requirement.
- **Impact:** affected fields, dependent calculations, stages, universe snapshot, and scope query.
- **Execution:** no, selective, or full rerun; owner; completion condition; and dashboard label.
- **Tests and decision:** boundary examples, expected outcomes, reviewer, approver, and exception path.

Keep the released version immutable. A correction is a new version or an annotated erratum, with a link to what it supersedes. The Semantic Versioning specification explicitly keeps released contents fixed, and W3C provenance describes revised entities as linked to predecessors. Both are useful models for reconstructing what was actually in force. (semver.org) (www.w3.org) A saved file date alone is insufficient if a score formula, exclusion lookup table, or evidence policy could have changed without that date changing.

Classify Changes by Their Decision Effect

Classify by possible effect on decisions, not by how small the edit looks on the page:

- **Clarification:** wording changes while the accepted cases and decision logic stay the same. Require paired examples showing equivalence.
- **Source rule:** a field now requires a more authoritative or fresher source. Prior answers may need evidence review even if the numerical threshold is unchanged.

- **Threshold:** a size, revenue, earnings, concentration, or score cutoff moves. Recompute candidates near the boundary and any score derived from that field.
- **Logic:** AND, OR, exception precedence, or weight changes. Recompute every record whose outcome could depend on the altered expression.
- **Geography:** country, headquarters, operating footprint, or platform versus add-on rules change. Define the geographic unit before mapping records.
- **Sector taxonomy:** categories merge, split, or acquire new exclusions. Reclassify records whose labels or source evidence cross the affected branches.
- **New disqualifier:** an existing pass can become fail or unknown. Search beyond the current shortlist because a disqualifier can also affect priority and outreach status.
- **Removed criterion:** some old fails may qualify. Revisit the excluded population rather than screening only the current pass list.

The classes are not mutually exclusive. A changed threshold plus a new source rule can require both recalculation and evidence collection. Real fund criteria illustrate why a universal “size” field is inadequate: **Edgewater Funds** states a \$3 million to \$40 million EBITDA range (www.edgewaterfunds.com), while **One Equity Partners** states a \$100 million to \$750 million revenue range (www.oneequity.com). Those measures cannot be swapped without a documented definition and evidence rule. **Southfield Capital** lists no size parameter for add-ons, and **Granite Bridge Partners** similarly publishes no minimum add-on size. A platform threshold should therefore not silently carry into an add-on screen. (www.southfieldcapital.com) (granitebridge.com)

The change request should also state whether it is prospective or retroactive. A prospective effective date governs future decisions while retaining old screens as historical. A retroactive restatement is a separate analytical run under the new rules using a specified historical evidence cutoff. Mixing those two exercises produces a false narrative about what a researcher could have known earlier. W3C distinguishes generation and invalidation times, a useful way to represent when versions and decisions become available or cease to be current. (www.w3.org) (www.w3.org)

Run Impact Analysis Before Choosing a Rerun

A re-screen plan begins with a **dependency graph**: changed rule to required field, field to evidence and transformations, transformation to eligibility and score, decision to pipeline action. Include manual overrides, imported source fields, deduplication, and any downstream shortlist or outreach list. Data build tooling illustrates the underlying idea: dbt's manifest records first-order parent and child relationships so affected resources can be identified. That is a technical analogy, not a claim that dbt defines PE screening practice. (docs.getdbt.com)

For each change, query the old universe, not just the current shortlist. Count records with a value inside the changed band; missing or stale values; unresolved identity or taxonomy matches; and records at every pipeline stage. Check whether a newly required field exists in retained evidence. A schema update alone does not fill historical rows in incremental data systems, a useful warning against assuming an added column is populated retrospectively. (docs.getdbt.com)

Table 1 is an impact matrix. The final column is a starting scope, subject to the dependency-map and evidence tests below.

Change class	Dependencies to inspect	Records to test first	Likely action
Clarification with proven equivalent interpretation	Rule text, examples, reviewer guidance	Boundary examples and prior exceptions	No rerun; log the equivalence test
Threshold or weight	Raw value, unit, period, score formula, evidence date	Changed interval and all records with missing values	Selective recomputation if the map is complete
New evidence standard	Source class, excerpt, collection date, review state	Every answer relying on a now-ineligible source	Selective evidence review; expand if source coverage is unknown
Geography or sector taxonomy	Headquarters and operating geography, category crosswalk, branch logic	Ambiguous and reclassified branches, including old fails	Selective reclassification or full rerun if crosswalk coverage is incomplete
New or removed disqualifier	Field availability, exception precedence, pipeline action	Entire universe, including excluded and active records	Full eligibility pass; collect missing fields separately

The matrix is a scoping tool, not permission to skip a record that the graph does not cover. A selective rerun is defensible only when the team can show that unaffected records do not depend on a changed input and that retained evidence is sufficient for the new rule. When a formula changed but stored historical rows still reflect old transformations, a rebuild may be necessary; dbt's incremental-model documentation describes that divergence in its own domain. (docs.getdbt.com) If the dependency map is incomplete, use a full rerun or a documented conservative expansion of scope.

Choose the rerun mode after a small **impact sample**. Test a prior pass, prior fail, unknown, exception, and record at each important pipeline stage. A no-rerun decision requires proof that semantics are unchanged, not merely a declaration that the edit is “minor.” A selective rerun needs a reproducible query or list of record identifiers and a reason each excluded record is unaffected. A full rerun evaluates the whole frozen universe against the new specification, then separately notes newly discovered companies outside that universe. This distinction keeps a mandate change from being mistaken for a change in market-map coverage.

Execute the Re-Screen and Preserve Decision History

A screening run should produce **new decision events**. Each event contains target identifier, mandate version, evidence cutoff, data snapshot or source references, execution time, result, score or reason codes, unknown fields, analyst or process identity, reviewer, exception approval, and link to the prior event. W3C PROV-O supports derivation and attribution relationships that can model this chain. (www.w3.org) (www.w3.org) The Public Company Accounting Oversight Board's (PCAOB) audit documentation standard asks auditors to record procedures, evidence, conclusions, performers, and reviewers. That is documentation inspiration here, not an audit obligation for PE origination. (pcaobus.org) (pcaobus.org)

A practical run sequence is:

1. **Freeze inputs:** store the old version, target-universe snapshot, old decision IDs, and evidence cutoff.
2. **Approve the diff:** record requester, rationale, approver, decision date, effective date, and whether historical restatement is authorized.
3. **Test the rule:** execute positive, negative, boundary, missing, contradictory-source, and exception cases before touching the full universe.
4. **Select scope:** run the dependency query and preserve both included and excluded record counts.
5. **Collect evidence:** fill newly required fields; preserve unsuccessful searches and their dates rather than writing a bare null.
6. **Compute outcomes:** apply disqualifiers before scores, then calculate priority under the new version.
7. **Review transitions:** inspect changed decisions and a sample of unchanged records for mapping errors.
8. **Publish a new view:** keep the prior events available; point the current view at the approved new event.
9. **Reconcile stages:** review active opportunities and outreach queues with the responsible owner before changing their operational status.

Docket's product page describes separate Triage, Scout, and Audit stages, source excerpts and collection dates, and preserved historical entries. Those are first-party product claims, useful context for a team comparing how to implement the procedure. (docket.capital) (docket.capital) Docket's evidence guide distinguishes a value with a URL from a value accompanied by a dated supporting excerpt and retained copy, and says a reviewer still must assess whether the excerpt supports the conclusion. (docket.capital) No product claim substitutes for the team's own mandate approval or exception policy.

Choose an Execution Arrangement and Handle Exceptions

Table 2 compares operating arrangements for the same re-screen control. The brand appears as a provider row because it offers target research and criteria management; its published model is described only to the extent the first-party page supports it.

Arrangement	How the change is executed	Record and review requirement
Internal spreadsheet	Freeze a copy, calculate affected-row query, append result columns or an event sheet	Versioned formulas, source references, reviewer signoff, and a retained prior workbook
Internal research database	Store immutable mandate and decision tables; run dependency query and new decision insert	Evidence snapshot, version foreign key, run ID, exception log, and dashboard comparability flag
Docket	Published product description says criteria, thresholds, disqualifiers, and new questions can be updated across a target list. (docket.capital)	Published description says source pages, excerpts, dates, and correction history are retained; verify the exact export and approval workflow for the mandate. (docket.capital)

The table compares control needs rather than price or investment results. Docket's published self-serve and managed delivery modes differ in who runs the work, while the version, evidence, and approval questions still belong to the investing team. (docket.capital) An internal implementation can meet the same record standard if it captures the fields and review events. The decisive question is whether an auditor of the team's own process can reconstruct why a target was included or excluded under each version.

Handle unknowns and exceptions explicitly

A new criterion often creates a third state, **unknown**, rather than an immediate pass or fail. If ownership is newly required and no acceptable source establishes it, mark unknown: ownership evidence pending. Record where the team looked, when, and the evidence rule still unsatisfied. Docket's screening guide warns that incomplete records and complete records can appear misleadingly comparable when missing answers are silently scored as neutral. (docket.capital) PostgreSQL's null-comparison behavior is a simple formal analogy for keeping unknown separate from false. (www.postgresql.org)

An approved exception is not a hidden rule edit. Store its target ID, version, reason, approver, duration, and scope. A single target exception should not change the global threshold; a global threshold change should not be disguised as a run of exceptions. Invest Europe's professional standards handbook describes checking proposals against restrictions in fund documents and retaining a written record of information considered in an investment decision. The practical lesson is to keep the fund's actual restrictions and decision evidence visible in the same review packet. (www.investeurope.eu) (www.investeurope.eu)

Preserve Comparability, Approval, and Reporting Context

A dashboard needs a **comparability label** on every funnel cohort. At a minimum, store mandate version, universe snapshot, evidence cutoff, missing-data policy, stage definition, and whether the cohort was prospectively screened or retrospectively restated. Label pairs comparable, restated comparable, or not directly comparable; include the reason. The UK Office for Statistics Regulation advises producers to assess the materiality of method changes, and the UK statistics code calls for explaining changes and creating a consistent back series when possible. These are analogies for honest pipeline reporting, not rules imposed on PE firms. (osr.statisticsauthority.gov.uk) (osr.statisticsauthority.gov.uk)

The **same-version comparison** is simplest: compare cohorts evaluated under identical rules and equivalent evidence cutoffs. A **restated comparison** replays both cohorts under one version and a specified data snapshot, showing what would qualify on that common basis. A **historical decision comparison** asks what each team actually decided at the time; it is valid as a process history but is not a like-for-like eligibility-rate comparison when criteria differ. Keep these views separate. W3C's provenance constraints describe multiple identified entities and links among document versions, a model for preserving both historical and restated screens. (www.w3.org)

Approval should precede the new version's effective date. The change log should include a short rationale, alternatives considered, impacted fields, planned rerun mode, acceptance tests, expected dashboard treatment, and an owner for unresolved research. Project Management Institute material on scope control recommends impact analysis before updating a plan and describes approvals as documented endorsements. That pattern is useful for mandate governance, but a PE team should tailor who can approve to its own fund documents and investment committee process. (www.pmi.org) (www.pmi.org)

Test cases should accompany the version, especially when boundaries or precedence change:

- **Just below lower bound:** fail the size rule under the new version.
- **Exactly on lower bound:** pass only if the rule explicitly includes the endpoint.
- **Just above lower bound:** pass the size rule, subject to other criteria.

- **Exactly on upper bound:** follow the stated inclusive or exclusive convention.
- **Missing value:** unknown and research queued, never quietly converted to fail.
- **Conflicting values:** apply the evidence hierarchy or mark unresolved.
- **Platform versus add-on:** apply the correct branch before testing size or geography.
- **Prior exception:** preserve old approval and require a new-version applicability decision.
- **Removed disqualifier:** retest old fails as well as old passes.
- **Duplicate target identity:** link records before counting universe transitions.

A validation suite can store such assertions alongside a fixed data batch. Great Expectations documents both grouped expectation suites and a validation definition linking a batch to a suite; these are technical examples of making tests repeatable. (docs.greatexpectations.io) (docs.greatexpectations.io) The business test is stronger than a green software check: a reviewer must confirm that each expected result follows the approved written rule.

Data Analysis and Evidence

The quantitative problem is to separate **rule effects** from **coverage effects**. For a worked **fictional example**, start with 500 unique target records frozen at cutoff (T_0). Version 1.0.0 accepts annual revenue from \$5 million through \$25 million, both endpoints inclusive; version 2.0.0 raises only the lower bound to \$8 million. Other criteria are held constant for this illustration. These numbers are example parameters, not observed market results. The live data must supply the counts.

For a record with an accepted revenue value, the changed interval is \$5 million inclusive to below \$8 million. Old passes in that band may become new fails, subject to other criteria. Prior fails below \$5 million cannot become passes solely because the floor rose. A missing revenue value remains unknown unless the old decision had another independent disqualifier that still resolves the result. The precise treatment of missing values and endpoint inclusion must be explicit, as database comparison semantics make clear.

(www.postgresql.org) (www.postgresql.org)

Table 3 is the transition ledger to fill after the re-screen. Let N_{eligible} be the number of the 500 records for which both versions can produce a comparable decision from qualifying evidence. Let $N_{\text{total}} = 500$. Do not enter a percentage without its denominator.

Transition	Count to enter after run	Percentage formula	Interpretation
Pass to fail	n_{pf}	$100 \times n_{\text{pf}} / N_{\text{eligible}}$	Rule or evidence change removed prior passes
Fail to pass	n_{fp}	$100 \times n_{\text{fp}} / N_{\text{eligible}}$	May be zero for this single upward-floor change; investigate any positive count
Unchanged pass or fail	n_{same}	$100 \times n_{\text{same}} / N_{\text{eligible}}$	Comparable classifications retained
Now unknown	n_{unknown}	$100 \times n_{\text{unknown}} / N_{\text{total}}$	New evidence or changed validity rule prevents a current decision

Transition	Count to enter after run	Percentage formula	Interpretation
Outside comparable denominator	500 - N_eligible	$100 \times (500 - N_{\text{eligible}}) / 500$	Missing or noncomparable records; explain each reason

The first three rows partition comparable decided records when the old and new results are both pass or fail; thus $n_{\text{pf}} + n_{\text{fp}} + n_{\text{same}} = N_{\text{eligible}}$. The now-unknown row is reported against all 500 records and must not be added to the first three percentages. If unknowns were already present under the old version, record unknown to pass, unknown to fail, and unknown unchanged in a supporting reconciliation rather than hiding them inside n_{same} . This preserves the 500-record control total and makes the denominator auditable. A changed evidence rule can move a record to unknown even when its stated revenue is unchanged. (docket.capital)

Published population data also warn against treating a market-map row count as a market-size estimate. The U.S. Census Bureau reported **5.58 million employer firms with fewer than 500 employees in 2023**, but those firms span sectors and sizes beyond any one acquisition mandate. (www.census.gov) An establishment is one physical location; deduplicate rows to company entities before computing transitions. (www.census.gov)

Report three companion measures. **Re-screen coverage** is completed new-version screens divided by the 500 frozen records; the numerator includes documented pass, fail, and unknown outcomes. **Evidence completion** is records with every newly required field supported divided by records that require that field. **Decision movement** is $n_{\text{pf}} + n_{\text{fp}}$ divided by N_{eligible} , with the numerator and denominator shown. If a sector crosswalk or ownership source is incomplete, publish the affected count and defer any claim that selective rerunning found every changed decision. The statistics code's emphasis on completeness, coherence, and comparability is a useful external frame for these disclosures. (osr.statisticsauthority.gov.uk)

Implications and Future Directions

A mature re-screening process turns a strategy change into a controlled analytical event. It does not erase the old market map. It preserves the earlier version, records the requester and approver, shows the scope query, reuses evidence only where still valid, and creates new decisions with links back to prior ones. W3C's model of revision, derivation, and attribution provides a compact conceptual architecture for this chain. (www.w3.org) (www.w3.org)

The largest practical risk in selective reruns is an incomplete dependency map. A changed ownership definition might touch identity resolution, parent information, an exclusion, a score, and an outreach decision. If the map contains only the score column, unchanged-looking records can be skipped incorrectly. A full rerun is the defensible choice when the team cannot demonstrate coverage of those links. A second risk is evidence vintage: a retained value may be precise yet too old for the new rule. The new version must state whether it can use old evidence, and dashboards must retain the cutoff used for each cohort. (docs.getdbt.com) (osr.statisticsauthority.gov.uk)

The organization should periodically test change controls on a small set of reviewed targets. Compare the expected transitions to the new run, sample unchanged outcomes, and reconcile the count of all frozen records. Preserve reasons for manual overrides and review whether recurring exceptions signal an ambiguous global criterion. **Versioning improves traceability; it does not cure ambiguity.** The criterion must still define observable facts and decision rules, as Docket's rubric guide emphasizes. (docket.capital)

As research systems add more automation, the useful design choice is an event log in which each answer, screen, review, and exception has a version and provenance link. An internal database can implement this pattern. The team should test exports and historical reconstruction against its own requirements before relying on any tool for committee materials.

Frequently Asked Questions (FAQs)

Should every mandate edit trigger a full re-screen?

No. A clarification can receive a documented no-rerun decision if test cases prove it leaves outcomes unchanged. A selective rerun is suitable when the rule dependencies are complete and all excluded records can be shown unaffected. A new disqualifier, removed criterion, incomplete crosswalk, or unknown dependency can justify a full pass over the frozen universe. The decision belongs in the approved change log, not in an analyst's memory. (csrc.nist.gov) (www.pmi.org)

What happens to targets already in the pipeline?

Keep their historical screening decision and stage history. Create a new-version result and route any transition to the responsible deal owner for review. A new fail may change current priority, but it does not retroactively rewrite the earlier research. Published fund processes show that screening can hand opportunities between sourcing and transaction teams, making stage ownership part of the impact analysis. (www.hkwinc.com) (www.hkwinc.com)

How should a team prioritize targets after a strategy change?

First separate eligibility from ranking. Revisit newly eligible old fails and newly ineligible passes, then queue unknowns for the evidence that could resolve them. Within the new eligible set, apply the approved scoring weights and evidence-completeness policy. Keep legacy priority visible for historical explanation but label it with its old version. Docket's rubric guidance explicitly distinguishes disqualifying conditions from weighted scores and warns against comparing incomplete records without showing completeness. (docket.capital) (docket.capital)

Can the earlier and later pass rates be compared?

Only with a comparability statement. If the target universe, criteria, evidence cutoff, and missing-data treatment differ, raw pass rates answer different questions. Restate both cohorts under one version and cutoff when a like-for-like estimate is needed, and display the original historical decisions separately. Statistics guidance recommends explaining material methodology changes and constructing a consistent back series where possible. (osr.statisticsauthority.gov.uk) (osr.statisticsauthority.gov.uk)

What belongs in an audit trail for acquisition target screening decisions?

Record the mandate diff, requester, rationale, approver, effective date, impacted fields, scope query, test cases, run ID, evidence cutoff, source references, old and new decisions, reviewer, and any exception. The PCAOB's documentation principles provide an analogy for recording procedures, evidence, conclusions, performer, and reviewer; they are not PE-specific requirements. (pcaobus.org) (pcaobus.org)

Conclusion

Re-screening begins by treating the investment criteria as a versioned decision specification. Freeze the old mandate and universe, describe the change at rule and evidence level, map its dependencies, and choose the smallest rerun whose coverage can be demonstrated. If the dependency map or required evidence is incomplete, expand the run and show what remains unknown. The old result stays in the record; the new result is a dated event with a link to it.

The working artifacts are a mandate diff, impact matrix, re-screen plan, boundary-case tests, transition ledger, and dashboard comparability flag. The fictional 500-record example deliberately leaves outcome cells empty because only a real run can supply pass-to-fail and other transition counts. Percentages should carry their denominators, and unknowns should remain visible rather than being forced into pass or fail.

A changed shortlist does not establish that the earlier team screened incorrectly. It can simply reflect a changed mandate, a changed evidence standard, or new information. A reviewable record lets an origination team explain which of those occurred, decide what to research next, and compare cohorts only on a stated common basis. It also gives the next mandate revision a tested starting point instead of an undocumented reconstruction of the last one.

Source links

Links cited in this article, in order of first appearance. Full addresses are provided for readers using extracted text.

1. https://csrc.nist.gov/glossary/term/baseline_configuration
2. <https://www.w3.org/TR/prov-o/>
3. <https://www.gppfunds.com/private-equity/investment-criteria/>
4. <https://www.mpepartners.com/criteria/>
5. <https://www.southfieldcapital.com/investment-criteria>
6. <https://www.postgresql.org/docs/current/functions-comparison.html>
7. <https://osr.statisticsauthority.gov.uk/guidance/guidance-for-producers-when-making-changes-to-statistical-methods/>
8. <https://osr.statisticsauthority.gov.uk/publications/the-code-of-practice-for-statistics/pages/5/>
9. <https://www.w3.org/TR/prov-constraints/>
10. <https://docket.capital/screening-criteria>
11. <https://csrc.nist.gov/pubs/sp/800/128/upd1/final>
12. <https://semver.org/>
13. <https://support.datacite.org/docs/versioning>

14. <https://www.edgewaterfunds.com/investment-criteria/>
15. <https://www.oneequity.com/approach/>
16. <https://granitebridge.com/strategy>
17. <https://docs.getdbt.com/reference/artifacts/manifest-json>
18. <https://docs.getdbt.com/docs/build/incremental-models>
19. <https://docket.capital/articles/acquisition-target-research-freshness>
20. <https://pcaobus.org/oversight/standards/auditing-standards/details/AS1215>
21. <https://docket.capital/>
22. <https://docket.capital/evidence-standards>
23. <https://www.investeurope.eu/professional-standards-handbook/planning-and-making-investments/>
24. <https://www.pmi.org/learning/library/scope-control-projects-you-6972>
25. https://docs.greatexpectations.io/docs/core/define_expectations/organize_expectation_suites/
26. https://docs.greatexpectations.io/docs/core/run_validations/create_a_validation_definition/
27. https://www.census.gov/library/stories/2026/05/small-business-week.html?trk=article-ssr-frontend-pulse_little-text-block
28. <https://docket.capital/articles/target-audit-sampling-worksheet>
29. <https://www.hkwinc.com/investment-criteria/>

THE PUBLISHER

About Docket

Docket provides deal-origination research software and managed research for private equity firms. We help investment teams investigate acquisition targets using structured screening criteria, retained sources and reviewable company evidence. Teams can work through a self-serve platform or use managed research, depending on how they want research delivered.

Research against a defined mandate

A useful target list needs more than company names. Docket focuses on the questions that determine whether a company fits an investment mandate, including the evidence needed to support or qualify each answer. Our research approach makes the connection between screening criteria, source material and conclusions visible to the team reviewing the work.

Triage, Scout and Audit

Docket's three named research agents perform complementary tasks. Triage resolves company identity and screens fit. Scout collects sourced answers against the mandate. Audit checks retained evidence, addresses contradictions and leaves unsupported answers visibly unresolved. This structure helps reviewers distinguish established findings from missing information and questions requiring further investigation.

Evidence that supports investment-team judgment

Our research library covers market mapping, screening criteria, private-company data, succession and ownership, source evaluation and evidence standards. These resources explain the methods and limitations behind origination research. Findings support a team's judgment; they do not establish that a company is for sale or guarantee a transaction or investment outcome.

Work with Docket

Visit Docket to explore the platform, managed research and the current contact path. Read about screening criteria, evidence standards, AI research agents and market mapping.

Public examples are illustrative unless explicitly identified otherwise. Research preparation and authorized outreach are separate activities; confidential target lists and customer outcomes should never be inferred from an educational example.

Website: <https://docket.capital>

This article is educational. Verify consequential medical, legal, financial, regulatory and technical decisions with appropriate professionals and the cited primary sources. Company services are described separately from the article's research findings.