

Relativity vs Reveal: A Buyer's Guide to Data Scoping, ECA, Scale and Active Hosting

A practical comparison of Relativity and Reveal across data scoping, early case assessment, scale and active hosting, with published performance figures and a clear verdict on which platform suits which matter.

CONTENTS

- 01 Data scoping
- 02 Early case assessment
- 03 Scale
- 04 Active hosting
- 05 The verdict
- 06 Sources

Choosing between Relativity and Reveal is rarely a question of which platform is "better". It is a question of which platform is better suited to the work in front of you. The two have converged on many features in recent years, but they remain architecturally different in ways that matter when you are deciding where to host a large investigation, how early you want analytical insight, and how much of that insight you need before a single document is reviewed.

This guide compares the two platforms on the four dimensions that most directly affect cost, timeline and defensibility: data scoping, early case assessment, scale and active hosting.

KEY QUESTION

Are you buying a platform that helps you understand your data before review begins, or one that helps you manage review once it is already underway? Reveal and Relativity answer that question differently.

Data scoping: where each platform starts

Data scoping is the process of understanding what you have before committing to full processing and review. Done well, it reduces the review population, controls cost and produces a defensible record of the decisions made before the first document was coded. Done poorly, it results in over-collection, inflated hosting volumes and budget surprises.

Reveal: analytics before the first decision

Reveal's architecture is built around the principle that analytical insight should be available from the moment data lands in the platform. Rather than requiring processing to complete before analysis can begin, Reveal runs an extensive set of automated operations on ingestion, so that visualisations, concept clusters and communication maps are accessible on day one.

The practical effect is significant. A solicitor or case manager can open a new matter and immediately interrogate the data using plain-language queries, Reveal's concept search, or its Dashboard, Clusters, Heatmap and Communications views. Scoping becomes a genuine analytical exercise rather than a keyword-guessing exercise conducted in advance of processing.

What this means for buyers: if your matter requires rapid risk assessment, particularly in regulatory investigations where the client needs to understand exposure before making strategic decisions, Reveal's front-loaded analytics compress that timeline materially.

Relativity: scoping as a configured workflow

Relativity's approach is more modular. The platform offers powerful tools, including dtSearch and conceptual search through its analytics suite, but these typically require a deliberate configuration step before they are available on a new workspace. Scoping in Relativity tends to be a workflow that practitioners build, rather than a capability that is on by default.

That is not a weakness in all contexts. For firms with established Relativity workflows and experienced administrators, the modular approach provides greater control over how scoping is conducted and documented. The audit trail and saved search infrastructure are mature and well understood by courts and regulators.

- Relativity strengths for scoping: granular saved search logic, mature audit trails, a broad ecosystem of third-party analytics apps through the RelativityOne Marketplace.
- Reveal strengths for scoping: immediate visual analytics on ingestion, plain-language AI querying, no separate configuration step.

The gap narrows considerably for teams with strong Relativity administration capability. For teams without that resource, the difference in time to insight is real.

Early case assessment: depth, speed and strategic value

Early case assessment is where the two platforms diverge most sharply in philosophy. The question is not just which one has ECA features, but what ECA actually means in each context.

Reveal: ECA as a strategic starting point

Reveal positions ECA as the entry point for the whole matter strategy, not a preliminary technical step. Its AI, including generative search and agentic models, allows lawyers to query a dataset in plain language and receive concise, sourced answers rather than a list of documents. Reveal's published data reports 75% faster fact-finding in complex investigations and roughly a 50% reduction in documents requiring manual analysis.

The reusable AI model library is a differentiator worth noting. Firms can build classifiers once, save them to an organisational library and deploy them at the start of a new matter without retraining. For teams handling recurring matter types, such as financial misconduct investigations or employment disputes, this compounds in value.

Reveal is also building toward AI agents that can categorise matter severity, gather data across multiple systems and produce a summary report for a senior investigator before any human reviewer opens a document. For buyers evaluating platforms over a three to five year horizon, that trajectory is relevant.

Relativity: ECA that is structured and auditable

Relativity's ECA capability centres on its analytics engine, including email threading, near-duplicate detection, concept searching and clustering. These are mature, well-validated tools with a long track record in litigation. The platform does not natively offer the same generative AI querying layer, though third-party integrations can extend it.

Where Relativity has a genuine advantage is the auditability of ECA decisions. Every saved search, filter and analytical operation is logged in a way that is straightforward to export and present to opposing counsel or a regulator. For matters where the ECA methodology itself may be scrutinised, that record is valuable.

Dimension	Reveal	Relativity
Time to first analytical insight	Immediate on ingestion	After workspace configuration
Plain-language AI querying	Native, generative	Through third-party integration
Reusable AI models	Built-in model library	Requires manual project setup
ECA auditability	Good	Excellent, mature audit trail
Agentic AI roadmap	Active, publicly stated	Developing

The honest summary: Reveal delivers more analytical value earlier, with less configuration overhead. Relativity delivers more control over the ECA process and a more established defensibility record. Neither is the wrong answer; the right choice depends on which of those properties your matter demands first.

Scale: what each platform can actually handle

Scale is where Relativity has the most publicly documented, verifiable data. Its Active Learning performance baselines are published in detail, which makes a specific conversation possible about what the platform can and cannot do at volume.

Relativity's published scale limits

According to Relativity's Active Learning performance documentation, the recommended limits for a single Active Learning project are:

- Maximum documents in the classification index: 9 million
- Maximum coded documents: 1 million
- Maximum concurrent reviewers: 150

At one million documents the platform processes at approximately 304,878 documents per hour during index population, with an initial index build time of around 13 minutes. Model rebuilds occur every 20 minutes during active review, so the classification stays reasonably current without manual intervention.

For matters exceeding nine million documents, Relativity's own guidance recommends a slicing approach: splitting the population into subsets and running parallel Active Learning projects. That is workable, but it adds configuration complexity and requires experienced administration.

THE PRACTICAL IMPLICATION

Relativity at scale is powerful but not self-managing. Large matters require deliberate architectural decisions about how Active Learning projects are structured, and those decisions have downstream consequences for validation and defensibility.

Reveal's scale posture

Reveal does not publish equivalent benchmarks in the same granular format. Its stated capability covers matters of any size, with processing supporting over 900 file types and a supervised learning architecture that allows classifiers to be deployed at the start of ingestion rather than after a review population is defined.

The more meaningful scale consideration with Reveal is not raw document count but the speed of analytical insight at volume. Because Reveal's AI operations run automatically on ingestion, a large dataset does not require a separate setup phase before analysis can begin. For investigations where the data volume is uncertain at the outset, such as multi-custodian regulatory matters, this reduces the risk of being caught flat-footed during early strategy discussions.

Where Relativity wins on scale: deeply documented, court-tested performance at very large review populations.

Where Reveal wins on scale: faster time to analytical insight at any volume, without upfront architectural decisions about how to structure the review.

Active hosting: cost, control and ongoing review management

Active hosting refers to the ongoing management of a live review environment: queue management, reviewer monitoring, coding propagation and the operational overhead of keeping a hosted matter running efficiently. This is where platform choice has the most direct impact on day-to-day cost.

Relativity's Review Centre

Relativity's Review Centre is the operational hub for active hosted review, providing project-level monitoring of reviewer throughput, queue refresh rates and coding activity.

- Saved search queues refresh every 15 minutes when there is active coding activity.
- Prioritised review queues refresh when 20% of documents in the queue have been coded, or after 8 hours of activity, whichever comes first.
- Coverage Mode triggers a queue refresh every 100 documents coded, or when 5% of documents have been coded positive or negative.
- Review speed reporting is available in 15-minute increments, giving project managers granular throughput visibility.

This operational granularity is a genuine Relativity strength. For large review teams running multiple simultaneous queues, the ability to monitor and adjust at that precision is material, and it means cost tracking can be tied to actual reviewer activity rather than estimated throughput.

Reveal's active review environment

Reveal's active review environment follows the same AI-first principle as its ECA capability. Pre-built AI models can be deployed to push the highest-scoring content to the front of the queue automatically, reducing manual queue management. Reviewers work within a flexible interface that stays connected to the platform's visual analytics, so a reviewer can pivot from coding to cluster analysis without leaving the review environment.

The reusable model library matters here too. If a firm has already built and validated a relevance classifier for a matter type, that model can be deployed at the start of active review on a new matter, compressing the time before AI-assisted prioritisation is operational.

The key hosting cost consideration is data volume management. Both platforms charge on a per-gigabyte active hosting basis, but the ability to cull aggressively before data reaches the review population directly affects that cost. Reveal's earlier analytical insight tends to produce smaller review populations; Relativity's more controlled workflow tends to produce more precisely defined ones. The downstream difference depends almost entirely on how well the scoping and ECA phases were executed.

Cost driver	Reveal approach	Relativity approach
Pre-review culling	AI-driven, automatic on ingestion	Configured keyword and analytics filters
Queue prioritisation	AI model-driven, reusable classifiers	Active Learning, requires project setup
Reviewer monitoring	Integrated with analytics dashboard	Dedicated Review Centre with granular metrics
Model reuse across matters	Native AI model library	Manual project recreation
Administration overhead	Lower for new matters	Higher, but more configurable

The practical difference in hosting cost is rarely about the per-gigabyte rate. It is about how much data reaches active hosting in the first place, and how efficiently reviewers are directed to the highest-value documents once it does.

The verdict: which platform for which matter?

Choose Reveal if

- The matter is analytically front-loaded. Regulatory and internal investigations, and matters where the client needs to understand exposure before committing to a strategy, benefit from immediate analytical access.
- Your team handles recurring matter types. The reusable AI model library compounds in value across similar matters.
- You want AI-driven culling before review begins. A smaller population reaching active hosting directly reduces cost.
- The matter is time-critical. Faster first insight means faster strategy decisions and a shorter overall matter.

Choose Relativity if

- The matter is large and operationally complex. Documented scale limits, queue controls and Review Centre monitoring provide operational precision that is hard to match.

- Defensibility of the review methodology is paramount. The audit infrastructure is court-tested over many years.
- Your firm already runs Relativity. Existing workflows, trained administrators and integrations represent investment worth preserving where the platform is fit for purpose.
- You need granular reviewer throughput data. Fifteen-minute increment reporting and queue-level monitoring give managers real-time cost control.

Where it genuinely does not matter

For mid-size matters with a defined review population, competent administration on either platform and a standard relevance-plus-privilege review structure, both platforms will produce comparable outcomes. The choice should then be driven by existing firm investment, administrator familiarity and managed service provider capability rather than platform features.

A NOTE ON MANAGED SERVICES

The platform is only part of the decision. How it is administered, configured and monitored matters as much as the underlying technology. A well-run Reveal environment will outperform a poorly configured Relativity environment, and the reverse is equally true. If you are evaluating platforms as part of a managed service arrangement, the provider's methodology and experience count as much as the platform specification.

Understanding the scoping and early data assessment decisions that precede platform selection is often where the most significant savings are made, before a single document reaches either environment. See our guide on early data assessment before disclosure costs escalate.

Sources

All sources were verified as of September 2026.

- Relativity Analytics overview, dtSearch, conceptual search and analytics suite capabilities in RelativityOne.
- Active Learning performance baselines (Server 2026), published scale limits, index build times and processing rates.
- Review Centre: monitoring a project, queue refresh logic, coverage mode behaviour and throughput reporting intervals.
- How AI agents are moving eDiscovery work upstream: from scoping to strategy, Reveal Data, September 2026.
- eDiscovery AI: agentic review reshapes legal teams, Reveal Data, July 2026.
- Reveal introduces Reveal AI: agentic eDiscovery from preservation to case development, Reveal Data, August 2026.

Frequently asked questions

What is the main difference between Relativity and Reveal?

Reveal front-loads analytics, running automated AI operations on ingestion so visualisations, clusters and plain-language querying are available on day one. Relativity treats analytics and Active Learning as configured workflows, which gives experienced administrators more control but delays first insight until the workspace is set up.

Which platform is better for early case assessment?

Reveal generally delivers analytical value earlier, with generative and agentic querying that returns sourced answers rather than document lists. Relativity offers a more auditable ECA record, with every saved search, filter and analytical step logged for export. Choose Reveal for speed of insight and Relativity where the methodology itself may be scrutinised.

How many documents can Relativity Active Learning handle?

Relativity's published baselines recommend a maximum of 9 million documents in a classification index, 1 million coded documents and 150 concurrent reviewers per Active Learning project. Larger populations are handled by slicing the data into parallel projects, which adds configuration and validation complexity.

Does Reveal publish comparable scale benchmarks?

No. Reveal does not publish performance baselines in the same granular format. It states support for matters of any size and processing of over 900 file types, with classifiers deployable at ingestion rather than after the review population is defined.

Which platform gives lower active hosting costs?

Both charge on a per-gigabyte active hosting basis, so the cost difference comes from how much data reaches review. Reveal's automatic pre-review culling tends to produce smaller populations; Relativity's configured filters tend to produce more precisely defined ones. Execution during scoping and ECA matters more than the headline rate.

How do the two compare on review management?

Relativity's Review Centre provides granular operational control, with saved search queues refreshing every 15 minutes during active coding, prioritised queues refreshing at 20% coded or 8 hours, and review speed reporting in 15-minute increments. Reveal relies on AI models to prioritise the queue automatically, reducing manual queue management.

When does the choice between them not matter?

For mid-size matters with a defined review population, competent administration and a standard relevance-plus-privilege structure, both platforms produce comparable outcomes. The decision should then rest on existing firm investment, administrator familiarity and the capability of the managed service provider.

Related guides

Early data assessment before disclosure costs escalate - How scoping decisions upstream of platform selection shape the overall cost of disclosure.

(e-discovery.uk/library/early-data-assessment-before-disclosure-costs-escalate)

How to scope an eDiscovery collection exercise - Defining custodians, date ranges and sources before data reaches a hosted review platform.

(e-discovery.uk/knowledge/how-to-scope-an-ediscovery-collection-exercise)

Preparing large data sets for legal review - Culling, de-duplication and processing choices that decide how much data you pay to host.

(e-discovery.uk/knowledge/preparing-large-data-sets-for-legal-review)

TAR and AI in UK eDiscovery - How technology assisted review and AI prioritisation are used and defended in this jurisdiction. (e-discovery.uk/knowledge/tar-and-ai-in-uk-ediscovery)

eDiscovery cost in the UK - Indicative cost drivers across collection, processing, hosting and review. (e-discovery.uk/ediscovery-cost-uk)