

WHITEPAPER

Legacy Classification is Losing the AI Race. Here's How to Win It.

This paper is for security leaders and data security practitioners evaluating DSPM solutions, especially teams that have watched data classification underdeliver for years. It is also for leaders who understand that AI does not just create more data risk. It multiplies the reach of data that was never found and data that was misclassified because it was never fully understood. The AI tools your organization deploys to accelerate work are, by design, seeking out data to learn from. Classification governs what they can access and what they cannot.

Open the
classification
case file →



Table of Contents

1. The Scenario	03

2. The Legacy Gap	04

3. What Cyera's Classification Actually Delivers	05

4. How Cyera Classifies Data	08

5. Deep Dive #1: Finding What You Didn't Know to Look For	12

6. Deep Dive #2: Moving From Labels to Business Decisions	14

7. Risk Evaluation in Action	17

8. Performance at Enterprise Scale	18

9. Closing the Case	19

10. Case Closed	20

Glossary	21

Interactive Menu

ToC

- 1
- 2
- 3
- 4
- 5
- 6
- 7
- 8
- 9
- 10
- Glossary



1. The Scenario

An Internal Label, a Codename, and a Hidden Exposure – A Classification Case File

There is a file sitting in one of your company's cloud storage buckets right now. It is a PowerPoint deck with a codename: **Project Orca**. It carries an internal Confidential label that someone manually applied. And it contains your complete mergers and acquisitions (M&A) strategy, including acquisition targets, financials, negotiation notes, and valuation models.

The file is currently accessible to a finance contractor provisioned six months ago and has been shared with outside legal counsel for an unrelated matter. Because the file lives in cloud storage and your organization's AI agent indexes it, it is now within that agent's reach as well. That agent relies on an enterprise indexing layer to retrieve and summarize relevant files, which means everything in that folder can become material the agent retrieves, references, or surfaces in response to an employee's question, with no deliberate decision to share it further.

Now here's the critical distinction: a legacy classification tool does not “understand” that deck is confidential just because a human marked it so.

At best, it can repeat the tag it sees. It has no idea what **Project Orca** means. It does not interpret M&A intent. It does not connect that codename to a live acquisition, deal timing, or material business impact. That gap between a superficial label and the meaning of the content is the problem.

In practice, conventional classification technology would likely overlook the PowerPoint as a critical risk — not because no one ever typed the word Confidential, but because legacy methods cannot reliably **1. detect the sensitive business reality inside the deck (see tables below)** and **2. translate that finding into action (see tables below)**. A tool built on pattern matching cannot interpret a project codename. It cannot infer that **Project Orca** represents a live acquisition that, if disclosed prematurely, creates deal risk, regulatory exposure, and material business impact. It sees, at most, a label someone applied months ago, and then it moves on.

This failure plays out in two distinct ways.



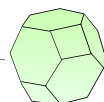
1. Detection failure.

Regex, regular expression pattern matching, will not reliably surface M&A intent. A file named **Project Orca** contains no structured data pattern to match: no Social Security Numbers, no credit card formats, no predictable string of characters that maps to a classification rule. And even when a human has applied a Confidential marking, legacy classification still lacks the business context to recognize why it is confidential and what kind of confidentiality it represents (M&A strategy vs. HR vs. customer data). Generic LLM models will not identify it as sensitive M&A material without business-specific configuration, configuration that requires ongoing manual maintenance most security teams cannot sustain.



2. Action failure.

Even if the file is flagged as Confidential / High Risk, the classification tool's job ends there. It tags it with the same label and stops. It does not tell you that access has been granted to external parties. It does not surface the fact that an AI agent or GenAI tool can index the folder where the deck lives, and pull its contents into summaries and answers. It cannot explain what **Project Orca** represents, who has been accessing it, whether that access is appropriate, or what to do next.



Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

Without classification that goes beyond assigning sensitivity labels, here is where the investigation ends. A legacy DSPM tool scans the file, repeats the Confidential label, and moves on. The finance contractor retains access. The legal vendor retains access. The AI agent continues retrieving and summarizing content from the folder. Nothing triggers a review. No one knows what **Project Orca** represents. The case is never opened.

Data like the **Project Orca** deck does not live in isolation. Today's sensitive data spans structured databases, unstructured documents, presentations, contracts, machine-generated artifacts, and AI-generated content, spread across clouds, SaaS platforms, and endpoints. And unlike the data challenges of five years ago, the risk compounds faster:

every AI tool deployed across your organization is actively discovering and querying that data at a speed no human team can audit. A classification tool that stops at assigning a sensitivity label cannot keep up. To understand why, you have to look at how most classification works today and where legacy methods break down.

2. The Legacy Gap

Why Classification Keeps Failing

The problem with legacy approaches to classification is that they were not built for the environment security teams face today: distributed hybrid infrastructure, AI agents multiplying how data is accessed and shared across the organization, and security teams already operating at or beyond capacity. Regex remains valid for what it was designed to do: scanning for structured data with well-defined formats like credit card numbers, Social Security Numbers, or API keys in source code. For these tasks, regex is fast, deterministic, and auditable.

But it breaks entirely when meaning depends on narrative context, organizational terminology, or a project codename. It can only find what it was told to look for.

Generic AI-powered classification tools promised to close this gap. These tools can read meaning in unstructured text where regex fails. But without business-specific training and appropriate guardrails, LLM-based classification at enterprise scale introduces its own problems:



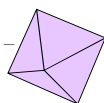
False positives at volume.

Excessive noise buries real risk. Pattern noise and hallucinations generate alert floods that crowd out genuine findings. A string that looks like a Social Security Number in an unrelated document triggers the same alert as an actual employee SSN in an HR database.



No autonomous adaptation.

Regex requires constant tuning as business language evolves. Generic LLMs require ongoing prompt engineering to stay current. Neither adapts autonomously, and neither was designed to recognize a project codename without being explicitly told what it means.



Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

The cost math doesn't work either. Running general-purpose LLMs across petabytes of enterprise data is expensive and produces inconsistent results. One-size-fits-all approaches just don't hold at these volumes. More fundamentally, neither generic LLM-based classification nor regex, alone or together, can deliver real understanding. Both return sensitivity labels. Neither knows what **Project Orca** means to your organization, who should have access, or whether the access that exists right now is a problem. And the data that matters most, like M&A planning materials, board resolutions, acquisition codenames, internal product strategies, has no regulatory marker. Legacy methods don't find these types of materials because they were never taught to look for them.

The goal was never a better label. It was always understanding what data means, and what should happen because of it.

3. What Cyera's Classification Actually Delivers

Actionable Data Intelligence

Cyera's approach to data classification starts from a different premise: classification is not the end state. It is the input to something more valuable.

Actionable Data Intelligence is the outcome you get when classification results are enriched with business context, identity and access permissions, and environmental signals. It answers not just what is this data? But what that data represents, who can reach it, whether that access is appropriate and, crucially, what to do next. The result is not a static label. It is a continuously updated, evidence-backed understanding of data that tells security teams what they need in order to act with confidence.

Actionable Data Intelligence means understanding: what data is, how it's used, who can access it, and why it matters, so teams can prioritize risk and take confident action.

The method Cyera uses to get to this state of understanding varies deliberately by data type. For structured data with well-defined patterns, deterministic regex with LLM validation is the right tool. For ambiguous or semi-structured content, a hybrid approach performs better. For unstructured documents like presentations, contracts, board materials, Cyera applies **Learned Classification**: semantic AI models that evaluate intent and full-document context in addition to token patterns.

Learned Classification is Cyera's approach to identifying sensitive data that has no predefined regulatory marker and no pattern to match.

Rather than searching for a keyword or format, it evaluates what a document represents: acquisition planning, due diligence preparation, valuation analysis. It applies to both structured and unstructured data, recognizing meaning within organizational context rather than relying on a known schema or format to signal sensitivity. This matters not just for files you know to protect, but for the ones you did not know to look for. **Learned Classification** surfaces those unknown unknowns, the sensitive data that would never appear on a predefined watch list but carries real business risk. That is what makes it the right method for a file named **Project Orca**.

Applying **Learned Classification** to our case scenario, the enrichment chain looks like this:

Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

Business context.

The document contains acquisition target financials, negotiation strategy, and valuation guidance. Exposure of this type of content represents material legal, regulatory, and reputational risk for the company.

Identity and access.

A finance contractor and an outside legal counsel firm currently have read access rights. Neither should. An AI agent can also retrieve and summarize content from the folder where the deck lives.

Exposure and environment.

The file sits in a broadly accessible cloud storage bucket, not a restricted executive workspace. It is co-located with routine finance files rather than isolated in a dedicated M&A repository. The output is not a label. It is an evidence-backed case for action: revoking access from the contractor and the legal vendor, restricting the AI agent's access path to sensitive directories, and confining access to a defined list of executives.

That is not a classification result. That is Actionable Data Intelligence.

What This Looks Like in Cyera's DSPM Platform

Actionable Data Intelligence surfaces in three distinct product experiences that correspond to the three questions security teams actually need to answer:

Topics Explorer

TOPIC	RECORDS	OBJECTS	ISSUES BY SEVERITY			
M&A planning: project Orca	148.2K	498	C 47	H 44	M 27	L 31
Product formulas	148.2K	1.8K	C 39	H 32	M 0	L 29
Customer contracts	148.2K	1.2K	C 37	H 36	M 23	L 22
Clinical trial data	148.2K	1.2K	C 0	H 34	M 21	L 22
Employee records	148.2K	1.2K	C 34	H 0	M 19	L 0

- **What does this data represent?** Classification results are mapped to business concepts using the **Topics** layer within the Cyera platform (more on this below). Instead of a list of file-level labels, you see "M&A Planning: **Project Orca**, 47 documents, 3 exposed to external parties." The business concept is the organizing unit. You act on that, not on individual findings.

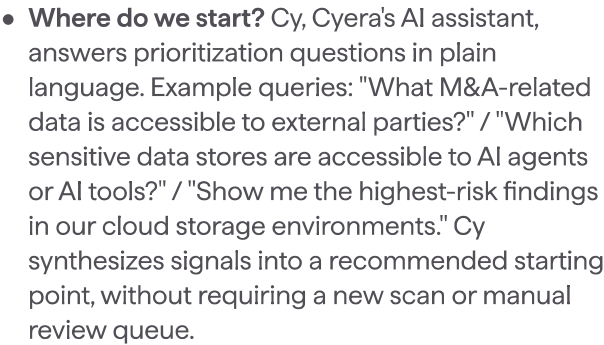
- **How serious is it, and why?** Every issue includes a severity score and an explanation of the signals that drove it: what content was found, who has access, what the exposure vector is, and what changed recently. You can trace the finding back to specific evidence, file path, identity, and access condition. And because each finding is tied to the context that produced it, the risk is actionable from the moment it surfaces, with no manual triage, no follow-up investigation required to understand what to do next.

Datastore Summary
AI

47 M&A planning documents discovered in a Sharepoint library belonging to ACME Inc., accessible to 3 external parties. For example, we found valuation models and negotiation notes from June of 2026 containing information about an upcoming acquisition

Interactive Menu





What Kind of Bad Outcomes Were Prevented by Stopping the AI Agent's Retrieval Path?

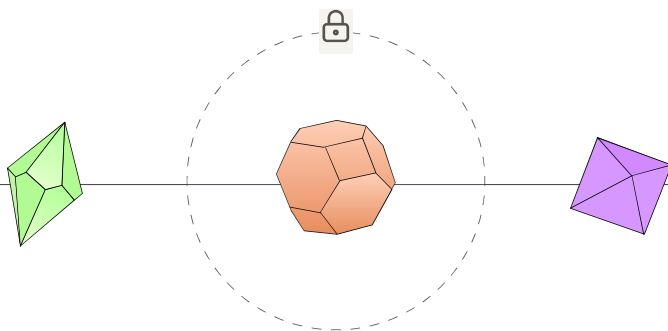
The AI agent in the **Project Orca** scenario is not a malicious actor. It is an authorized enterprise tool doing exactly what it was designed to do: retrieve internal content, summarize it, and enable team members to work faster. That is precisely what can make it risky.

If permissions are not strictly mirrored and enforced end to end, sensitive content can become discoverable through the agent's interface in ways that expand beyond the original file's intended access boundaries.



Stopping the retrieval path before it reached the **Project Orca** folder meant: no employee could inadvertently surface M&A strategy through a routine work query. No acquisition target name entered an agent workflow response. No negotiation position became accessible through a retrieval-augmented generation (RAG) pipeline. The risk extends beyond the original file permissions: when AI tooling can retrieve a document, that content can become discoverable through the tool's query interface, especially when permission mirroring and index-level auditing are not strictly enforced.

Sensitive data that is correctly labeled and access-controlled at the file level can still become broadly retrievable once AI tooling can access it, if the tool does not enforce the same access boundaries. Cyera's classification surfaces this relationship before the risk materializes, giving security teams the evidence to act before a query retrieves what should never have been findable.



4. Overview of How Cyera Classifies Data

Method Selection and the Seven Signals

A key engineering principle underpins Cyera’s entire architecture: **no single classification method works well across all data types**. Cyera selects and applies the method that best balances precision, speed, and cost for each specific dataset it encounters. This is not a compromise. It is what makes high-accuracy classification possible at enterprise scale without creating unsustainable inference costs or sacrificing throughput.

How Cyera routes each classification decision:

Data Type	Method Applied	Why
Structured data with defined patterns	Deterministic regex + LLM validation	Fast, auditable, and precise; LLM validation checks surrounding context before a flag is raised, materially reducing false positives.
Structured data with no predefined pattern	Semantic AI models (Learned Classification)	Surfaces sensitive data types with no regulatory marker or pattern to match; the unknown unknowns that regex alone would never find.
Ambiguous or semi-structured data	Hybrid: regex + contextual analysis	Establishes meaning where structure alone falls short, reducing the gap between what data contains and what it represents.
Structured and unstructured documents, contracts, presentations	Deterministic regex + LLM validation for defined elements; Semantic AI models (Learned Classification) for undefined elements	Evaluates intent and full-document context to identify what a document represents to the business, not just what it contains; critical for sensitive materials with no predefined pattern.



A critical detail in this architecture: when regex detects a potential match, Cyera applies LLM validation to check surrounding context before classifying. This step materially reduces the false positive rates that have historically plagued pattern-matching tools, where format-matching strings in irrelevant documents generate the same alerts as genuine sensitive data.

Interactive Menu

Cyera's classification architecture operates in three distinct stages that build on one another:

Stage	What It Does
Categorization	Establishes what data exists and where it lives; discovering data stores, file types, and basic attributes across the environment.
Classification	Determines what the data represents its type, content, and sensitivity; using the method appropriate to the data type (regex, hybrid, or Learned Classification).
Risk Evaluation	Correlates classification results with identity and access, environmental exposure, behavioral signals, and organizational context to produce a prioritized, evidence-backed risk picture.

Once classification is complete, the next step is determining sensitivity, and this is where most conventional tools stop short. Classification tells you what a piece of data is. Sensitivity determination tells you what it means: how much risk it carries, for whom, and in what context. Cyera uncovers this crucial information by correlating classification results across multiple dimensions simultaneously: what the content represents, where it lives, who has access to it, how it has been used, and critically, what it reveals in combination with other data nearby.

A financial projection is not inherently sensitive. An M&A target list is not inherently sensitive. But a financial projection in the same folder as an M&A target list and a deal timeline, accessed by someone outside the deal team, forms a toxic combination: actionable trading intelligence that creates material legal, regulatory, and reputational risk. Sensitivity is not a property of a single file. It is an attribute colored by everything around it.



Cyera builds a complete picture of risk by reading seven distinct layers of signal simultaneously for every piece of data it encounters. Several of these layers, including contextual associations, tenant context, behavioral signals, and environmental exposure, go beyond what conventional classification tools evaluate. Cyera also applies advanced semantic understanding to content, turning classification from pattern matching into business-aware interpretation. **(table below)**

Interactive Menu

Signal	What It Captures	What It Revealed in the Project Orca Scenario
Content	What is inside the data. Cyera uses semantic models, natural language processing (NLP), and named entity recognition to evaluate meaning, not just pattern matches. This is where many legacy classification approaches primarily focus, and often where they stop. For Cyera, it is the starting point.	Identified the file as M&A planning material, an acquisition strategy deck, based on document intent and full-context evaluation. Not keywords. Not a label. What the document represents to the business.
Contextual Associations	What the data is grouped with and surrounded by. It reveals risk that is not visible in any single file, because sensitive business activity like M&A is often expressed across a cluster of related documents, not one file.	The Project Orca deck had no obvious regulated-data patterns. But it was co-located with target company financials and a deal timeline in a shared folder and co-exposed via the same permissions to a finance contractor and external legal counsel. This is an M&A deal-room context that sharply increases exposure risk and elevates remediation priority.
Tenant Context	Tenant name, country, industry, and description. Tenant topic taxonomy configuration.	Identified users or groups with access patterns that might cause concern, such as the finance contractor and external legal vendor. Tenant context also clarifies organizational context, helping determine which risks to prioritize for mitigation.
Sensitivity	Building on contextual findings, correlates content with attributes such as access, usage, ownership, and relationships. Transforms raw findings into prioritized intelligence to inform security posture actions.	Correlated the M&A content, the contractor's access level, the external legal vendor's permissions, and the AI agent's retrieval path to produce a composite risk picture, surfacing the file as Critical Exposure and prioritizing it for immediate action above thousands of other findings.
Metadata	Attributes about the data, including column names, file schemas, and existing tags and labels. Often reveals sensitivity without reading a single byte of content.	The file name (Project Orca), folder path, and existing Confidential label together flagged the file for prioritized inspection and triage.
Behavioral	The breadth of access rights on the data. Who can reach it, and whether that scope is appropriate for the sensitivity of what it contains.	The M&A materials were accessible to a significantly larger group than the deal team. That access scope mismatch was a risk indicator that classification surfaced without needing to track individual access events.
Environmental	Encryption status, network exposure, AI model access, publicly accessible storage, and connections between data stores across the enterprise.	Confirmed the folder was being actively indexed by a deployed GenAI tool — creating retrieval risk for anyone querying the tool, regardless of their permissions on the original file.

Interactive Menu

[ToC](#)
[1](#)
[2](#)
[3](#)
[4](#)
[5](#)
[6](#)
[7](#)
[8](#)
[9](#)
[10](#)
[Glossary](#)

The Advantage of a Layered Approach

To see why this layered approach matters in practice, return to the M&A scenario. The CFO maintains a file called `targetfinancials.csv` containing proprietary financial data on the acquisition target. A conventional content-oriented classification tool tags it 'Financial Data,' assigns it a Medium Risk sensitivity label, and moves on. That's technically accurate. But it tells security teams almost nothing about risk. With Cyera's **Learned Classification**, once a file is flagged, a series of follow-on steps then paints a far more vivid and nuanced picture of the risk posed by that file (`targetfinancials.csv`) and the larger digital locale within which it is situated. These steps align with our Signal table above:

- **Add contextual associations:** the file lives in a shared folder accessible to the entire company, not just the finance team, and it is grouped with related deal artifacts such as valuation models and a deal timeline. In isolation, any one of these items might not trigger a critical alert. Together, the cluster matches a high-risk M&A deal-room profile because target information combined with valuation and timing can create actionable trading intelligence, the same general pattern behind well-known insider trading cases.
- **Add tenant context:** geographical and organizational signals, organizational taxonomy, and identity and access permissions all set a framework for understanding the content's specific level of sensitivity.
- **Add sensitivity scoring:** using multi-dimensional intelligence, Cyera correlates content with a range of factors to transform raw findings into prioritized intelligence.
- **Add metadata:** column headers name the target's founders, investors, equity stakes, and revenue estimates.
- **Add behavioral signals:** downloaded seventeen times in two weeks by people outside finance.
- **Add environmental signals:** the storage bucket also feeds downstream systems such as enterprise search, GenAI indexing, or applications that expand where the data can be surfaced.

Suddenly, the same file goes from 'Financial Data' to 'Critical Exposure: sensitive M&A records broadly accessible, with anomalous access patterns and environmental risk.' That's the difference between a label and intelligence.

"The Cyera Classification Engine offers a greater degree of quality than other competitors in the field. They don't rely on pure regex, nor do they try and manually enumerate the entire corpus of data elements that exist in the entire world (two fundamentally flawed approaches in other systems). Their ability to detect and classify business context (aka not just "name", but "customer name") and their ability to automatically detect and create new data element types opens up new avenues for compliance and governance workflows."

— Software Developer, Healthcare and Biotech, [Gartner Peer Insights](#)

Interactive Menu

ToC

1

2

3

4

5

6

7

8

9

10

Glossary

5. Deep Dive #1

Finding What You Didn't Know to Look For

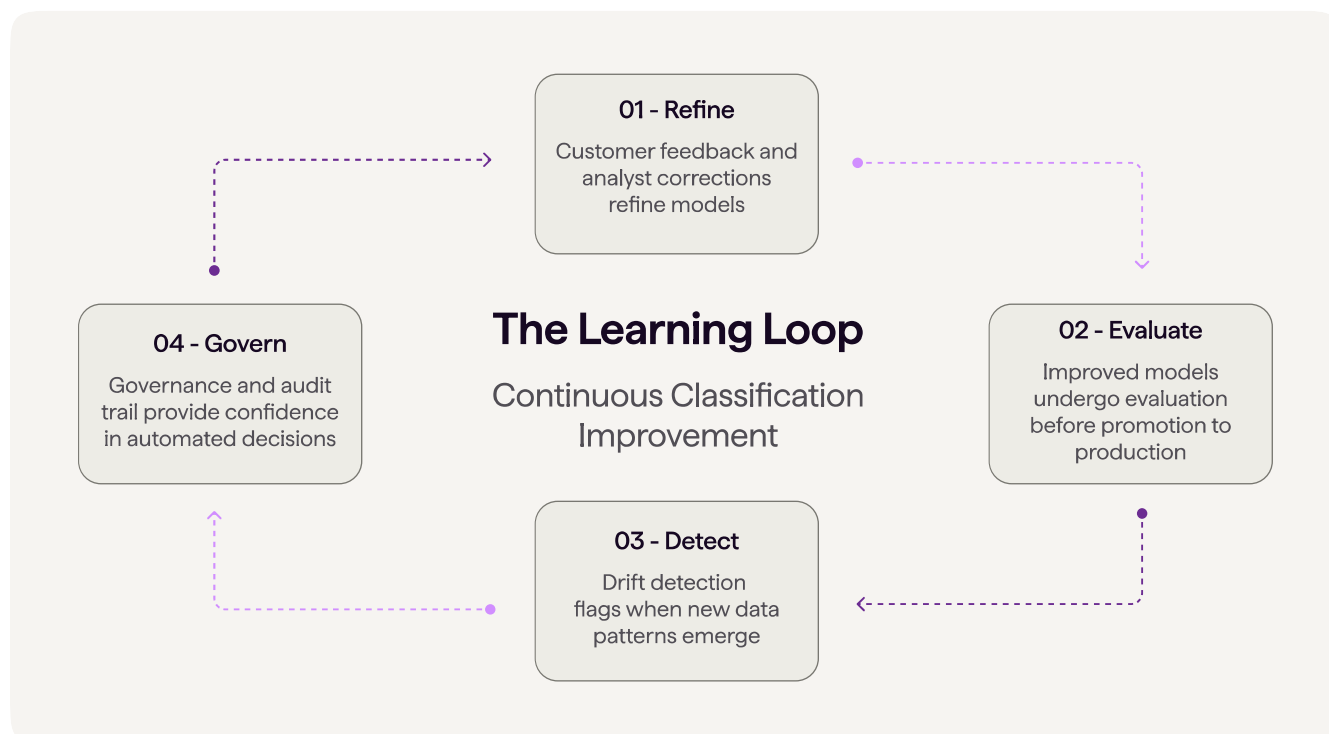
Every organization's sensitive data includes a category that does not appear in any regulatory framework or predefined taxonomy. Project codenames. Proprietary product formulas. Board resolutions. Deal structures. The data that matters most to the business often has no label anyone thought to create because no one anticipated it would need one.

Cyera's **Learned Classification** addresses this directly. Its AI-native classification models continuously improve their understanding of each customer's environment, learning the specific terminology in use, the naming conventions for sensitive projects, and the proprietary data types that carry significant business value without any regulatory definition. Over time, the platform gets better at finding exactly what matters to each individual organization.

Important: These models are isolated per customer. Learnings from one customer's data are never used to train models for other customers. Classification telemetry remains within each customer's environment.

This adaptive approach matters especially during periods of organizational change, M&A activity, product launches, regulatory shifts, or any event when entirely new categories of sensitive data become business-critical overnight. Cyera adapts without requiring a rescan. Once data has been processed, new classification parameters apply retroactively to existing data stores, not just forward to new content. The Cyera platform can also be customized on the fly to address shifting business processes and priorities.

The learning loop in practice:



Interactive Menu

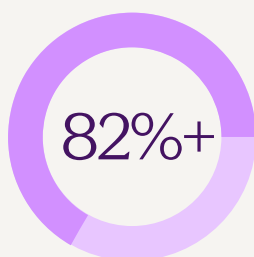
[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

A practical requirement for any DSPM program is keeping classification relevant as the environment changes. Cyera is built to continuously enrich classification results as signals evolve. Changes in permissions and access, ownership, behavioral access patterns, and exposure or environmental conditions (for example, encryption status, network exposure, or AI tool connectivity) update the risk picture over time without requiring new processing just to reflect those changes.

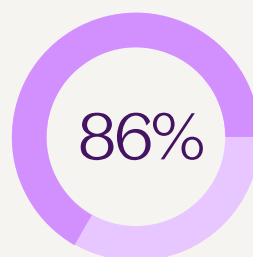
Cyera also supports retroactive updates to previously processed data when business context changes. For example, when teams define new **Topics**, update a business taxonomy, or apply a new prioritization lens, those definitions can be applied back to existing results so teams do not have to start over. In the same way, **Topics** enables users to customize their level of risk acceptance and their definitions around sensitivity as business focus evolves over time.

Net-new data still has to be processed. When new files appear, or when a data store is connected for the first time, Cyera must ingest and classify that content before the same enrichment signals and **Topics** logic can be applied.

Learned Classification produces the underlying signal layer and meaning. That architecture matters because traditional classification methods simply never find a large portion of sensitive, business-critical data. Two metrics illustrate the gap. In one documented deployment, 82% of the data classes Cyera identified had no equivalent in the customer's prior classification coverage: no pattern to match, no regulatory marker to trigger a flag. They were business-specific, and legacy methods were never taught to look for them. The scale of that gap holds across Cyera's broader customer base: based on aggregated, anonymized telemetry from 1,000+ customers, 86% of data in highly collaborative environments like Google Workspace, Microsoft 365, and Box is classified using **Learned Classification**. Without it, that data would be invisible, too unique in form and language for regex or generic LLMs to surface without customer-specific training that most teams cannot sustain.



of the data classes Cyera identified in a documented deployment had no equivalent in the customer's prior classification coverage: no pattern to match, no regulatory marker to trigger a flag.



of data in highly collaborative environments like Google Workspace, Microsoft 365, and Box is classified using Learned Classification.

This is based on aggregated, anonymized telemetry from 1,000+ customers

Interactive Menu

ToC

1

2

3

4

5

6

7

8

9

10

Glossary

6. Deep Dive #2

From Labels to Business Decisions: The Topics Layer

Learned Classification and **Topics** are not the same thing. **Learned Classification** determines what a piece of data represents, its intent, and meaning, when no pattern or label provides that answer. It operates at the content level. **Topics** works above that layer, using both the classification label and a direct view of the content to get full context into account. It does two distinct things. First, it translates classification results into the business concepts that security teams and leadership use for governance and prioritization: M&A Planning, Customer Contracts, IP Strategy, Board Meeting Materials.

Second, it enables teams to create custom classifications in plain language – defining new data types specific to their organization, including internal terminology and project codenames, that Cyera then applies consistently across the environment. A team can define a **Topic** called '**Project Orca**' and surface every related document without writing a single rule.

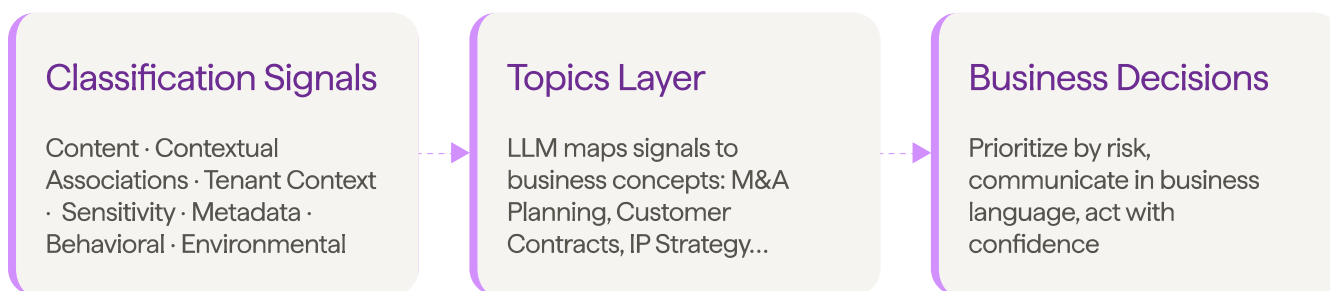
They can also create a custom classification for a data type that exists nowhere in a predefined taxonomy and have it applied retroactively across existing data stores.

The two capabilities are complementary: **Learned Classification** is how **Project Orca** is identified as M&A material; **Topics** is how it surfaces as 'M&A Planning: **Project Orca**' in the security team's queue – and how the team can build a custom classification around it that persists as the deal evolves.

Classification can succeed technically and still fail operationally. A team can accurately identify millions of sensitive data elements, PII, financial records, regulated health data, and still be unable to answer the question leadership needs: What is our highest security priority right now?

Security programs don't stall at discovery. They stall after classification, when teams have findings but still can't answer what matters most: what's urgent, and what gets fixed first.

Most teams know this feeling: the classification run finishes, the export lands in someone's inbox, and then there's a meeting where nobody can agree on which finding to start with. The volume isn't the problem. It's that findings written in technical label language don't translate cleanly into decisions a security team can act on.



Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

The distinction matters. **Topics** is not a label-grouping exercise. A document that mentions revenue figures, competitor names, and pricing details is not just "Financial Data." Within the right organizational context, it is "Competitive Intelligence," and that changes who should have access to it, what controls apply, and what the remediation priority is.

Teams define **Topics** in plain language, writing a natural-language description of the concept they want to find. This can include internal terminology, project codenames, and business-specific language. **Topics** then applies those definitions consistently across the environment, retroactively and without requiring a rescan. This is how the platform surfaces M&A Planning materials even when documents are not explicitly labeled "M&A" and how it identifies **Project Orca** files even before a security team knows to look for them.

Topics in Action: Project Orca Enters Due Diligence

As the M&A process matures, the legal team begins due diligence. The deal is still **Project Orca** internally, but the security team now needs to surface all related materials across a sprawling enterprise environment, fast. A security admin creates a new Topic with a plain-language prompt:

The screenshot displays a user interface for creating a Topic. At the top, a label reads "M&A Planning: Project Orca". Below this, the "Intent" section contains a text box with the prompt: "Complete mergers and acquisitions strategy. Includes acquisition targets, financials, negotiation notes, and valuation models". A downward arrow indicates the flow to the "Matching Logic" section. This section features two rows of logic rules. The first row has "Content Matching" (with a star icon) and "Strategic Business Plans" (with a document icon), separated by an "Or" button. The second row has "Purchase Price" (with a green square icon), "Buyer Data", and "Revenue Forecasts" (with a green square icon), with "Or" buttons separating the first two and the last two items.

Interactive Menu

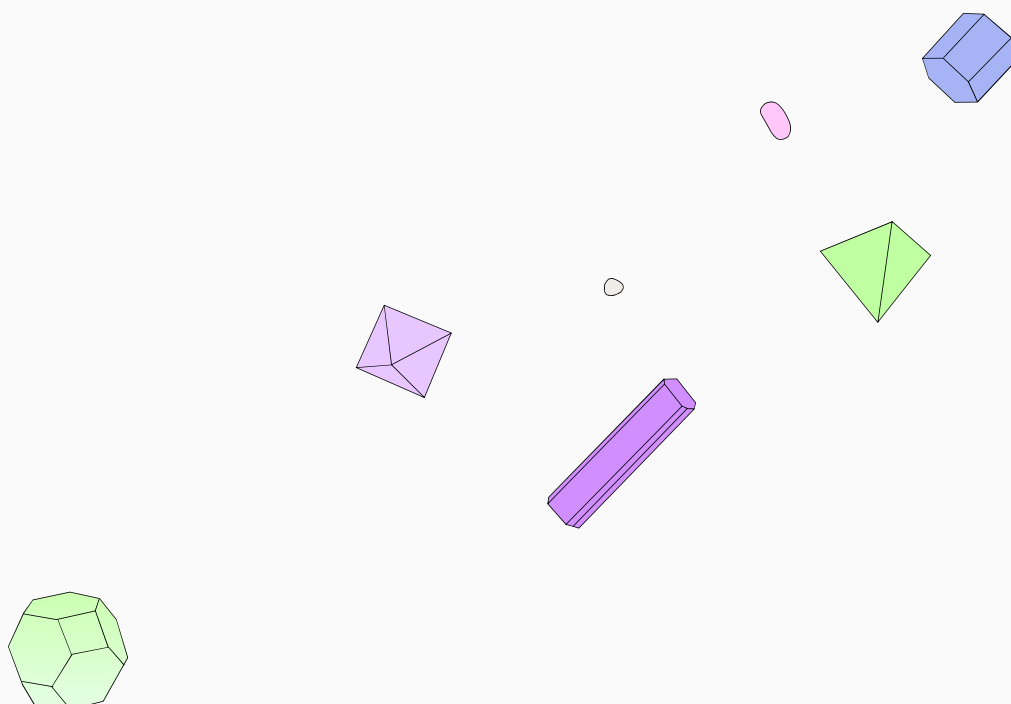
[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

Within hours, the platform surfaces all related materials across the enterprise, retroactively, without manual rule-building, without a rescan. The team does not have to sift through thousands of individual classification results. They see a business concept, M&A Planning: **Project Orca**, and see that a subset of those documents is publicly accessible. They revoke access before the merger announcement is public. They acted on the business concept, not on a label.

What **Topics** actually changes is where decisions get made. Instead of sorting through label counts, teams compare exposure across business concepts and act on the ones with real consequences right now: active M&A planning, clinical trial documentation, board materials. Policies get simpler too: anchoring controls to a stable business concept rather than a drifting list of labels means the policy doesn't break every time document types evolve. And when it comes time to explain risk to a leadership team, 'M&A planning documents are accessible to external parties' is a sentence that creates urgency. '37,000 documents labeled High Risk' usually just creates a follow-up question about what that means.



A Cyera customer with a live M&A transaction underway discovered that merger-related documents were accessible via public OneDrive links. The files had no special label triggering review. **Learned Classification** and **Topics** identified them as M&A Planning materials. The platform automatically surfaced the project codename across the environment, without any manual rule creation, without a rescan. The CISO revoked access and closed the exposure before the public merger announcement.



Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

7. From Topics to Priorities

Risk Evaluation in Action

Surfacing every document related to **Project Orca** is only half the job. The other half is telling the team which risks to act on first, and why. Once **Topics** identifies the M&A Planning: Project Orca bucket, Cyera doesn't return a flat list of findings. It correlates every document against the full signal set (content, contextual associations, tenant context, sensitivity, metadata, behavioral patterns, and environmental exposure) and produces a ranked risk view. The team sees the same business concept, but with a severity gradient that tells them exactly where to start.

Here is how that prioritization played out for **Project Orca**:

Document	Risk Signals	Severity	Action
Project Orca acquisition strategy deck (full)	Full deal financials, target valuations, negotiation positions. Accessible to finance contractor + external legal vendor. AI agent retrieval path active. 17+ downloads by out-of-scope parties.	Critical	Immediate: revoke contractor and vendor access, restrict agent retrieval path.
Due diligence checklist and letter of intent (LOI) drafts	M&A-related content, incomplete deal information. Accessible to an external legal vendor. No agent retrieval confirmed.	High	Within 24 hours: scope external access to named individuals only.
Finance team correspondence referencing Project Orca	References project codename but contains no deal financials. Internal access only, no external sharing detected.	Medium	Monitor; restrict codename access to deal team members.
Routine finance files co-located in the same folder	No M&A content, but physical proximity to sensitive materials and the same access group means exposure risk if folder permissions are ever modified.	Low	Flag for folder restructuring; separate M&A materials from routine content.

Most of these documents carry the same internal Confidential label and would receive a similar sensitivity label from a legacy tool. Underneath those labels, the risk could not be more different. Without the full signal context, a security team working from labels alone would have no basis for deciding which to remediate first, or for explaining that decision to leadership. With it, the answer is obvious before anyone opens a ticket.

This is the moment where classification stops being a finding and starts being an answer. Not 'these documents are sensitive.' But 'this document, with this access, at this exposure level, is your first call this morning.'

Interactive Menu

ToC

1

2

3

4

5

6

7

8

9

10

Glossary

8. Performance at Enterprise Scale

The capabilities described in this paper are meaningful only if the underlying tool performs at the scale modern enterprises require. Technical credibility demands evidence, not feature lists. Cyera operates across SaaS, IaaS, DBaaS, and on-premises environments. Coverage spans structured and unstructured data alike, from database tables to documents, images, presentations, and machine-generated artifacts. Unlike traditional tools that require manual reconfiguration as environments change, Cyera's understanding of your security posture persists and improves without manual tuning.

Cy, Cyera's AI assistant, surfaces these findings as actionable intelligence, explaining what data is sensitive, why it matters, and what to remediate first, in plain language, without requiring teams to re-initiate scans or rebuild taxonomies.

"We found with Cyera we could get full coverage across our entire ecosystem because the data classification system was second to none. The agent-less implementation allowed us to scale very quickly. We moved from purchase to implementation to full operationalization within months."

— Pete Chronis, CISO, [Paramount](#)

95%+ Precision

Maintained at enterprise scale across all data types

86%+ Unique Finds

A large percentage of data classes found would be missed by traditional tools

28.5TB In 17 Hours

Classification throughput in a single customer deployment

\$227K Saved Annually

One enterprise- scale Cyera customer saved over a quarter million dollars by reducing manual classification labor and eliminating manual data review

Interactive Menu

ToC

1

2

3

4

5

6

7

8

9

10

Glossary

9. Closing the Case

How the Team Remediated Project Orca Risks

Within hours of the risk evaluation completing, the security team had a clear picture: three categories of exposure, ranked by severity, with the evidence already assembled. They didn't need to investigate further. They needed to act.

- ✓ **First priority — the acquisition deck itself.** Access was revoked from the finance contractor and the external legal vendor. The AI agent's retrieval path to the **Project Orca** folder was restricted. Both actions were completed before the end of the business day. No manual triage. No ticket queue. Cyera routed the issue to the security team lead with the evidence chain already populated: which identities had access, for how long, and what they'd accessed.
- ✓ **Second priority — due diligence materials.** External access to LOI drafts and the due diligence checklist was scoped from open folder access to named individuals on the legal team. Two sharing links were revoked. One document was moved to a restricted M&A repository with explicit access controls.
- ✓ **Third priority — codename containment.** Internal correspondence referencing **Project Orca** was flagged for monitoring. The security team established a **Topics** alert: any new document mentioning the codename in a non-restricted location would trigger an immediate notification. The folder structure itself was reorganized, with M&A materials separated from routine finance files, so that future co-location risk was eliminated.

When the merger was announced publicly three months later, every sensitive document had been locked down. There was no scramble, no incident response, no leak review. The deal team hadn't noticed anything had changed, which was exactly the point.

None of it required writing a new rule. The security team didn't need to know in advance that **Project Orca** was sensitive. They didn't need to anticipate the codename, configure a regex pattern, or manually review every file in the finance folder. The **Learned Classification** understood what the data meant. The **Topics** layer made it findable. The risk evaluation told them where to start taking action. And the workflow told them exactly what to do.

The case closed because the intelligence was already there.
The team just had to act on it.

Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)

10. Case Closed

Project Orca did not become a headline until the deal was officially announced. By the time the deal progressed to due diligence, security already knew where every related document lived, who could reach it, and what to protect. None of that happened because someone wrote a better regex rule. It happened because Cyera's classification engine understood what the data meant and produced data intelligence that made action obvious. This is the version of AI governance that does not appear on vendor slides: not a policy, not a framework, but a continuous, evidence-backed understanding of what data your AI can reach.

The organizations getting this right are not just finding sensitive data faster. They are operating with a clearer picture of what they have, what is exposed, and what to do about it. In a world where the biggest risk is not that AI will generate sensitive data it already has, the question is whether your classification tool can keep up.

A Better Way Forward

What this paper established: Legacy classification fails in two specific ways. It cannot detect sensitive data expressed through business intent and organizational language: a project codename, a deal structure, proprietary negotiation terms. Legacy classification fails here because those signals have no pattern to match. And even when these legacy tools find something, they stop at a label. They do not know who has access, whether that access is appropriate, what downstream tools are ingesting the data, or what the business consequence of exposure actually is.

The question to take into your next vendor evaluation is not whether a tool can classify data. Every DSPM solution classifies data. The question is whether it can tell you what that data means to your business and produce the evidence needed to act on that understanding without manual interpretation, without a rescan, and without leaving your highest-value data invisible because no one thought to write a rule for it.

"The result is we will have all the visibility into the data that we need to make key decisions that result in reduced security and compliance risk, lower cost of operations, and improved data management capabilities."

— IT Director, Manufacturing, [Gartner Peer Insights](#)

Because in a world where AI learns from your data, what you understand about your data defines your security posture.

Ready to see Cyera's Classification in action?

Schedule your private demo to see the Cyera platform in action, and learn more about how our classification capabilities can give you the data clarity needed for safe AI enablement.

Book a demo →

Interactive Menu

ToC

1

2

3

4

5

6

7

8

9

10

Glossary

Glossary



Learned Classification

Learned Classification is Cyera's AI-native approach to identifying sensitive data that has no predefined regulatory marker and no pattern to match. Rather than searching for a known format or keyword, it evaluates what a document represents — acquisition planning, due diligence preparation, IP strategy, board materials — within the full context of the organization.

This matters because the data that carries the most business risk often has no label anyone thought to create. Learned Classification surfaces those unknown unknowns: the sensitive data that would never appear on a predefined watch list but carries real consequences if exposed.

Learned Classification models are isolated per customer. Learnings from one customer's environment are never used to train models for another. And because Learned Classification continuously improves its understanding of each customer's environment, it gets better at finding exactly what matters to that organization over time — without requiring manual rule-writing or ongoing configuration.

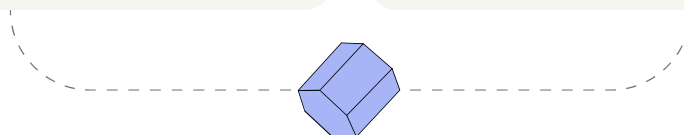


Topics

Topics is an LLM-powered business context layer that uses classification results as signals, then applies intent and full document context to map data to the business concepts that matter.

The result is a curated set of meaningful business concepts — M&A planning documents, product formulas, customer contracts, clinical trial data, financial planning and forecasting — that security teams can prioritize, govern, and remediate.

Topics is also customizable by design. Instead of forcing every organization into a fixed set of categories, Topics lets customers bring their business taxonomy to life inside Cyera. Teams can define the business concepts that matter to them in their own language, and Topics adapts as priorities change. When a new priority appears, teams can define it in plain language and make it a recognized concept across the platform.



About Cyera

Cyera is the Data and AI Security Platform built for the age of agents. Named a Leader in The Forrester Wave™: Sensitive Data Discovery & Classification Solutions, Q2 2026, Cyera helps enterprises like Paramount, Chipotle, and Valvoline control exactly what data their AI can reach—and govern what happens next. The platform secures data at rest, in motion, and in use.

Learn more at www.cyera.com.

Interactive Menu

[ToC](#)[1](#)[2](#)[3](#)[4](#)[5](#)[6](#)[7](#)[8](#)[9](#)[10](#)[Glossary](#)