

Explainable Scoring Benchmark:

What Enterprise RevOps Teams Should Demand From Account Intelligence Platforms

A transparency benchmark across the account-intelligence category, and why the gap between intent data adoption and intent data ROI traces back to scoring you can't audit

Published August 2026

The Core Finding

Across the account-intelligence platforms enterprise RevOps teams evaluate today, scoring transparency is the exception, not the rule. Of the five major platforms benchmarked for this report, only one publishes a fully auditable scoring methodology that a RevOps team can see into and adjust; the rest deliver either partial explanations or none at all. This matters because the industry's own data shows a widening gap between how much intent data organizations collect and how much of it they can actually trust. When 87% of organizations report their intent signals are unreliable or inflated, and only 26% of those signals convert into qualified opportunities, the problem isn't a shortage of signals. It's that most teams cannot see, trust, or act on the scores those signals produce.

87% vs. 26%

87% of B2B organizations report their intent data produces unreliable or inflated signals, but only **26%** of those signals convert into qualified opportunities, according to DemandScience's 2026 State of Performance Marketing report.

Context: What the Market Shows

Intent and account-scoring platforms have become close to mandatory infrastructure for enterprise revenue teams. Nearly every RevOps and sales operations function now runs some combination of fit scoring, intent scoring, and engagement scoring to prioritize which accounts get outbound attention, which get marketing spend, and which get escalated to sales leadership.

The adoption curve is not the problem. The trust curve is. Teams aren't short on signals, they're short on signals they can act on with confidence.

Talk to enough RevOps leaders and the same complaint surfaces regardless of which platform they run: reps don't trust the scores. Not because the underlying signals are wrong, but because nobody can explain, at the record level, why an account scored an 82 instead of a 41. When a rep can't see the "why," the score becomes a suggestion they can ignore rather than a prioritization tool they act on.

KEY TAKEAWAY

The intent data category has largely solved for signal volume. It has not solved for signal legibility. That distinction is where the adoption-to-ROI gap lives.

This is not a fringe concern. It shows up in vendor-published documentation, in support FAQs, and in the professional services line items that platforms charge when a customer wants to understand or change how their own model works. A scoring engine that requires a paid services engagement just to explain itself is, functionally, a black box with a price tag attached.

The Evidence

To quantify how the category handles transparency, this report benchmarks scoring explainability across five platforms commonly evaluated by enterprise RevOps teams: HG Insights, 6sense, Demandbase, ZoomInfo, and Bombora.

Each platform was rated Full, Partial, or None based on whether it discloses which signals drove a given account's score, whether operators can adjust the weighting of those signals without a professional services engagement, and whether the explanation is written into the CRM record itself or trapped in a separate UI.

PLATFORM	SCORING TRANSPARENCY RATING	WHAT THIS MEANS IN PRACTICE
HG Insights	Full	Signal-level explanation and operator-adjustable weighting are delivered as structured CRM fields, not a separate report a rep has to go find.
Demandbase	Partial	Model-level explanations exist and are relatively sophisticated, but they are read-only. Users can select training inputs, not adjust scoring weights, and explanations may not sync to the CRM fields reps actually work from.
ZoomInfo	Partial	The AI-driven fit score itself is not explainable. A separate, manually configured rules engine is transparent by design, but it's a different system running alongside the black-box model, not a window into it.
Bombora	Partial	Topic-level intent scores are visible, but the signal provenance behind them stays hidden. Bombora offers no self-service weight adjustment and no native fit scoring; it's a signal provider layered into other platforms' models, not a standalone scoring engine.
6sense	None	The platform's own support documentation states that fit scoring cannot be adjusted by the customer. Changing how a model weighs signals requires a separate paid services engagement.

The pattern that emerges is not "one good vendor and four bad ones." It's a spectrum of how much control an organization retains over a system it depends on daily. Only one of five platforms evaluated lets an operator see every input to a score and change how those inputs are weighted without engineering or paid services help. Three platforms offer partial visibility that stops short of operator control.

1 of 5PLATFORMS RATE FULL ON
SCORING TRANSPARENCY**3 of 5**RATE PARTIAL: SIGNALS
EXPOSED, WEIGHTS ARE NOT**~\$12K**TYPICAL SERVICES COST TO
REQUEST A MODEL CHANGE ON
THE LEAST TRANSPARENT
PLATFORM

Why does this specific failure mode, no operator control, matter more than raw signal volume? Because the RevOps function's job isn't to collect signals. It's to convert signals into actions a sales team will actually take. A score a rep cannot interrogate gets treated the same way as a recommendation from a system they don't trust: acknowledged, then ignored. Multiple independent analyses of AI-driven account fit scores have documented this exact failure pattern, describing black-box outputs as "nearly impossible to disprove," which cuts both ways. If a rep can't disprove a bad score, they also can't validate a good one, and the entire model loses credibility by default.

"When your ABM platform tells reps 'this account is hot' but can't show why, the data gets ignored."

Independent industry analysis of black-box account scoring adoption

The mechanism behind the one "Full" rating in this benchmark is worth unpacking briefly, because it illustrates what operator-level transparency actually requires architecturally. In HG Insights' implementation, every scored account pushes four structured fields into the CRM record: a plain-language segment label, a numeric 0-100 score, a visual indicator, and a signal list that operators configure to explain each score's key positive and negative factors to reps. Segment thresholds and model weighting can be adjusted directly by an operator in the platform's model-building tool, without an engineering ticket or a services invoice, and scores automatically normalize into the 0-100 range reps see in the CRM. That is the structural difference between a platform that explains itself and one that merely reports a number.

KEY TAKEAWAY

Explainability is not a UI feature. It is an architectural decision made at the point a scoring model is built. Platforms that bolt on explanations after the fact tend to land in "Partial": visible but not adjustable, or adjustable but not visible in the CRM where reps actually work.

A Secondary Signal: Contact-Level Data Completeness

Scoring transparency is the primary finding of this benchmark, but it sits inside a broader question RevOps teams are asking about platform depth: once an account is correctly scored, can the platform also deliver a complete, usable profile of the humans inside that account?

This is a supporting data point, not a separate finding of equal weight. Raw contact database size remains a category where longer-tenured providers hold a volume advantage; platforms built primarily around contact databases have had more time to accumulate verified phone and email records at scale. What is less understood is that contact enrichment architecture and contact volume are two different things, and a platform can rate strongly on one while trailing on the other.

Waterfall enrichment, which pulls the best available contact fields from multiple provider sources rather than relying on a single database, is increasingly the architecture of choice across the category precisely because no single vendor has complete contact coverage on its own. Evaluated on architectural maturity rather than raw record count, waterfall-based approaches rate on par with established category leaders, even when the underlying contact volume is still scaling toward parity.

The practical implication for buyers evaluating multiple platforms: don't conflate "how many contacts does this vendor claim to have" with "how good is the methodology for filling gaps when the primary source comes up short." Those are different questions, and the second one often matters more once a team is past initial account prioritization and into execution.

"Having contact-level data would be the holy grail. It would give us a complete profile for every campaign."

Amanda Mock, Senior Channel Marketing Manager, NiCE

That quote captures the practitioner reality well: contact completeness isn't a nice-to-have layered on top of good scoring, it's the thing that turns a correctly prioritized account into a campaign that can actually be executed against verified titles, functions, and direct numbers.

Implications: What This Means For You

If your team is currently running a platform rated Partial or None on this benchmark, the adoption-to-ROI gap identified in DemandScience's research is not an abstract industry statistic. It is very likely showing up in your own pipeline reviews as reps quietly reverting to gut instinct over model-driven prioritization, or as a RevOps team that spends more time defending the scoring model in QBRs than improving it.

The counterintuitive part is that this problem rarely gets flagged as a scoring problem. It gets misdiagnosed as a rep adoption problem, an enablement problem, or a data quality problem, and each of those gets its own remediation budget. Meanwhile the actual cause, a model nobody can see inside, goes unaddressed because "buy a new scoring model" is a bigger conversation than "run another enablement session."

KEY TAKEAWAY

If reps distrust your account scores, run the diagnostic before you run the training. Ask five reps to explain, in their own words, why their top three accounts scored the way they did. If they can't, the fix is architectural, not behavioral.

For teams currently mid-contract with a platform that rates Partial or None, the near-term move is not necessarily a full displacement. It's building an internal audit process that treats scoring transparency as a renewal criterion with the same weight as pricing or contact volume, so the decision gets made with eyes open rather than by default at auto-renewal.

What To Do About It

1. **Run a five-rep transparency audit before your next renewal cycle.** Pick five reps and five scored accounts each. Ask each rep to explain the score without looking anything up. If explanations are vague, generic, or absent, you have quantified evidence of a transparency gap you can bring into a vendor conversation or a budget request.
2. **Separate "signal volume" from "signal legibility" in your evaluation criteria.** Most RFP scorecards weight the number of signals a platform ingests. Add a line item scoring whether those signals are explained per record, in the CRM, without a services request. Treat it as a first-class criterion, not a footnote.
3. **Price out the cost of opacity, not just the license.** If a platform charges a separate professional services fee to explain or adjust its own model, that fee is a recurring tax on your ability to operate the tool. Model it into total cost of ownership rather than treating it as an occasional line item.
4. **Distinguish contact volume from contact architecture in vendor comparisons.** When comparing platforms on contact data, ask specifically how gaps get filled when a primary source misses a record, not just how many total contacts a vendor claims. A strong waterfall methodology can offset a volume disadvantage faster than raw database size alone suggests.
5. **Make explainability a named criterion in your next RFP, not an assumed feature.** Write it into the requirements document explicitly: per-record signal explanation, CRM-native delivery, and operator-adjustable weighting without paid services. Vendors that can't check all three boxes should have to say so in writing.

Methodology & Sources

This report synthesizes competitive functionality and messaging research on account scoring and data enrichment platforms conducted in Q2 2026, alongside third-party industry research on intent data adoption and ROI. Platform ratings (Full, Partial, None) reflect published vendor documentation, support center content, and independent third-party analysis as of the research date; vendor capabilities evolve, and ratings should be revalidated before use in a live evaluation. This analysis reflects market conditions as of mid-2026. Recommended refresh: quarterly, given the pace of change in AI-driven scoring features across the category.

1. DemandScience. "2026 State of Performance Marketing: Exposing the Marketing Data Mirage." December 2025 (n=750 senior B2B marketing leaders).
2. HG Insights Competitive Intelligence. "Account Scoring: Aggregated Competitive Functionality Report." May 2026.
3. HG Insights Competitive Intelligence. "Aggregated Competitive Intelligence: Data Enrichment." May 2026.
4. Vendor-published support documentation and FAQs (platform-specific scoring adjustability policies), as cited in the above reports.
5. NiCE customer interview, HG Insights customer case study compilation.

Ready to see how explainable scoring could change what your reps actually trust and act on? Contact HG Insights at hginsights.com or speak directly with your HG representative.

[Request a Demo](#)

ABOUT HG INSIGHTS

HG Insights delivers AI-powered Revenue Growth Intelligence (RGI) solutions that modernize GTM strategy and activation, enabling B2B companies to prioritize, target, engage, and convert the best opportunities faster. HG Insights' platform analytics and agents turn deep market, account, technology, spend, intent, and customer data into actionable insights and automated workflows that accelerate pipeline and enhance predictability. That's why 95% of Fortune 1000 B2B tech companies and all major hyperscalers rely on HG Insights to grow revenue, boost efficiency, and improve retention. For more information, visit www.hginsights.com.