



GLASS LEWIS

The Architecture of Investor-Grade AI

Lessons From Climate Intelligence



Arik Brutian, Senior Vice President, Artificial Intelligence & Data
Kaveh Beigi, Vice President, Digital & Content Strategy

September 2026

Contents

Executive Summary	3
Key Concepts in This Paper	4
Section 1: How Methodology Expands in an AI-Powered System	5
The Depth of the Methodology Layer	6
Building the Governing Layer	7
The Hybrid AI Ecosystem: Neuro-Symbolic by Design	8
The Data Foundation	9
Section 2: The Human-Centric AI Production Philosophy	10
Implementation Through Design: A Hybrid AI Architecture and Human Oversight	11
Output-Level Oversight: Human-in-the-Loop	12
Process-Level Oversight: Human-on-the-Loop	13
Architecture-Level Oversight: Human-in-Command	14
Section 3: Glass Lewis' Climate Intelligence	15
What Climate Intelligence Delivers	15
The Transition Assessment Framework	16
Inside the Human-Centric AI Production Architecture	17
The Three Levels in Climate Intelligence	18
What Stewardship Teams and Portfolio Managers Receive	19
The Five Questions, Answered	20
Section 4: The Standard Beyond Climate Intelligence	22
The Glass Lewis Approach	22
Notes and References	23
Glossary of Terms	24

Executive Summary

When a provider says its AI system has human oversight and methodology-driven outputs, what should that actually look like in operation, and how would an investment team recognize it?

This paper answers that question from inside a production system operating at scale today.

It starts with methodology. In an AI-powered governance research system, institutional methodology expands in function: it becomes the structured governing layer that determines what the AI is permitted to conclude, from what evidence, and by what reasoning path. How deeply that layer has been built is the most consequential architectural choice an AI governance research provider makes. It is the difference between outputs that reflect institutional expertise and outputs that reflect a general-purpose model's training data.

It then sets out the human oversight that governs that layer in production, operating at three levels — output, process, and architecture — each governing a different scope of the system. This distinction matters when evaluating a provider. When organizations describe human curation of AI, in the large majority of cases they mean oversight at the output level alone: human verification and checking of finished work. That is necessary. It is not sufficient for investor-grade research.

The paper then uses Glass Lewis' [Climate Intelligence](#), launched in April 2026, as a working example: what the product delivers, the methodology that governs it, how the production architecture operates, and how stewardship teams and portfolio managers use the outputs. Section 3 closes by answering the five evaluation questions posed earlier from *AI and the Fiduciary Test*.

This is the third paper in the Glass Lewis AI series. [Part one](#), *AI and the Fiduciary Test*,¹ set out the standard institutional investors should apply when evaluating any AI-enabled governance system: a methodology framework built by domain experts that governs the operations and applications of AI, human oversight embedded in the production process at multiple levels, and full source traceability from disclosure documents to final output. [Part two](#)² turned to the data that AI reasons over, setting out the standards investor-grade data must meet, especially when produced through AI: accuracy, consistency, completeness, traceability to its sources, validity, and timeliness. This paper shows what the standard looks like in production.

Key Concepts in This Paper



1. METHODOLOGY AS GOVERNING ARCHITECTURE.

In an AI-powered system, institutional methodology expands in function: it becomes the structured governing layer for the operations and applications of AI, and an amplifier that extends expertise across more outputs than analysts could cover alone.



2. THREE LEVELS, THREE SCOPES.

HCAI operates at the output, process, and architecture levels. Each governs a different scope of the system and performs a governance function the other two cannot.



3. A PERPETUAL DISCIPLINE.

Governance of AI-powered research is not a one-time build. As technologies, regulations, and disclosure regimes evolve, so do the architecture and the implementation work behind it. HCAI is an AI philosophy embedded by design into the architecture, regardless of how the underlying technology advances.

Section 1: How Methodology Expands in an AI-Powered System

In a traditional governance research workflow, methodology lives partly in the documents that codify it and partly in the heads of the analysts who apply it. Experienced analysts fill in the methodology's gaps with judgment, resolve corner cases by precedent, and recognize patterns the formal methodology never named. When AI enters that workflow, the inherited, intuitive portion cannot carry across. Everything must be made explicit, structured, and machine-readable for the AI system to apply. That requirement is what makes methodology's function expand when AI becomes part of the production process.

The function methodology retains is familiar: guiding how analysts think about a problem and how they design and oversee the work. The function it adds is new. Methodology becomes the governing architecture that the AI systems operate within. It also becomes an amplifier: institutional expertise extends across more outputs, more documents, and more companies than analysts could cover alone, while remaining in charge of the operations and applications of AI. Both functions operate at the same time.

The implementation work this expansion requires is qualitatively new. It is different from anything governance research has had to do before, and it calls for new skills within the methodology team: comfort with large-scale generative and reasoning models, fluency in how protocols and heuristics shape AI behavior, and judgment about which architectural decisions belong to methodology and which to engineering. This is not a one-time build, and it does not stand still.



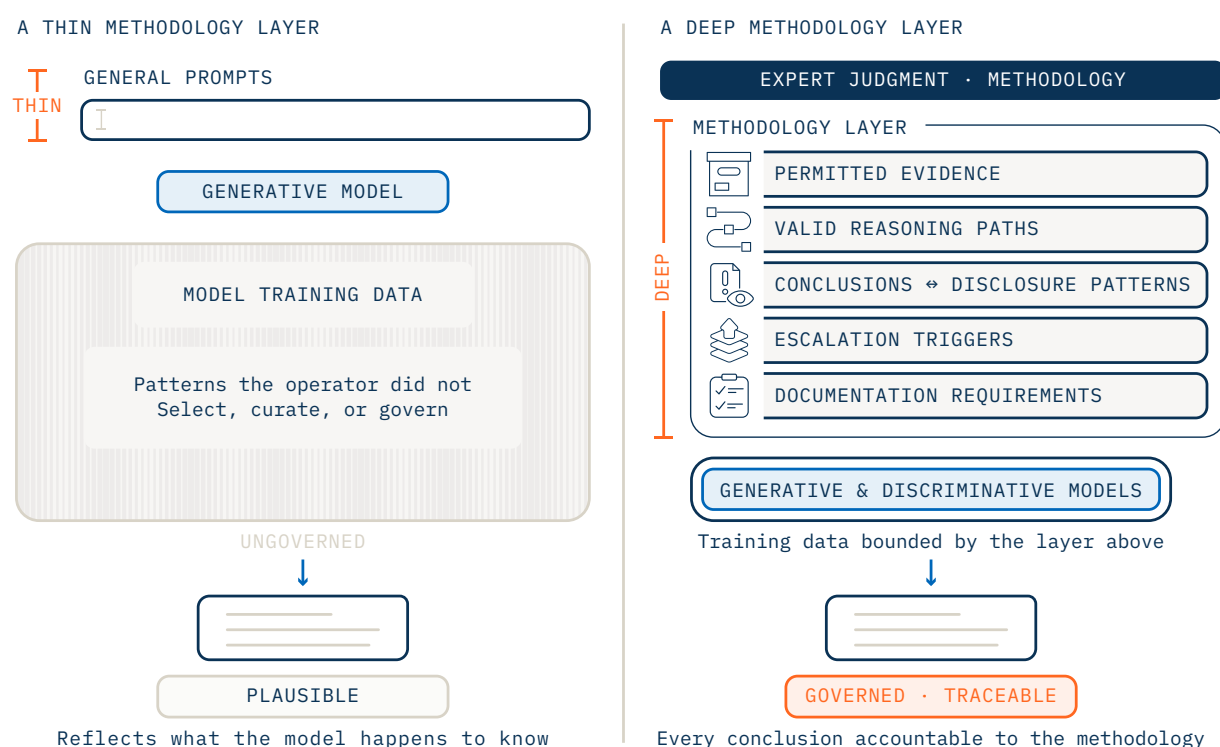
In a traditional workflow, methodology guides how analysts think. In an AI-powered system, methodology must explicate the analyst's implicit thinking, create new AI-specific reasoning protocols, and define what a valid AI output is. Everything the analyst applies intuitively must be made explicit and machine-readable. That expansion in function is the most consequential design decision in any AI governance research system.

The Depth of the Methodology Layer

Any organization building AI in governance research will have a methodology in some form. A major consideration is how deeply that methodology has been operationalized for the AI system it governs. In an AI context, what methodology is becomes even more important to be precise about. Methodology is three things at once: 1) the logic of research design, 2) a shared language for talking about research independently of any single subject, and 3) a set of evaluative standards for what counts as a sound conclusion.³

A system relying primarily on general prompts and the model's training data has a shallow methodology layer (Figure 1). Prompts cannot carry research design on their own. They are not a shared language for reasoning about research, and they capture evaluative standards only partially. The prompts shape the AI's focus, but the model's outputs remain shaped by patterns it learned from data its operators did not select, curate, or govern. The reasoning behind the output is plausible rather than governed. It reflects what the AI happens to know from training on internet-scale content, not what the institution's methodology requires.

Figure 1: Visual Comparison of Shallow vs. Deep Methodology Layers



A concrete example illustrates why this matters. In climate analysis, water risk and climate mitigation are often discussed together in general media and corporate communications. The associations are present throughout an AI model's training data. When a general-purpose AI is asked to assess a renewable energy deployment, it may surface water-risk considerations, because that is what the training data suggests is relevant. But water

risk does not apply to most renewable energy deployments. Identifying the false positive requires domain expertise, and pushing the AI away from the spurious association requires methodology rules that explicitly bound what the AI is permitted to conclude from a given evidence pattern.

A system built on a deep methodology layer is different. The layer translates institutional expertise into structured rules, heuristics, and processes the AI operates within. For each task the AI performs, it specifies what evidence is permitted, what reasoning paths are valid, what conclusions are supported by what disclosure patterns, what gets escalated to analyst review, and what documentation each output must produce (Figure 1 above). The AI reasons within a framework that experts defined. Its reasoning is bounded, traceable, and accountable to the methodology, which functions as a pre-defined body of expert judgment that governs the work before the AI ever runs.



Glass Lewis is not interested in the AI's opinion. Glass Lewis is interested in its reasoning. The methodology layer is what makes that distinction possible in production.

The difference between shallow and deep is the most consequential architectural choice any AI governance research provider makes. It determines whether the outputs reflect institutional expertise or only the patterns available in source documents and the model's general training.

Building the Governing Layer

The translation from institutional knowledge to an operational governing layer is the labor-intensive part of building this kind of architecture. It is an upfront investment of time and expertise. It begins with the methodologists and analysts who have spent years covering a governance domain. They specify the analytical framework: which questions matter, what evidence is material, what disclosure patterns warrant which inferences, where the judgment calls live, and what reasoning logic should bring questions, evidence, disclosure, and judgment together. They identify the corner cases accumulated across years of research and the principles that resolved them.

That framework is then translated into a structured, cohesive rule set the AI operates within: rules, heuristics, and processes that specify, for each task the AI performs, what sources it may use, what reasoning approach it should apply, what model is appropriate for the task, and what the output must contain. The translation requires close collaboration between methodologists, AI architects, research analysts, and data engineers. None can do it in isolation, because the depth of the methodology layer depends on understanding both what the research has to support and what the AI can and cannot reliably do. Data engineering, often less visible than methodology or AI work, is a substantial share of the

build: the infrastructure that moves documents through the pipeline reliably is foundational to everything that runs on top of it.

Two further points about the translation are worth noting. First, machine reasoning is not necessarily just an explication of human reasoning. Reasoning logic that works well for human analysts can fail with large language models, and reasoning logic that works well for the models can be unintuitive to humans. The implementation team has to discover what works through testing rather than assume from analogy. Second, the work is bounded by stable rules of engagement: methodology changes happen rarely and only when many companies are affected. Most adjustments in production are to the implementation layer, the heuristics and protocols that shape how methodology is applied, not to the methodology itself.

This collaboration is what makes the architecture hard to replicate. The governing layer reflects institutional knowledge accumulated over decades, translated with operational specificity for an AI production system. The time and expertise required to build it cannot be compressed or substituted by technology alone. And because the work is continuous, even a competitor that copied the current state would find themselves behind as soon as the next set of conditions emerges.

The Hybrid AI Ecosystem: Neuro-Symbolic by Design

In an AI-powered governance research system, the choice of which AI technology to apply at which stage is itself an act of methodology, with consequences for the defensibility of the final output. Technologically, Glass Lewis has implemented its human-centric AI ecosystem as a neuro-symbolic system,⁴ a design that keeps human agency, expert guidance, and contextual judgment in the leading role.

Generative AI systems are, on their own, heavily black-boxed: to a human observer, the machine's reasoning is not transparent. A neuro-symbolic design augments the generative models with a rule-based component, made up of formal logic, expert heuristics, and explicit constraints. That rule-based component acts as the executive function of the system. The useful way to picture it is the human brain: the generative model is the part that recognizes patterns and reacts quickly, while the rule-based layer acts like the prefrontal cortex, supplying the logic, the filtering, and the control that keep the system on track.

Within that design, each technology is used where it is genuinely strong. Large language models are powerful for pattern detection, synthesis, and structured output generation, and, governed by the rule and heuristics layer, they are well-suited to the inferential work that comes later in a production pipeline. Specialized classifiers handle document categorization reliably across governance-specific document types. Extraction tools are built for textual, tabular and numerical content. Rule-based and symbolic systems produce deterministic validation checks and oversee the whole life cycle. Matching each tool to the stage it serves best, and matching the oversight approach to each tool's strengths and limits, is what makes the hybrid architecture more than the sum of its parts.

For investment professionals evaluating AI-enabled research, the choices a provider has made about which tools to use, where, and why, reveal how deeply the system was designed for the accountability context it operates in.

The Data Foundation

The methodology layer governs the AI's reasoning, but the quality of that reasoning also depends on the governed data it reasons over. A methodology layer operating on unreliable, inconsistent, or incomplete data produces governed conclusions from ungoverned evidence. The data architecture, the standards, quality controls, and governance that ensure the underlying data is investor-grade, is the foundation the methodology layer sits on.

That data architecture is the subject of part two of this series.⁵ For the purposes of this paper, the point is that the methodology layer and the data architecture are complementary: one governs the reasoning, the other governs the evidence. Both must be investor-grade for the outputs to also be investor-grade.

Section 2: The Human-Centric AI Production Philosophy

When an investment professional defends a voting decision to a client, a board, or a regulator, what they are defending is the reasoning behind that decision: the evidence it rested on, the methodology that governed it, and the judgment that produced it. That standard does not change when AI enters the production process. What changes is where in the system the reasoning, the methodology, and the judgment need to live.

Human-centric AI is where they live. HCAI is a design and production philosophy in which institutional expertise governs the AI both by upfront design and throughout the production process. Rather than building powerful machines as an end in themselves, HCAI builds reliable, safe, and trustworthy systems that make the people who use them, governance researchers and investors, more effective and better informed.⁶ Putting humans at the center is also what makes AI suitable for a regulated industry: safe, transparent, and accountable.

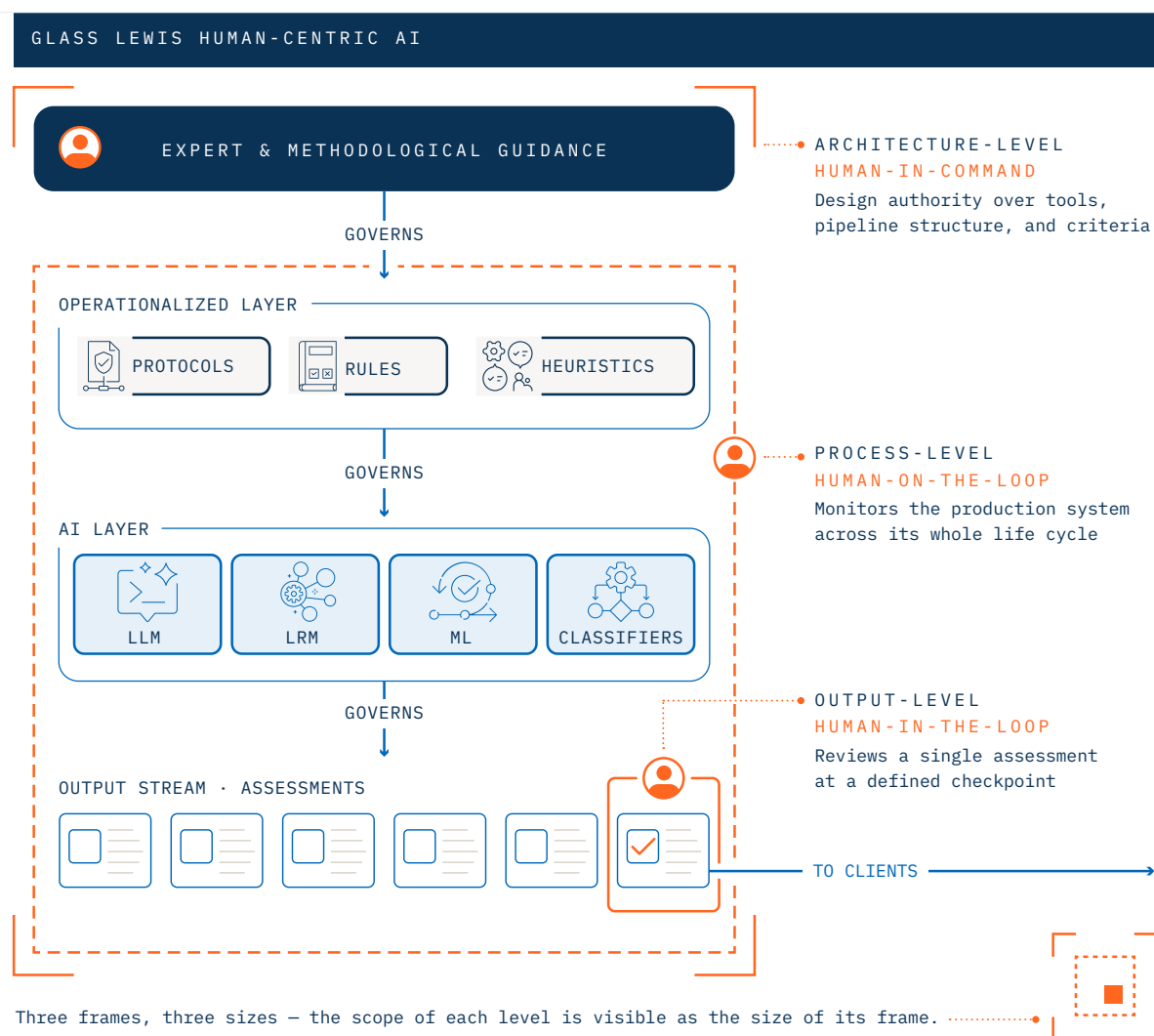
As a philosophy, the human-centric AI approach rests on a set of core principles for how AI should be governed:

- **Human control:** Experts and clients remain in charge of the process and can override, direct, or set aside the AI's judgment at any point.
- **Transparency:** The architecture is designed for both explainability and interpretability. Explainability means the system can give the evidence and reasons behind its outputs in a way users can follow, and those explanations reflect how the system actually reached them.^{7,8} Interpretability means the internal structure of the AI system and the causes of its decisions are open and clearly understandable by humans.
- **Ethical alignment:** The system's objectives are held to core values, including fairness, data privacy, and the prevention of bias.

Implementation Through Design: A Hybrid AI Architecture and Human Oversight

These principles are realized through two components. The first is the hybrid AI architecture described in Section 1,⁹ in which rule-based systems carry analyst and methodological logic and govern the generative models. The second is a human oversight model operating at three distinct levels. Each of the three levels governs a different scope of the system, looks at a different unit of analysis, and performs a governance function the other two cannot (Figure 2).

Figure 2: Human-Centric AI Combining the Three-Level Oversight Model With Hybrid AI Architecture



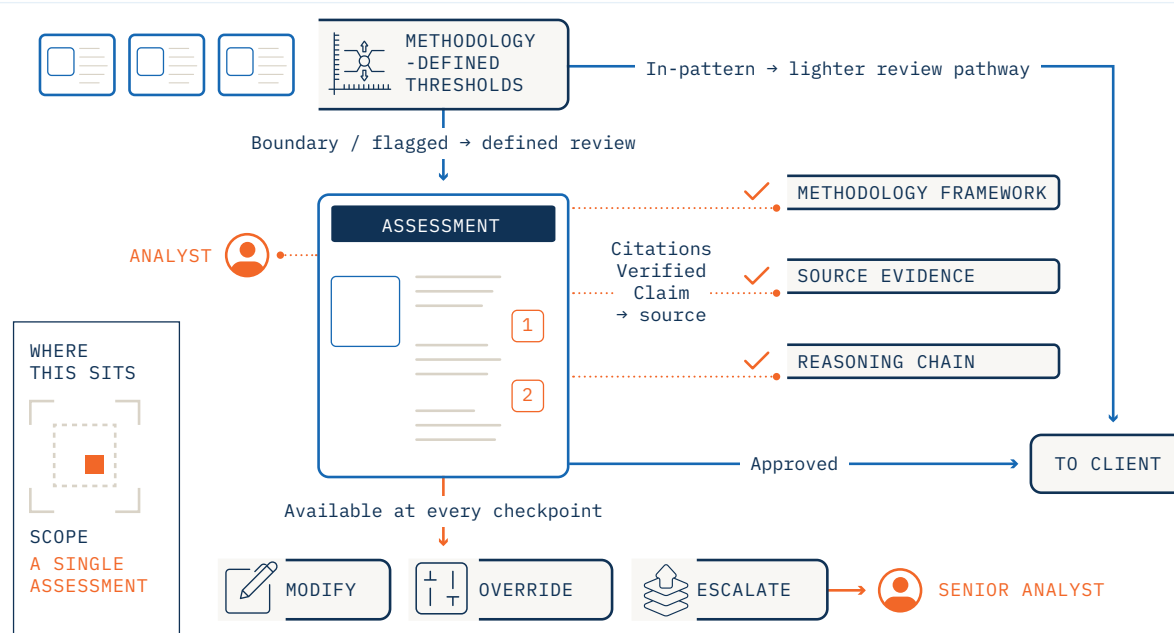
Source: Summarized by Glass Lewis and informed by the EU High-Level Expert Group Ethics Guidelines for Trustworthy AI and the NIST AI Risk Management Framework.

Output-Level Oversight: Human-in-the-Loop

The scope of this level is the individual high-stakes analytical output: a specific assessment reviewed by an analyst before it reaches a client.

At this level, an analyst reviews an AI-produced assessment against three things: the methodology framework that governs the research, the source evidence the AI relied on, and the reasoning chain the AI followed to reach its conclusion (Figure 3). The review is substantive. The analyst can see which documents informed the output and which methodology rules governed the inference. Modification, override, or escalation is available at each checkpoint.

Figure 3: Visual Overview of Output-Level Oversight



Citations are part of how the review works. The AI generates a citation for every claim in an assessment. As part of the review process, analysts verify the accuracy of those citations against the source documents, confirming that the reasoning chain is traceable from claim to evidence and that the cited material genuinely supports the conclusion drawn.

The depth of review varies with what the assessment requires. Reviewers check that the reasoning holds together, that the evidence supports the conclusion, and that the result sits where the methodology would expect it to. Assessments that raise a question — whether through a logical gap, a drift off topic, or a result that looks out of step with the evidence — receive closer attention. This is what allows AI to handle the production workload while preserving substantive, human judgment where it matters most.

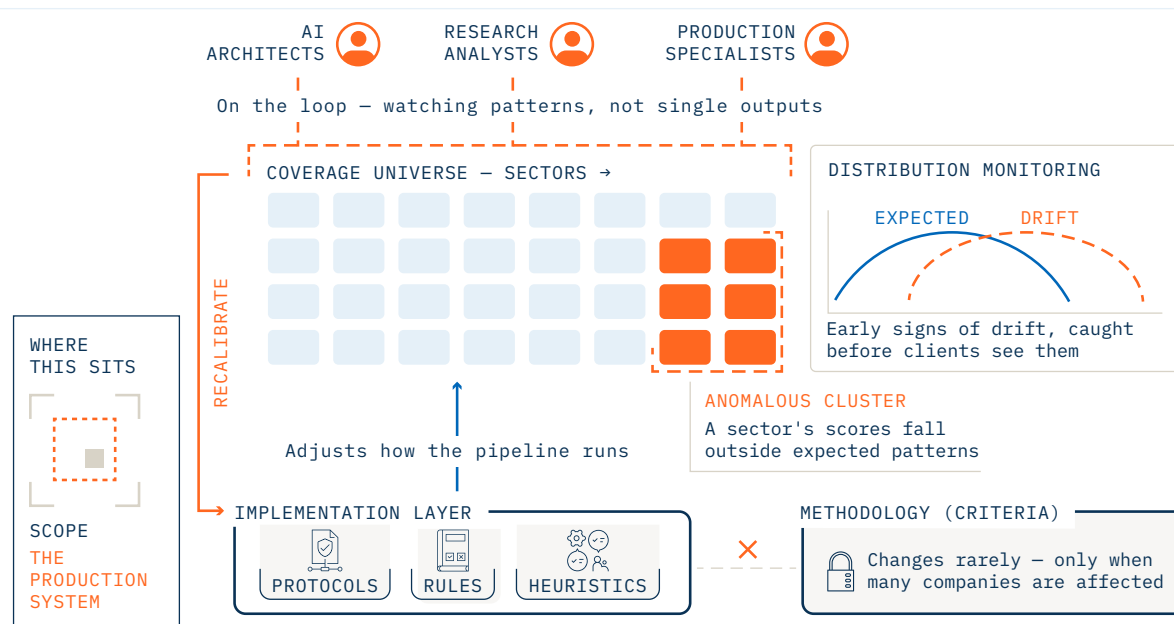
When organizations describe human curation of AI, in the large majority of cases they mean oversight at this level alone: human verification and checking of output. That work is necessary and important for producing high-quality output and building trust. It is not sufficient on its own for investor-grade data and insight. More is needed, and the next two levels are where this comes from.

Process-Level Oversight: Human-on-the-Loop

The scope of this level is the production system across its entire life cycle, a bird's eye view of how the system is performing, distinct from reviewing any individual assessment.

At this level, AI architects, research analysts, and production specialists monitor streams of outputs for patterns that indicate something has shifted: quality drift, anomalous score distributions across a sector or region, or a class of assessments consistently triggering escalation (Figure 4). These patterns are invisible at the level of any individual review, regardless of how thorough that review is. They are the patterns that matter most for the consistency and reliability of what clients receive across the full coverage universe.

Figure 4: Visual Overview of Process-Level Oversight



Research analysts contribute substantively at this level. They monitor for places where the methodology is applied unevenly across sectors or company types, document the exceptions that arise, and feed those signals back to the AI system so its operations can be adjusted. When anomalies surface, the response is recalibration: adjusting protocols and heuristics, modifying how the pipeline handles a particular document type, or flagging a category of assessments for elevated review at the output level. Most adjustments at this level operate on the implementation, not on the underlying methodology.

This level matters for a further reason. AI systems degrade over time through data drift and concept drift: the inputs and the underlying concepts move away from what the models were built and tested on, and models can overfit to the conditions they were tuned against. Process-level oversight is what catches the early signs of that drift, before it reaches the output a client sees.

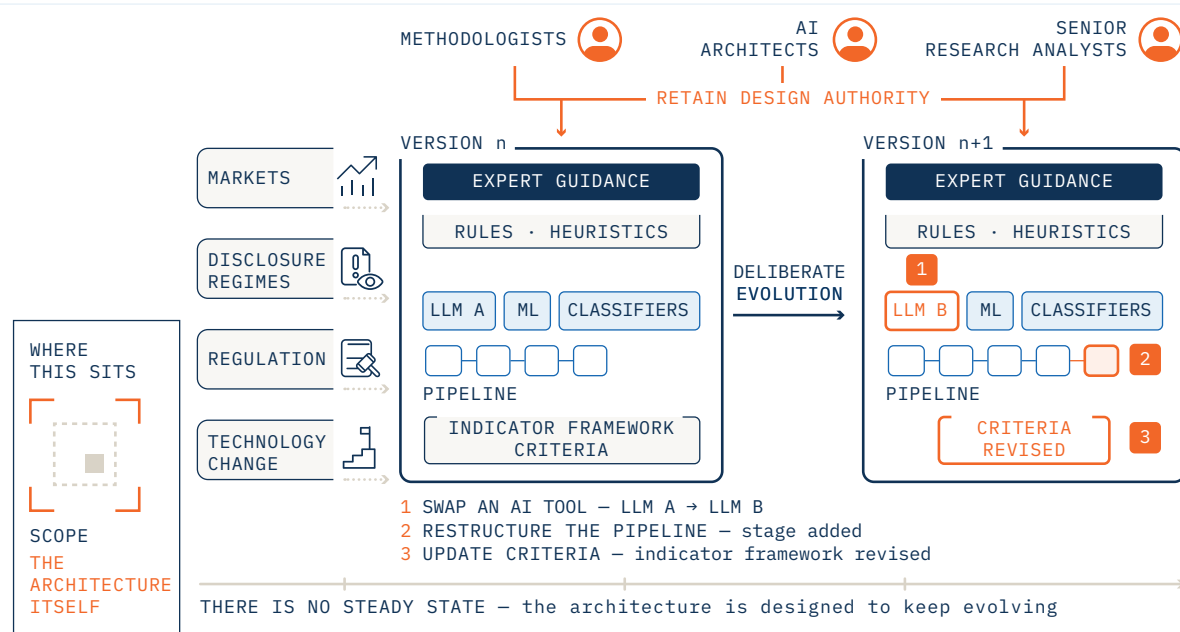
Architecture-Level Oversight: Human-in-Command

The scope of this level is the design and fitness of the particular AI approach itself. At issue is whether the architecture as a whole is built, and remains, fit for the work it supports.

Governance of AI-powered research is not a one-time build. Markets evolve, disclosure regimes change, regulatory expectations shift, and the AI technologies themselves change rapidly. Today's large language models will not be the systems we run on in two years. The architecture that is fit for today's production system will need to transform, proactively, as conditions evolve. There is no steady state.

At the architecture level, methodologists, AI architects, and senior research analysts retain the authority to determine whether the AI tools selected, the pipeline design, the indicator framework, and the implementation rules remain fit for the work (Figure 5). The production system is designed to evolve in response to the dynamics it covers. Architecture-level oversight is the governance mechanism that ensures it does so deliberately, through expert judgment, rather than through unexamined drift.

Figure 5: Visual Overview of Architecture-Level Oversight



The three levels operate concurrently. Output-level review happens at methodology-defined checkpoints. Process monitoring runs continuously across the coverage universe. Architecture governance is exercised both proactively and reactively, whenever the fitness of the approach comes into question. Together, they provide oversight across individual analytical judgments, production system health, and the design integrity of the architecture itself.

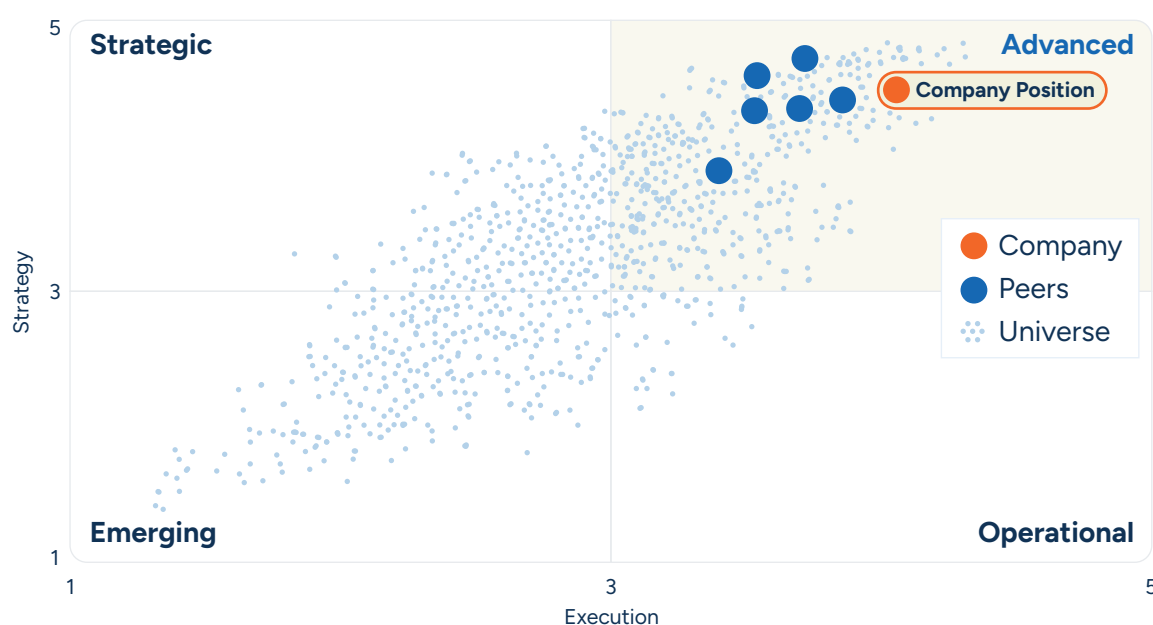
Section 3: Glass Lewis' Climate Intelligence

Glass Lewis launched its new Climate Intelligence product in April 2026. It is the most direct, current illustration of what the human-centric AI production philosophy produces in a complex, forward-looking governance domain. This section walks through what the product delivers, the methodology that governs it, and how the production architecture works.

What Climate Intelligence Delivers

Climate Intelligence is a forward-looking transition assessment platform for institutional investors, investment teams, and portfolio managers. It evaluates how companies are managing the financial materiality of the climate transition — both the strategic direction they have set and how effectively they are executing against it.

Figure 6: Visual Overview of Climate Intelligence Assessment Framework



The output is a structured analytical assessment for each company in the coverage universe.¹⁰ Each assessment includes: a strategy score and an execution score, individual indicator-level scoring across 20 indicators, peer comparison within the company's subindustry, a financial materiality assessment, identification of strengths and gaps with specific recommendations, and detailed analytical commentary with full source traceability to the corporate documents the assessment drew from.

Investment professionals and portfolio managers use these assessments to prioritize engagement, identify advanced and emerging companies within sectors, inform voting decisions on climate-related shareholder proposals, integrate climate context into broader governance and investment analysis, and report to stakeholders and clients with defensible, evidence-based analysis.

The distinction from disclosure-based ESG data is specific. Disclosure-based approaches summarize what companies report about themselves, which means companies that do not articulate their strategy in climate-specific language can fall off the map. Climate Intelligence does not require companies to frame their disclosures around climate. The methodology translates generic financial disclosures into what matters for the transition: capital allocation, revenue model, operational strategy, and governance structure. The assessment evaluates whether a stated strategy is credible, and whether execution evidence supports the claim (see Figure 7 below). The first answers what the company says. The second answers what investment teams need to know to make decisions they can defend, and shows where value is likely to be created or eroded across the transition.

The Transition Assessment Framework

The methodological foundation of Climate Intelligence is the transition assessment framework developed by the climate research team. The framework evaluates companies along two axes: Strategy (i.e., the direction and credibility of the transition plan) and Execution (i.e., the evidence and capability behind implementation).

Figure 7: Example Climate Intelligence Assessment



The two axes produce four positions: Strategic Advanced (strong strategy and strong execution), Strategic (strong strategy, execution developing), Operational (execution underway, strategy incomplete), and Emerging (limited progress on both). This framework translates climate into a governance and performance issue, assessed with the same analytical rigor as any other material driver of long-term value.

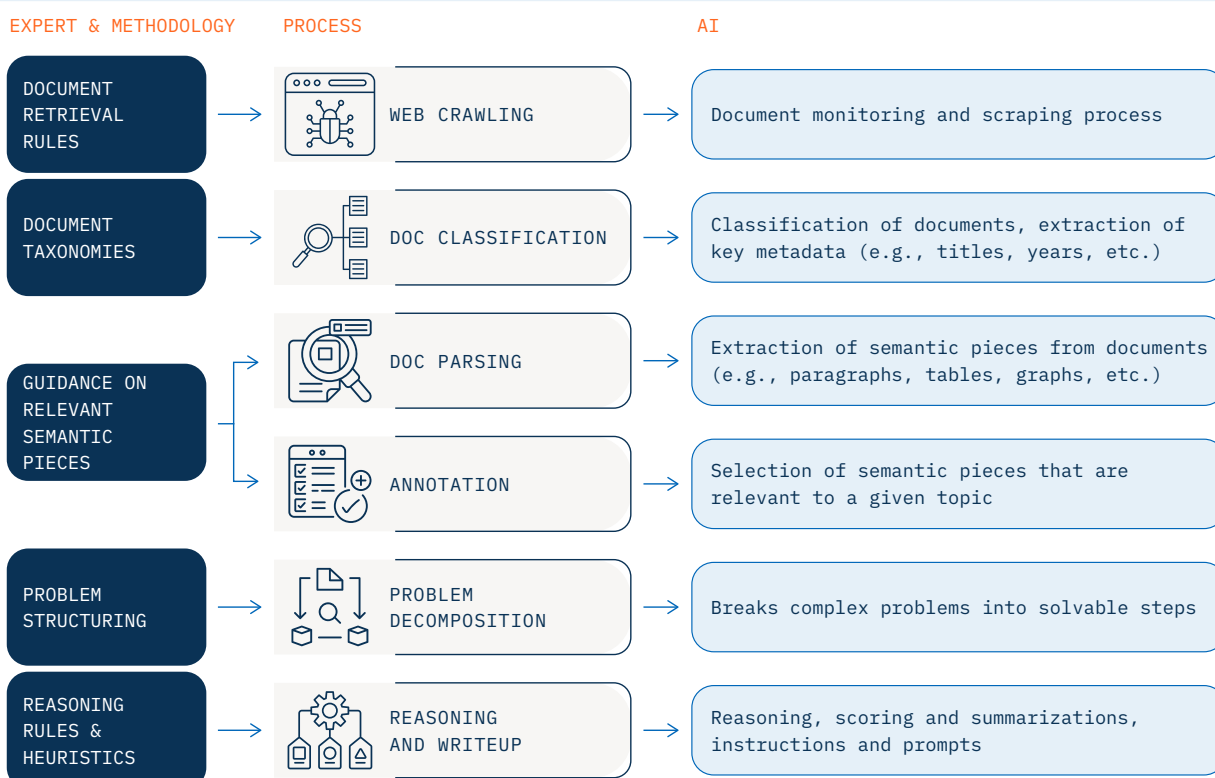
Underneath these positions, the framework operationalizes the assessment across 20 indicators, organized into themes covering opportunity and risk identification, governance and board oversight, strategy development and scenario analysis, executive leadership and execution capability, encompassing organizational capacity, talent, assets, supply chain, market traction, funding, and financial performance.

This framework is the methodology layer the AI operates inside. It came from the climate research and methodology team. It defines what the AI is permitted to conclude.

Inside the Human-Centric AI Production Architecture

The Climate Intelligence production pipeline runs through six stages, from document acquisition through reasoning and writeup. At each stage, expert guidance shapes what the AI is permitted to do, the AI's operation is transparent and inspectable, and each stage produces outputs that can be modified or overridden. The AI's work is bounded by the methodology layer throughout (Figure 8).

Figure 8: AI Powered Climate Intelligence, Six Stage Pipeline With Expert Guidance Flowing Into Each Stage



The principle running through the architecture is simple: the AI does not interpret governance concepts on its own. At every stage, it operates within the rules, heuristics, and reasoning paths the methodology layer defines. The output it produces can be traced back through those rules to the specific evidence in the source documents that supported each inference.

The Three Levels in Climate Intelligence

The three levels of human-centric AI from Section 2 operate in the Glass Lewis Climate Intelligence production environment as follows.

At the output level (human-in-the-loop), climate analysts review individual company assessments at methodology-defined checkpoints. A company whose transition position falls at the boundary between framework categories, an assessment that triggers an engagement priority, a finding with significant implications for a voting recommendation: each triggers a defined review process. The system is designed so that the analyst can see the AI's reasoning chain, the evidence it selected, and the methodology rules that governed the inference, with this visibility deepening as the production system matures.

At the process level (human-on-the-loop), AI architects, research analysts, and production specialists monitor the system in two directions: horizontally, across the full coverage universe, and vertically, along the AI-powered reasoning for each individual company. When score distributions across a sector or region fall outside expected patterns, or when a class of disclosures is consistently triggering escalation, the rule set or pipeline configuration is recalibrated. Most adjustments at this level work on the implementation: protocols, heuristics, retrieval, and aggregation.

Changes to the methodology itself happen rarely and only when many companies are affected. Here is an illustrative example: in early production, analysts noticed a pattern of lower-than-expected scores for certain types of companies. The foundational issue was traced to insufficient document acquisition for those companies and the classification errors that followed. Some of the resolutions were central improvements to the crawler and classification rules. Some were manual, uploading important documents for companies that had been missed.

At the architecture level (human-in-command), methodologists and AI architects exercise ongoing design authority over whether the AI tools, pipeline structure, indicator framework, and implementation rules remain fit for what the assessment requires. This is where the continuous nature of the work shows up most clearly. As disclosure regimes evolve and the climate transition shifts, as new AI technologies become available and existing ones are deprecated, and as data and concepts drift, the architecture governance function ensures the system adapts deliberately.

For example, in early production, two indicators related to talent and governance disclosure showed a floor effect. The framework was designed to measure these at the level of detail investors needed, but disclosure norms across companies were not yet at that level

of fidelity, so most companies received low scores that were not decision-useful. The response was a criteria update that acknowledged the disclosure environment as it actually is, while preserving the ability to recognize companies that disclose at the higher fidelity.

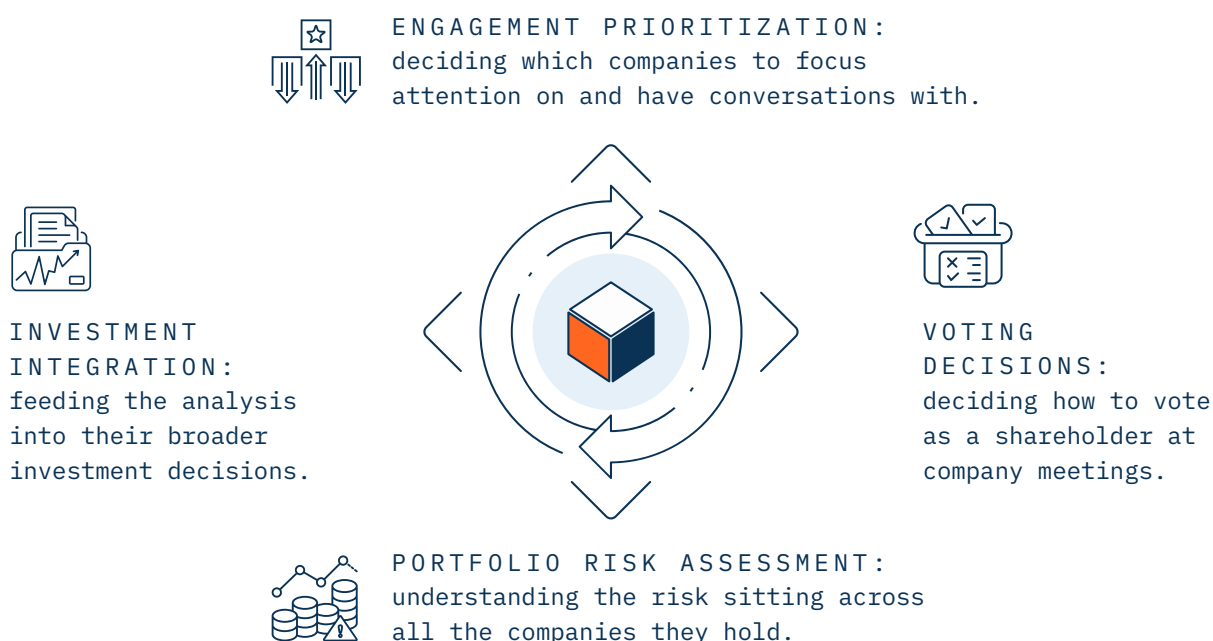
The team behind Climate Intelligence is a working example of co-design. Climate researchers, methodology specialists, product researchers, AI designers and architects, research analysts, and data engineers designed the system together. The same team monitors and refines it.

What Stewardship Teams and Portfolio Managers Receive

Stewardship teams and investment professionals receive transition assessments organized by company and sector. Each assessment places the company in the framework based on a structured evaluation across the 20 indicators, with individual indicator scores, subindustry averages and peer comparison, financial materiality assessment, strengths and gaps analysis, detailed analytical commentary, and a full list of the source documents the assessment drew from.

For investment stewardship teams, this changes the engagement conversation. Rather than starting from disclosure language, teams start from a structured, evidence-grounded view of where each company sits, what is driving that position, and what would move it. Engagement priorities follow from the analysis, not from the volume of disclosure available (Figure 9). Sector benchmarking becomes analytically grounded through comparable assessments across a consistent methodology, and voting on climate-related shareholder proposals can be informed by structured analytical evidence rather than disclosure summaries alone.

Figure 9: Climate Intelligence in Stewardship and Investment Workflows



For portfolio managers, the same assessments support investment decisions with analysis that is comparable across the universe and defensible at the holding level. Specific data points investment teams care about, including green capital expenditure, are standardized across industries in a way that ad-hoc disclosure analysis cannot easily produce. Portfolio-level transition risk reporting becomes possible with comparable inputs: asset owners overseeing external managers, and asset managers reporting to their own clients, can ground reporting in structured assessments that hold up to scrutiny. The financial materiality framing in the assessment is what makes the outputs usable in investment workflow, not just in stewardship workflow.

The outputs are forward-looking, focused on trajectory. They are comparable across companies and sectors because they apply a consistent methodology. They are traceable, with each inference linked to specific source evidence. The architectural properties that produce these outputs are: governed methodology, three levels of human oversight over the AI, source traceability, and a neuro-symbolic AI ecosystem. These are the properties Glass Lewis is building toward across its governance research.

The Five Questions, Answered

In *AI and the Fiduciary Test*,¹¹ Glass Lewis introduced five questions any asset manager can use to evaluate an AI-enabled governance research provider. Climate Intelligence is a working answer to all five, and walking through them is the most direct way to see what the architecture this paper has described produces in practice.

What governs the underlying data?

Document acquisition is governed by expert-defined taxonomies that specify both the document types relevant to climate research (i.e., annual reports, sustainability disclosures, transition plans, regulatory filings) and the substantive content to retrieve from them (i.e., climate-related disclosures, but also generic financial content the methodology translates into climate-relevant signals). The taxonomies serve two purposes: funneling information into the system in ways that reflect its financial materiality (e.g., CFO-attested content like annual reports is weighted differently from sustainability marketing material that may or may not be financially relevant), and ensuring coverage is highest for the most material disclosures. Classification, parsing, and annotation rules operate on top of this foundation, all governed by the methodology before any inferential reasoning occurs.

What is the human actually doing?

Three levels of oversight operate concurrently and complement one another. At the output level (human-in-the-loop), analysts review assessments at methodology-defined checkpoints, including verifying citations against source documents. At the process level (human-on-the-loop), AI architects, research analysts, and production specialists monitor system health across the coverage universe and along the full life cycle of the research, adjusting protocols when patterns shift. At the architecture level (human-in-command), methodologists and AI architects govern the fitness of the approach, with criteria changes reserved for situations affecting many companies. Each level is staffed, active, and documented.

How is investor-grade defined?

Climate Intelligence data and insights are validated against source evidence (accuracy), maintained through deterministic rules across sectors (consistency), verified against known coverage (completeness), traceable from source disclosure to final assessment (traceability), conformed to specified formats (validity), and delivered within the timeframes stewardship workflow requires (timeliness). Validation combines deterministic methods, including automated checks that every assessment must pass at defined points in the pipeline, with human review of selected outputs, where analysts work through the reasoning chains, validate sources, and feed any mistakes they find back into protocol corrections. Scaled validation also includes peer-comparison reasoning: examining whether higher and lower scores across a sector make economic sense, and tracing back when they do not.

What is the designed scope?

Climate Intelligence covers a defined universe specified by the methodology layer, which sets out what the system is designed to assess and where the boundaries are. Current coverage and the expansion schedule are maintained on the Climate Intelligence product page. The framework is sector-agnostic by design, evaluating companies through archetypal business activities rather than fixed sectoral criteria. Where a company's disclosure falls outside the framework's designed scope, the system identifies the gap rather than generating an unsupported conclusion.

What is the accountability structure for exceptions?

When the system encounters a non-routine situation, the escalation pathway routes to the domain expertise that applies. The accountability is layered. Routine anomalies surface through process-level monitoring and are addressed through protocol adjustments. Outputs flagged for substantive concern at the output level are escalated to senior analysts for validation, and outputs with multiple anomalies can trigger regeneration or further escalation. Where the issue is genuinely architectural, methodology and AI architecture leadership address it through criteria or pipeline changes.

Section 4: The Standard Beyond Climate Intelligence

Climate Intelligence is one production environment. The standard it illustrates runs across everything Glass Lewis is building. This closing section steps back from that specific application to the broader picture.

The Glass Lewis Approach

The architecture described in this paper is the architecture Glass Lewis is building toward across its governance products. Climate Intelligence is the first production environment in which it operates at scale. The governed data architecture examined in part two of this series is the foundation it sits on. It is also the foundation for future work, including early-stage assessment of how AI could support multiple, client-specific research perspectives.

The HCAI framework is the AI design and production philosophy that runs through all of it. Institutional expertise underlies the methodology. The methodology governs the AI. Human oversight operates at three levels in production. Source traceability is preserved from disclosure document to final output. The operational coverage that informs the methodology layer reflects more than two decades of governance research expertise across global markets, in more than 40 languages, within more than 100 distinct regulatory regimes. That operational presence is what gives the methodology layer the depth required to govern AI in a fiduciary context. The expertise came first. The architecture follows from it.

There is no finished version of this architecture. New technologies will arrive, regulations will evolve, and disclosure regimes will change. The expertise that governs it, and the discipline of applying that expertise as technologies, regulations, and disclosure practices change, is what gives the work its durability. The five questions in AI and the Fiduciary Test, and the architectural lens in this paper, are questions Glass Lewis can answer specifically, today and as the work evolves. The firm invites any client, prospect, or evaluator to ask them.

Notes and References

1. Brutian, A. and Beigi, K. *AI and the Fiduciary Test: A Guide for Institutional Investors in Evaluating AI Proxy Voting Solutions*. Glass Lewis. April 4, 2026. <https://www.glasslewis.com/ai/fiduciary-test>.
2. Gustitus, C. and Brutian, A. *The Foundation Matters More Than the Model: Why Trusted Data and Human-Centric AI Will Define the Next Era of Investment Stewardship*. Glass Lewis. June 10, 2026. <https://www.glasslewis.com/ai/the-foundation-matters-more-than-the-model>.
3. Krippendorff, K. *Content Analysis: An Introduction to Its Methodology*. Fourth Edition, Thousand Oaks, CA: SAGE Publications, Inc., 2019. <https://doi.org/10.4135/9781071878781>.
4. Garcez, A. D. and Lamb, L. C.. 2023. "Neurosymbolic AI: the 3rd wave." *Artif. Intell. Rev.* 56, 11 (Nov 2023), 12387–12406. <https://doi.org/10.1007/s10462-023-10448-w>.
5. See #2 above.
6. Shneiderman, B. 2020. "Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy." *International Journal of Human–Computer Interaction* 36 (6): 495–504. [doi:10.1080/10447318.2020.1741118](https://doi.org/10.1080/10447318.2020.1741118).
7. Phillips, P., Hahn, C., Fontana, P., Yates, A., Greene, K., Broniatowski, D. and Przybocki, M. 2021. Four Principles of Explainable Artificial Intelligence, NIST Interagency/Internal Report (NISTIR), National Institute of Standards and Technology, Gaithersburg, MD, [online], <https://doi.org/10.6028/NIST.IR.8312>, https://tsapps.nist.gov/publication/get_pdf.cfm?pub_id=933399.
8. Biran, O. and Cotton, C. 2017. "Explanation and Justification in Machine Learning: A Survey." *IJCAI-17 Workshop on Explainable AI*. https://www.cs.columbia.edu/~orb/papers/xai_survey_paper_2017.pdf.
9. Dellermann, D., Ebel, P., Söllner, M. et al. "Hybrid Intelligence." *Bus Inf Syst Eng* 61, 637–643 (2019). <https://doi.org/10.1007/s12599-019-00595-2>.
10. The coverage universe is defined by the Climate Intelligence methodology. Current coverage is maintained on the Climate Intelligence product page..
11. See #1 above.

Acknowledgements

The authors would like to thank the following contributors whose expertise and insights substantially contributed to the contents in this white paper:

Emil Moldovan, Senior Director, Climate

Alka Bhide, Senior Manager, Climate

Glossary of Terms

This glossary extends the appendix in AI and the Fiduciary Test. New terms introduced in this paper are defined below.

Methodology Layer

In an AI-powered governance research system, the structured operationalization of expert judgment that governs the operations and applications of AI. In a traditional workflow, methodology guides how analysts think and operate. In an AI system, it also defines the governing architecture and operational guidance that amplify institutional expertise across more outputs than analysts could cover alone.

Implementation Layer (Symbolic Layer)

The protocols, heuristics, and operational rules that flex the methodology to a specific AI system's strengths and constraints. The methodology itself (the criteria) changes rarely and only when many companies are affected. Implementation adjusts in real time as production patterns emerge.

Output-Level Oversight (Human-in-the-Loop)

The level of HCAI oversight at which analysts review individual high-stakes outputs against the methodology, source evidence, reasoning chain, and citations. The depth of review varies with what the assessment requires.

Process-Level Oversight (Human-on-the-Loop)

The level of HCAI oversight at which AI architects, research analysts, and production specialists monitor the production system across many outputs for quality drift, anomalous patterns, and distribution shifts, and for the signs of data drift and concept drift over time. Most adjustments at this level work on the implementation layer rather than the methodology.

Architecture-Level Oversight (Human-in-Command)

The level of HCAI oversight at which methodologists and AI architects retain the authority to determine whether the AI approach itself remains fit for its purpose, and to extend, restructure, or replace it. A strategic governance function exercised continuously, not a one-time design step.

Hybrid AI Ecosystem

An AI architecture combining different AI technologies (large language models, classifiers, extraction tools, rule-based systems) with embedded human expertise, matching each tool to the stage of production it serves best.

Neuro-Symbolic AI

The technical implementation of the hybrid AI ecosystem. Generative AI (machine-learning and large language model architectures) is combined with a rule-based component made up of formal logic, expert heuristics, and explicit constraints. That rule-based component acts as the executive function, or prefrontal cortex, of the generative system.

Transition Assessment

A forward-looking evaluation of how a company is managing the financial materiality of the climate transition, assessed along two axes: strategy (direction) and execution (delivery).



The Architecture of Investor-Grade AI

Glass Lewis is a leading independent governance research and proxy advisory firm serving more than 1,300 institutional investors globally.

Connect With Glass Lewis

Corporate Website | glasslewis.com

Email | info@glasslewis.com

Social | [!\[\]\(77e670be72de63f664b9f3cf25895195_img.jpg\)](#)