

Beyond the Universal League Table: A Structured Multi-Factor Framework for Industry Rankings

Keith Lee

Swiss Institute of Artificial Intelligence, Chaltenbodenstrasse 26, 8834 Schindellegi, Schwyz, Switzerland

Abstract

Many institutional and commercial rankings reduce a broad population of organizations to a single ordered list. The apparent simplicity of this approach conceals a demanding mathematical assumption: that one scalar function can represent excellence across organizations with different operating models, assets, capabilities, and markets. This article develops the mathematical rationale for a structured ranking network in which p heterogeneous inputs contribute to k related but non-identical ranking outputs. The proposed architecture combines shared and category-specific latent components, permits the same firm to appear in multiple rankings, and allows the relevance and weight of an input to vary across outputs. Hierarchical layers provide flexibility when new evidence or new ranking categories enter the system, while Restricted Boltzmann Machine modules and Gibbs updates offer one implementation for estimating relationships between adjacent layers. Information criteria can help compare candidate network structures without being confused with validation of the final rankings themselves. The framework also motivates tiered publication: when rankings are jointly estimated and neighboring scores cannot always be separated with high confidence, grouped recognition is often more defensible than a claim of exact ordinal precision. The result is neither an assumptionless algorithm nor a mechanical replacement for industry knowledge. It is a disciplined measurement architecture that combines mathematical structure, empirical evidence, and documented expert interpretation.

Keywords: industry rankings; factor analysis; latent variables; multi-output systems; Restricted Boltzmann Machine; Gibbs sampling; tiered rankings; ranking methodology

Methodological Scope Note: This paper was prepared to support public understanding of the methodological rationale underlying the ranking architecture jointly developed by the Swiss Institute of Artificial Intelligence (SIAI) and The Economy Network, particularly for Ranking News. Its purpose is to explain the system's basic architecture and mathematical intuition to a general professional audience. It is therefore a conceptual and practice-oriented exposition, rather than a complete technical specification, empirical validation

study, or exhaustive account of the models and procedures used in producing individual rankings. Certain implementation details, estimation choices, data-governance procedures, and category-specific adjustments are consequently presented only at a general level.

1 Introduction

Rankings are an efficient form of institutional communication. A large volume of financial, operational, and qualitative information is condensed into a format that readers can understand immediately. Yet the communicative strength of a ranking can conceal the difficulty of its underlying measurement problem.

A conventional ranking usually begins with multiple indicators and ends with one score for each organization. If $x_i = (x_{i1}, \dots, x_{ip})^\top$ denotes the evidence collected for organization i , a general single-index ranking can be represented as

$$s_i = \phi \left(\sum_{j=1}^p w_j T_j(x_{ij}) \right), \quad r_i = \text{rank}(-s_i),$$

where $T_j(\cdot)$ places indicator j on an appropriate measurement scale, w_j is its weight, s_i is the composite score, and r_i is the published position.

This structure is not necessarily incorrect. For a narrowly defined population pursuing a sufficiently homogeneous objective, one composite measure may be informative. The difficulty arises when a broad institutional label is treated as though it represented one market and one form of excellence.

Consider the financial sector. Commercial banks, investment banks, life insurers, non-life insurers, asset managers, private equity firms, and specialist private-capital organizations are all serious financial institutions. Nevertheless, their economic functions are not interchangeable. Total assets may be central to the operating model of a commercial bank or insurer, but it does not carry the same meaning for an advisory-led investment bank, an asset-light financial intermediary, or a private equity firm organized through externally committed funds. A universal financial-institutions ranking based heavily on balance-sheet size would not merely omit information. It would confer a structural advantage on a particular business model.

The same problem appears within advisory services. Procurement advisory, supply-chain advisory, business-process transformation, automation and productivity advisory, AI implementation, digital transformation, and cloud infrastructure are closely connected. Some firms operate across several of them. But the services remain different enough that one global weight vector cannot represent every market equally well.

The methodology developed here therefore replaces the conventional many-input, one-output structure with a many-input, multi-output system. Formally, n organizations are observed through p evidence variables and evaluated across k related ranking categories. The central question is not how to force every organization onto one line. It is how to represent several forms of sector-specific strength without losing their economic connections.

2 The Mathematical Cost of a Universal Ranking

2.1 One-dimensional reduction

Every universal ranking performs a dimensional reduction. A collection of different observations becomes one scalar score. Even where the published methodology is not based on principal component analysis, PCA provides a useful mathematical analogy.

For standardized evidence $\tilde{x}_i$, the first principal component is

$$z_{i1} = v_1^\top \tilde{x}_i,$$

where v_1 is selected to maximize the sample variance of the projected data. The share of total variation represented by that component is

$$\frac{\lambda_1}{\sum_{j=1}^p \lambda_j}.$$

Unless this proportion is overwhelmingly large, a one-component representation necessarily leaves substantial variation outside the published score. PCA creates successive uncorrelated components that maximize remaining variance; the first component is not presumed to represent the full structure of the data (Jolliffe and Cadima, 2016).

The analogy has an important boundary. Most rankings are not literally principal-component models, and the direction of greatest statistical variation is not automatically the direction of greatest economic importance. The lesson is narrower: a one-dimensional projection is a strong structural restriction, regardless of whether its weights are equal, judgmental, regression-based, or data-derived.

2.2 Fixed weights imply a universal meaning

Under a linear composite score,

$$s_i = \sum_{j=1}^p w_j x_{ij},$$

the marginal contribution of evidence item j is fixed:

$$\frac{\partial s_i}{\partial x_{ij}} = w_j.$$

The same indicator therefore carries the same marginal importance for every evaluated organization. If the candidate population contains several business models, the model assumes that the meaning of the indicator remains stable across them.

That assumption often fails. Geographic office coverage may be crucial for a client-facing implementation network but much less informative for a specialist strategic adviser. Assets under management are meaningful for an asset manager but inappropriate as a direct measure of private-equity operating expertise. The issue is

not simply that a researcher may choose an imperfect weight. It is that a single universally correct weight may not exist.

2.3 A ranking is a latent measurement problem

The distinction between regression and ranking construction is fundamental. In an ordinary regression,

$$Y_i = \beta_0 + X_i^\top \beta + \varepsilon_i,$$

Y_i is an observed dependent variable whose conditional relationship with X_i is estimated. In an institutional ranking, overall category strength is usually not observed independently. It is the latent construct that the researcher is attempting to measure.

The ranking problem therefore does not begin with a single observable Y for which one seeks a best unbiased estimator under a common conditional error law. It begins with heterogeneous evidence and a question about how many economically meaningful constructs are required to represent that evidence. This makes latent-variable and factor structures more natural than a conventional supervised regression.

2.4 Nonlinearity does not solve the classification problem by itself

A nonlinear function can fit one dataset more closely than a linear index:

$$s_i = g(x_i; \theta).$$

But a close in-sample fit does not establish that the same function will remain meaningful across industries or years. A nonlinear model may learn the peculiarities of one candidate universe, one information environment, or one period. Transfer to a different market requires stability in the underlying economic relationships, not merely greater functional flexibility.

This is why the principal design question comes before nonlinear estimation. If the underlying market contains several forms of excellence, a more elaborate $p \rightarrow 1$ function still produces only one ordering. The model becomes more flexible without addressing the substantive mistake of forced unidimensionality.

3 Three Sources of Distortion in Single-Output Rankings

3.1 Heterogeneous evidence

Ranking evidence can include financial values, operating counts, market coverage, professional authorizations, service descriptions, transaction records, text-derived measures, ordinal evaluations, and indicators of organizational specialization. These observations do not all have the same economic meaning or statistical behavior.

No single conditional distribution for a scalar ranking outcome is naturally implied by this evidence. More importantly, there is no observed ranking Y that supplies an external criterion for estimating one universally

correct set of weights. A single-output model may still be constructed, but its apparent precision should not be confused with identification of an objectively true ordering.

3.2 Measurement error, missing evidence, and feedback

Measured indicators are imperfect. Some capabilities are easier to observe than others. Large organizations generally leave larger public data footprints, while specialist firms may generate less visible but more relevant evidence. Several sources may also reproduce the same underlying information, causing one characteristic to enter the model repeatedly under different labels.

There can also be feedback between reputation and measurement. Prior recognition may increase media visibility, recruitment, client inquiries, and future information availability. Commercial success may increase both genuine operating capability and the volume of observable evidence. These relationships can create circularity even when the final ranking score is mechanically constructed from current inputs.

A multi-output architecture does not eliminate measurement error or feedback. Its contribution is more specific: it reduces the risk that evidence relevant to one type of activity is treated as a universal measure of every activity. Source review, temporal controls, duplicate-evidence checks, and expert examination remain necessary.

3.3 Excessive compensability

Composite rankings usually allow strength in one indicator to compensate for weakness in another. This may be reasonable within a coherent construct. It becomes questionable when unrelated capabilities are aggregated.

A very large firm may compensate for limited specialization through size, visibility, or geographic breadth. Conversely, a specialist may be penalized for lacking capabilities that are irrelevant to the service in which it competes. The result can be mathematically consistent with the chosen formula while economically misclassified.

4 From p Inputs to k Connected Rankings

4.1 Basic factor representation

Let $x_i \in \mathbb{R}^p$ represent the evidence vector for firm i , and let $f_i \in \mathbb{R}^k$ represent its latent strengths across k ranking categories. A basic multi-factor representation is

$$x_i = \mu + \Lambda f_i + \varepsilon_i,$$

where $\Lambda \in \mathbb{R}^{p \times k}$ is the loading matrix and ε_i contains evidence-specific variation.

Instead of receiving one universal score, firm i receives a vector:

$$f_i = (f_{i1}, f_{i2}, \dots, f_{ik})^\top.$$

Each component supports a different ranking output. The outputs are jointly connected because they are

estimated from overlapping evidence and may share latent capabilities. They are nevertheless distinct because each output can receive a different loading pattern.

4.2 Structured relevance and sparse connections

Not every input should affect every output. Define an admissibility matrix

$$M \in \{0, 1\}^{p \times k},$$

and write

$$\Lambda = M \odot B,$$

where $\odot$ denotes elementwise multiplication. If $M_{jr} = 0$, evidence item j cannot directly affect ranking category r . If $M_{jr} = 1$, its weight may be estimated from the data and the surrounding network.

This arrangement permits three economically different cases:

1. An indicator can be central to one category.
2. It can contribute with different weights to several related categories.
3. It can be excluded from a category to which it has no defensible relationship.

For example, cloud-infrastructure capability may be heavily relevant to cloud advisory, moderately relevant to digital transformation and AI implementation, and irrelevant to an unrelated advisory category. The architecture represents these differences directly instead of forcing a common weight across all outputs.

4.3 Connected does not mean identical—or orthogonal

Related ranking categories need not be statistically independent. Let

$$\text{Cov}(f_i) = \Psi,$$

where Ψ may contain non-zero off-diagonal elements. This allows two category factors to share meaningful variation.

The objective is therefore not to force complete orthogonality between connected markets. Luxury-yacht and superyacht capabilities, for example, may share design, engineering, supplier, and production foundations. The categories should be distinguishable, but their correlation is economically legitimate. A model that artificially eliminates all common variation could be as misleading as a model that collapses both markets into one.

4.4 Multiple rankings for a single firm

The multi-output structure has an important practical implication: a firm need not be assigned exclusively to one ranking category.

Define eligibility as

$$e_{ir} \in \{0, 1\},$$

where $e_{ir} = 1$ indicates that firm i conducts a material and relevant activity in category r . The eligible universe for ranking r is

$$\mathcal{U}_r = \{i : e_{ir} = 1\}.$$

There is no restriction requiring

$$\sum_{r=1}^k e_{ir} = 1.$$

A diversified firm may therefore appear in several rankings. Its position in each category is determined by the category-specific score and evidence weights, not by its placement elsewhere.

This is not duplication. It is a direct consequence of recognizing that firms can operate across multiple slices of an interwoven market. A consultancy may possess genuine capabilities in AI implementation, automation, digital transformation, and cloud infrastructure. Forcing it into one category would discard information; granting it one universal score would erase the differences among those capabilities. Multiple category appearances preserve both breadth and specialization.

5 Distribution-Agnostic Structure and Estimation Choice

5.1 What distribution-agnostic means

The structural $p \rightarrow k$ framework does not require the researcher to posit one Gaussian—or other common parametric—conditional distribution for a scalar ranking output. It is not organized around estimating $E[Y | X]$ for one observed Y . At that architectural level, it is distribution-agnostic.

The distinction between architecture and estimator is important. A linear factor structure can be estimated through covariance reconstruction and unweighted least squares. Minimum-residual factor analysis, for example, selects loadings by minimizing residuals in the reproduced correlation matrix. Unweighted least-squares factor estimation has been developed as a distribution-free alternative to maximum-likelihood factor analysis (Harman and Jones, 1966; Krijnen, 1996).

However, distribution-agnostic does not mean assumptionless. The selection of an objective function, factor dimension, admissible loading structure, and evidence transformation still matters. If a probabilistic RBM module is used between layers, that module defines its own energy-based probability model. The framework avoids one universal distributional assumption for the ranking problem; it does not claim that every possible implementation is free of modelling assumptions.

5.2 Restricted Boltzmann Machines between adjacent layers

The implementation considered here uses Restricted Boltzmann Machine modules to estimate selected relationships between adjacent layers. For visible state v , hidden state h , and weight matrix W , a basic RBM can be represented through the energy function

$$E(v, h) = -a^\top v - b^\top h - v^\top W h.$$

The associated joint probability is

$$P(v, h) = \frac{\exp[-E(v, h)]}{Z},$$

where Z is the partition function.

Because an RBM contains no within-layer connections, hidden states are conditionally independent given visible states, and visible states are conditionally independent given hidden states. This bipartite structure permits alternating block Gibbs updates:

$$h^{(t)} \sim P(h \mid v^{(t)}), \quad v^{(t+1)} \sim P(v \mid h^{(t)}).$$

Gibbs sampling is therefore not invoked as a general claim about every deep-learning architecture. It is used because the selected layerwise implementation is built on RBM structures for which alternating conditional updates provide a natural estimation mechanism. RBMs have long been used as two-layer generative learning modules and as components of deeper systems (Hinton, 2012).

5.3 Comparing candidate network structures

The number of output factors, intermediate nodes, and layers should not be selected through an unguided search over every possible architecture. Economic taxonomy supplies the initial structure: which services exist, which indicators can reasonably affect them, and which capabilities may be shared. Statistical model comparison then evaluates competing representations within that constrained space.

Where comparable likelihoods can be evaluated or consistently approximated, information criteria may be written as

$$\text{AIC} = -2\ell(\hat{\theta}) + 2d,$$

and

$$\text{BIC} = -2\ell(\hat{\theta}) + d \log n,$$

where $\ell(\hat{\theta})$ is the fitted log-likelihood and d is the effective number of estimated parameters. AIC and BIC balance fit against architectural complexity rather than rewarding the largest possible network (Akaike, 1974; Schwarz, 1978).

Their role here is specific. They can help compare whether, for example:

- an additional latent node improves the representation sufficiently to justify its complexity;
- an intermediate capability layer fits better than a direct input-to-output mapping;
- two proposed output categories are sufficiently distinguishable;
- a new branch improves the architecture more than it fragments the evidence.

Information criteria do not certify that the published ranking is substantively correct. Nor do they evaluate a Gibbs sampler merely because it uses different update steps. They compare candidate statistical architectures represented by different dimensions or restrictions. Convergence diagnostics, stability, interpretability, and industry coherence must accompany the numerical comparison.

6 Hierarchical Flexibility

6.1 Intermediate capability layers

A direct $p \rightarrow k$ factor system may be sufficient for a narrow market. A broader ranking network can use intermediate layers:

$$x_i \rightarrow h_i^{(1)} \rightarrow h_i^{(2)} \rightarrow f_i,$$

where the $h_i^{(m)}$ terms represent shared or partially shared capabilities.

In advisory services, intermediate factors might represent geographic execution, sector expertise, technical implementation, client acquisition, regulatory capability, transaction delivery, or operational infrastructure. Final ranking categories then depend on different combinations of these intermediate capabilities.

6.2 Adding new evidence

When a new data source becomes available, it does not have to connect directly to every final ranking. It can enter the intermediate capability to which it is economically related. The model can then estimate how that capability contributes to the connected outputs.

This does not make the system automatically more accurate. It makes the system easier to extend without rebuilding every ranking as an isolated formula. New evidence can be absorbed locally, its propagation can be controlled, and its broader effects can be examined across the connected network.

6.3 Adding or bifurcating outputs

Market development may also require a new ranking output. Suppose one initial marine-industry factor contains both luxury-yacht and superyacht capability. As evidence coverage improves, the upper layer can branch:

$$f_{i,Y} = \beta_Y h_{i,M} + \gamma_Y h_{i,L} + u_{i,Y},$$

$$f_{i,S} = \beta_S h_{i,M} + \gamma_S h_{i,C} + u_{i,S},$$

where $h_{i,M}$ represents shared marine capability, $h_{i,L}$ represents luxury-market capability, and $h_{i,C}$ represents complex bespoke-project capability.

Both outputs preserve their common industrial foundation while receiving different category-specific information. The separation is not created by arbitrarily changing final weights until two different lists appear. It is created by introducing a structured branch at the layer where the economic distinction arises.

6.4 Flexibility rather than automatic robustness

Additional layers increase representational flexibility. They can make it easier to incorporate new inputs, new intermediate capabilities, and new output categories. They do not guarantee greater robustness merely by existing.

Every additional node or layer increases the number of possible relationships and therefore the risk of overfitting or weak identification. The advantage of the hierarchical system lies in controlled extensibility. Whether a deeper structure improves the ranking network remains a question for model comparison, stability review, and economic interpretation.

7 Why Tiered Publication Follows from the Model

7.1 Exact positions can imply excessive precision

The latent score for category r may provide an ordering of eligible firms, but neighboring scores can be close. Their relative positions may also depend on evidence that contributes to several connected outputs. A change in one shared component can affect more than one category through jointly determined weights.

Publishing every firm as though the difference between positions 7 and 8 were exact and isolated would overstate what the system can establish. This concern becomes stronger when two related ranking categories share capabilities or when firms operate across both markets.

Tiered publication is therefore not merely a visual design choice. It is a methodological acknowledgement that the system measures relative zones of performance more reliably than every adjacent ordinal distinction.

7.2 Institutionally calibrated tier sizes

Ranking News commonly uses tier structures such as 5/10/5, while larger publications may use different group sizes, including formats such as 5/10/30. These numbers are institutionally determined from accumulated observations across category rankings. They are not presented as a universal mathematical ratio.

The appropriate partition can vary with:

- the size of the eligible market;
- the density of credible candidates;

- the degree of observable separation among firms;
- the maturity of the category;
- the amount and quality of available evidence;
- the intended breadth of published recognition.

The first tier identifies the most consistently differentiated group. Subsequent tiers recognize strong or credible participants without claiming that every internal position is precisely separable. A tier boundary remains a publication rule, not proof of a discontinuity in the underlying industry.

7.3 Tiering as methodological humility

The connected ranking system does not claim perfect decomposition of overlapping markets. Factor structures do not convert noisy industrial evidence into naturally discrete classes. Tiering accepts this limitation openly.

The methodology therefore makes a narrower and more defensible claim: the available evidence supports differentiated groups within a category, but it may not support strong conclusions about every one-place difference inside those groups. This form of restraint increases the credibility of the output.

8 Ranking Stability and Controlled Recalibration

Input comparability is a precondition for estimation. Once the evidence has been placed on appropriate and mutually interpretable scales, the methodological question becomes whether reasonable specification changes produce excessive positional movement.

The stability review should examine:

- alternative evidence windows;
- inclusion or exclusion of unusually influential sources;
- changes in candidate eligibility;
- reasonable adjustments to admissible connections;
- alternative numbers of intermediate nodes;
- missing-evidence treatments;
- year-to-year changes in the information environment.

The purpose is not to search for a specification that preserves a predetermined list. It is to identify whether a small methodological change causes disproportionate movement in tier membership or category placement. If it does, the output should be interpreted more cautiously, the architecture reconsidered, or the publication groups widened.

This review is especially important in a joint system. Because an input can influence shared intermediate capabilities, a model adjustment may propagate across several outputs. That propagation is not necessarily an error: it may reflect a genuine shared component. But it should be traceable. A connected ranking network is defensible only when researchers can distinguish an economically meaningful system-wide effect from an unintended consequence of model specification.

9 Factor Discovery and Human Interpretation

9.1 Why factors do not name themselves

Factor analysis groups shared variation. It does not determine the economically correct name for every hidden component. Two statistically similar loading patterns may admit different substantive interpretations, while rotationally related solutions may reproduce the evidence comparably well.

Industry experts must therefore examine:

- the evidence entering each component;
- the signs and relative sizes of its loadings;
- its relationship with neighboring factors;
- the firms receiving high and low factor scores;
- whether the inferred construct corresponds to a recognizable market activity.

This is not an optional editorial layer added after the mathematics. It is part of construct validation.

9.2 Two principal interpretation risks

Human interpretation creates at least two risks:

1. **Misnaming:** a statistically identified component is given a category name that does not accurately describe its evidence pattern.
2. **Misdefinition:** the category name may be reasonable, but the stated operating scope includes services that the component does not measure coherently.

Both risks can be reduced through documented category definitions, comparison of alternative factor structures, and review by specialists familiar with the relevant industry.

9.3 Expert review without arbitrary reweighting

Expert intervention should be structural rather than outcome-seeking. Appropriate interventions include:

- defining candidate eligibility;
- identifying economically impossible connections;

- specifying anchor indicators;
- distinguishing structural non-applicability from missing evidence;
- proposing a new factor or branch when the market has changed;
- correcting an interpretation contradicted by the loading structure.

Inappropriate intervention would include changing weights solely because a preferred firm did not receive the expected position. Every material structural adjustment should be documented and followed by renewed model comparison and stability review.

The separation between editorial ranking decisions and commercial licensing is equally important. Recognition licenses may govern how a result is communicated, but they must not determine eligibility, weights, tier placement, or continued editorial inclusion.

10 Structural Illustrations from Ranking Markets

10.1 Financial institutions

A universal financial-institution ranking would struggle to reconcile balance-sheet scale, underwriting capacity, transaction expertise, externally managed assets, fund deployment, geographic distribution, and advisory networks. A multi-output structure instead permits related rankings for banks, insurers, asset managers, private-equity firms, and other financial specialists. Shared financial capability remains visible, but each category receives its own evidence structure.

10.2 Operations and technology advisory

Procurement, supply chain, business-process transformation, automation, AI implementation, digital transformation, and cloud infrastructure form an interconnected advisory space. Firms may operate in one, several, or nearly all of these markets.

The $p \rightarrow k$ architecture permits a project-delivery indicator to affect several categories with different weights. It also allows specialized evidence—such as cloud architecture capability or procurement-market intelligence—to enter only the relevant branches. A diversified adviser may appear in multiple rankings, while a specialist can compete on the evidence appropriate to its own market rather than being overwhelmed by universal scale indicators.

10.3 Luxury yachts and superyachts

Luxury-yacht and superyacht markets illustrate why connected outputs should not be forced into either complete merger or complete independence. The markets share industrial capabilities and some firms operate in both. Yet bespoke project complexity, vessel scale, customer structure, engineering demands, and production economics can support separate evaluations.

A hierarchical branch preserves their common foundation and introduces distinct upper-layer components. When the evidence does not justify sharp separation, the methodology can retain correlated factors and broader tiers rather than impose arbitrary final weights. As coverage improves, the branch can become more differentiated without discarding the earlier network structure.

11 Limitations

The proposed architecture improves the representation of heterogeneous markets, but it does not transform ranking into an exact science.

First, the model remains dependent on evidence quality. Missing, duplicated, promotional, or systematically biased information can affect the latent structure.

Second, factor separation is incomplete by nature. Closely related economic activities may remain strongly correlated, and some hidden components may combine more than one substantive capability.

Third, architecture selection is not purely mechanical. AIC, BIC, or other fit measures can compare candidate structures, but economic interpretability and institutional judgment remain necessary.

Fourth, RBM and Gibbs implementations carry probabilistic and computational assumptions. They are estimation tools, not guarantees of an objectively true market taxonomy.

Fifth, additional layers increase flexibility and extensibility, but they may also increase overfitting, weak identification, and opacity.

Sixth, tier boundaries are institutionally calibrated publication decisions. They communicate groups supported by accumulated evidence, but they do not establish natural discontinuities between adjacent firms.

Finally, multiple appearances across rankings must be governed by substantive eligibility. Cross-category recognition is a strength only when it reflects genuine business activity rather than the mechanical advantage of corporate scale.

12 Conclusion

The central weakness of a universal league table is not that it necessarily uses the wrong formula. It is that it assumes one formula can represent organizations whose economic functions differ.

A structured $p \rightarrow k$ ranking network offers a more coherent alternative. It allows several rankings to share evidence and hidden capabilities without treating them as identical. The same input can receive different weights across outputs, irrelevant connections can be excluded, and firms with several genuine business lines can appear in multiple categories. Intermediate layers permit new evidence and new rankings to enter through controlled extensions rather than isolated redesigns.

The approach also changes how ranking outputs should be communicated. Because category scores are jointly determined and related markets cannot always be separated perfectly, tiered publication is often more defensible than exact one-place ordering. Institutionally calibrated tiers acknowledge the limits of the evidence while preserving meaningful recognition.

No mathematical architecture eliminates the need for judgment. Factor components must be interpreted, category boundaries must be defined, and structural changes must be reviewed. The purpose of the model is not to remove expert responsibility, but to discipline it. Mathematics determines how evidence can be connected; industry expertise determines whether those connections describe a real market.

The resulting ranking is neither a universal scalar nor a collection of unrelated lists. It is a connected measurement system: one designed to preserve the common foundations of an industry while recognizing that excellence can take more than one form.

References

- Akaike, H. (1974). “A New Look at the Statistical Model Identification.” IEEE Transactions on Automatic Control, 19(6), 716–723.
- Harman, H. H., and Jones, W. H. (1966). “Factor Analysis by Minimizing Residuals (Minres).” Psychometrika, 31, 351–368.
- Hinton, G. E. (2012). “A Practical Guide to Training Restricted Boltzmann Machines.” In Neural Networks: Tricks of the Trade.
- Jolliffe, I. T., and Cadima, J. (2016). “Principal Component Analysis: A Review and Recent Developments.” Philosophical Transactions of the Royal Society A, 374, 20150202.
- Krijnen, W. P. (1996). “Algorithms for Unweighted Least-Squares Factor Analysis.” Computational Statistics & Data Analysis, 21(2), 133–147.
- OECD and European Commission Joint Research Centre. (2008). “Handbook on Constructing Composite Indicators: Methodology and User Guide.” OECD Publishing.
- Schwarz, G. (1978). “Estimating the Dimension of a Model.” The Annals of Statistics, 6(2), 461–464.