

TECHNICAL SUPPLEMENT: SCENARIO CONTROL, STRUCTURAL ATTRIBUTION, AND CLOSURE TESTS

*The Mathematical Foundations for Migrating from Operational Industry Peer Rankings to
Bound Capital Market Valuation Regimes*

Hany Hassanien Badr

Independent Researcher in Corporate Finance & Valuation Systems | Banker
Cairo, Egypt

Scenario Architecture and Structural Attribution in Closed-Loop Peer Ranking

Operating Regimes, Valuation Boundaries, Required-Return Specifications, and Closure Tests

Why Scenario Control Is Necessary

The preceding analysis established RANK as a closed-loop framework linking operating-capital efficiency, economic value creation, cash-flow capacity, financing structure, and shareholder return. That framework can be interpreted correctly only when every reported output is attached to the economic regime that generated it. The present chapter therefore extends the book's central argument: structural ranking requires not only internally coherent calculations, but also disciplined attribution across alternative operating and valuation environments.

A reader encountering several ROIC sequences, WACC figures, IRRs, and per-share values may reasonably ask whether the model contradicts itself. The correct answer depends on scenario identity. The computational cases were developed for different analytical purposes: portability, baseline valuation, operating sensitivity, required-return control, and terminal-growth sensitivity. Their outputs should not be forced into numerical equality when their periods, sales paths, valuation dates, discount regimes, or terminal assumptions differ.

The fundamental error would be to compare a number from one scenario with a number from another while silently assuming that both measure the same economic experiment. The correct procedure is attribution: first identify the scenario; then identify the changed assumption; then compare only the outputs that remain on a common boundary.

Observed Difference

$$\begin{aligned} &= \text{Operating Effect} + \text{Timing Effect} + \text{Required} \\ &\quad - \text{Return Effect} + \text{Terminal} - \text{Assumption Effect} \\ &\quad + \text{Claim/Unit Effect} + \text{Interaction/Residual Effect} \end{aligned}$$

The chapter therefore serves the RANK framework directly. Structural ranking rewards explainable economic strength. The same standard of explainability must govern the model evidence supporting the rank.

Scenario attribution as an identification problem

A cross-scenario comparison is an identification exercise. An observed change in value or rank can be attributed to operating performance only when valuation timing, required-return specification, terminal growth rate (g), claim boundary, and unit scale are held constant. When several of these elements change together, the comparison remains informative, but its interpretation must be decomposed rather than assigned to a single cause.

Valid Operating Attribution requires $\Delta\text{Timing} = \Delta\text{WACC Regime} = \Delta g = \Delta\text{Claim Boundary} = \Delta\text{Scale} = 0$, with definitions and model structure held constant

This condition does not imply that every useful comparison must be single-variable. It means that causal language must match the controls actually preserved. The baseline-to-higher-sales comparison provides the clearest operating attribution because the sales path changes within the same broken-date floating-WACC architecture. The higher-sales broken-date case and the full-year closure benchmark principally identify timing and required-return effects because the improved operating path is retained while the valuation architecture changes. The independent reference dataset is a portability test, while the replicated case is a provenance control. The evidential contribution of the scenario family lies in this differentiated architecture, not in treating every computational file as an equal forecast of one case.

Analytical Scenario Architecture

Reader-facing regime	Analytical role	Period / boundary	Operating and valuation specification	Function within RANK
Independent portability case	Independent reference dataset	2011–2015; full-year; different scale	Fixed WACC 16.39%; $g=10\%$; distinct sales, capital, debt and equity scale	Tests structural portability
Broken-date baseline	Historical baseline	2005–2009; valuation begins 14 Sep 2005; 108/365 stub	Baseline sales; floating annual WACC; $g=7\%$	Establishes the reference operating path
Higher-sales sensitivity	Operating sensitivity	2005–2009; broken-date	Sales +5% versus baseline; floating annual WACC; $g=7\%$	Isolates operating-path improvement
Fixed-WACC growth sensitivity	Combined sensitivity	2005–2009; full-year	Baseline sales; fixed WACC 17.6754%; $g=10\%$	Tests fixed-WACC and growth assumptions
Full-year closure benchmark	Closure benchmark	2005–2009; full-year	Sales +5%; fixed WACC 17.6754%; $g=7\%$	Tests exact DCF–EVA convergence

Reader-facing regime	Analytical role	Period / boundary	Operating and valuation specification	Function within RANK
Replication case	Provenance control	2005–2009; broken-date	Replicates the higher-sales floating-WACC path	Confirms lineage; adds no new observation

Table 1. Reader-facing analytical architecture. The regimes are distinguished by their economic assumptions, valuation boundaries, and inferential functions rather than by internal workbook filenames.

Reader's One-Page Control Map

The following map gives the reader a safe order of interpretation. It is deliberately asymmetric: the regimes do not have equal evidential status and should not be read as competing valuations of one unchanged case.

Reader question	Primary regime	Comparison regime	Valid conclusion
Does the architecture transfer?	Portability case	Baseline or closure benchmark	Transfer is supported, but scale is not directly comparable.
What is the historical baseline?	Broken-date baseline	Higher-sales sensitivity	Baseline sales, floating WACC, and a 108/365 stub.
What changes when operations improve?	Higher-sales sensitivity	Broken-date baseline	The comparison attributes the higher-sales operating path.
What demonstrates exact closure?	Full-year closure benchmark	Higher-sales sensitivity	Fixed-WACC DCF–EVA equality at 335.0839/share.
What does the $g=10\%$ case show?	Fixed-WACC growth sensitivity	Baseline and closure benchmark separately	A combined assumption sensitivity, not the closure result.
Does the replicated case add evidence?	Replication case	Higher-sales sensitivity	No new economic observation; it preserves provenance.

Reader rule: identify scenario → identify changed assumption → confirm common boundary → compare output → interpret the RANK implication. Reversing this order invites false contradiction.

Independent Portability Case

The independent portability case applies the same 28-sheet closed-loop logic to a 2011–2015 dataset on a substantially different numerical scale. It tests whether the capital-return-income-cash-value architecture transfers to another period and magnitude. It does not supply the 2005–

2009 RANK figures and should not be placed inside the same annual ranking table without explicit normalization and a defined research purpose.

Metric	2011	2012	2013	2014	2015
Net sales	24,845,849	49,813,093	58,328,680	73,164,457	89,386,579
NOPLAT	2,232,620	6,033,950	6,545,333	8,825,094	11,149,245
ROIC	15.40%	54.41%	51.94%	67.72%	84.47%
FCFF	5,643,350	4,521,047	6,116,298	8,657,297	10,873,014

The portability case reports operating value of approximately 109,505,604.24, equity value of approximately 103,505,604.24, and per-share value of approximately 10.35056 in its own scale. These quantities are not directly comparable with the 2005–2009 model-unit values. Their evidential role is portability: they show that the architecture can be instantiated on another dataset, not that the case is another forecast of the same firm.

Broken-Date Baseline Regime

The broken-date case is the principal historical baseline for the RANK analysis. The valuation begins on 14 September 2005, leaving a first-period fraction of 108/365, or approximately 0.2958904. Annual WACC is floating; the 2005 annual WACC is approximately 15.7822%, while the first valuation-period capital charge is approximately 4.6698%. These two rates are not contradictory: one is annual, the other is the stub-period equivalent.

$$15.7822\% \times 108/365 = 4.6698\%$$

Metric	2005	2006	2007	2008	2009
Sales	21,804	27,996	32,918	36,258	40,549
ROIC	82.13%	80.08%	65.40%	58.71%	62.98%
Annual WACC	15.78%	16.83%	17.35%	17.64%	17.69%
FCFF	2,501	3,557	6,032	8,300*	10,391

*A financing-path schedule displays 8,299 in 2008 because of a one-unit capital aggregation convention (19,134 versus 19,135). This is controlled rounding, not an economic defect. One schedule must be used consistently within any formal proof.

The broken-date baseline does not achieve exact displayed DCF–EVA equality. Depending on the calculation path, DCF per share is approximately 320.9936–320.9963, while the Economic Value display produces approximately 321.7155. The gap is a temporal-closure diagnostic arising from the stub-period architecture and related capital-charge/display conventions. Accordingly, USD 321.71 should be identified as the EVA/economic-value display path, not described as the unique DCF answer.

Higher-Sales Operating Sensitivity under Floating WACC

The higher-sales sensitivity increases sales by five percent relative to the broken-date baseline while preserving the floating-WACC architecture and 7 percent terminal growth. Its purpose is operating attribution: it asks how higher sales and the associated operating path change ROIC, NOPLAT, FCFF, and value while the required return continues to move annually.

Metric	2005	2006	2007	2008	2009
Sales	22,894.20	29,395.80	34,563.90	38,070.90	42,576.45
NOPLAT	6,300.59	8,916.75	10,472.71	11,565.63	13,661.31
ROIC	93.26%	90.95%	74.29%	66.91%	71.40%
FCFF	3,252.59	4,622.75	7,284.71	9,717.63	12,001.31

The higher-sales case produces DCF operating value of approximately 82,188.37, DCF equity value of 81,638.37, and DCF per share of approximately 371.0835. Its EVA path produces approximately 371.8059 per share. These outputs are sensitivity results, not the basis of the full-year 335.0839 closure benchmark.

Fixed-WACC Regimes Must Remain Analytically Distinct

Baseline-sales fixed-WACC growth sensitivity

The baseline-sales fixed-WACC sensitivity applies a required return of approximately 17.6754% and terminal growth of 10 percent to the baseline operating path. It produces operating value of approximately 83,326.10, equity value of approximately 82,776.10, and per-share value of approximately 376.25. This is a legitimate combined sensitivity, but it is not the source of the USD 335.0839 full-year closure result.

Full-year higher-sales closure benchmark

The full-year closure benchmark combines the five-percent higher-sales path with a full-year boundary, fixed WACC of approximately 17.675365%, and terminal growth of 7 percent. Its operating path matches the higher-sales broken-date case. The comparison therefore isolates timing and required-return architecture more cleanly than a comparison that also changes the operating path.

Control	Full-year closure result
ROIC 2005–2009	93.2592%; 90.9501%; 74.2851%; 66.9075%; 71.3981%
FCFF 2005–2009	3,252.59; 4,622.75; 7,284.71; 9,717.63; 12,001.31
Fixed WACC	17.675365%
Terminal g	7.00%
Terminal FCFF	12,841.4017
Terminal operating value	120,290.0452
DCF operating value	74,268.4586

Control	Full-year closure result
DCF equity value	73,718.4586
DCF/EVA per share	335.0839028

The full-year benchmark provides the strongest closure test in the scenario family. DCF and EVA converge exactly at displayed precision: operating value 74,268.458619, equity value 73,718.458619, and per-share value 335.083902813. Its lower headline value relative to the higher-sales broken-date sensitivity is not evidence that better operations destroyed value. The valuation boundary and discount architecture differ; the regimes answer different questions.

Replication, Provenance, and the Avoidance of Double Counting

A separately archived workbook reproduces the sales path, WACC path, ROIC, FCFF, DCF/EVA outputs, and IRRs of the higher-sales broken-date sensitivity. Its distinct file identity does not create a distinct economic observation. The replicated case is valuable for lineage and reproducibility, but presenting it as an additional empirical result would overstate the evidence by counting one experiment twice.

Interpretive rule: a replicated computational file may corroborate implementation and provenance, but it must be excluded from the count of independent economic regimes.

Controlled Comparisons and Attribution Order

Comparison	What changes	What can be inferred	What must not be claimed
Portability vs 2005–2009 regimes	Period, scale, dataset, capital and claim structure	Structural portability of the closed-loop architecture	Direct performance superiority or rank
Baseline vs higher-sales sensitivity	Sales/operating path; floating-WACC architecture retained	Effect of operating uplift under moving required returns	Pure discount-rate effect
Higher-sales vs closure benchmark	Same operating path; broken-date/floating vs full-year/fixed WACC	Timing and required-return attribution	Additional operating improvement
Baseline vs fixed-WACC growth case	Timing, WACC and g all change	Combined valuation-architecture sensitivity	Single-variable causal attribution
Higher-sales vs replication case	No material economic change	Implementation replication and provenance	Two independent confirmations

Table 2. Correct comparison order. The baseline-to-higher-sales comparison is principally operational; the higher-sales-to-closure comparison principally identifies the valuation regime; the portability case tests transferability; and the replicated case establishes traceability.

Reconciling the Two ROIC/WACC Presentations in the RANK Book

The book contains two prominent ROIC sequences. The 82.13% to 62.98% sequence belongs to the baseline-sales broken-date regime. The 93.26% to 71.40% sequence belongs to the higher-sales operating path used in both the operating sensitivity and the full-year closure benchmark. These sequences are not alternative calculations of the same numerator on the same operating case. They reflect different sales and NOPLAT paths.

Year	Baseline-sales ROIC	Higher-sales ROIC	Full-year WACC fixed
2005	82.13%	93.26%	17.6754%
2006	80.08%	90.95%	17.6754%
2007	65.40%	74.29%	17.6754%
2008	58.71%	66.91%	17.6754%
2009	62.98%	71.40%	17.6754%

For interpretive completeness, every exhibit should identify the scenario, sales regime, valuation boundary, and WACC convention. A reader should never have to infer these attributes from the percentage sequence alone.

IRR and XIRR Are Convention-Specific Controls

The broken-date baseline reports model-period IRRs of approximately 77.5514% for total invested funds and 143.6988% for operating invested funds. They differ because the capital bases differ, while both calculations use the model's fractional-period timing convention. A legacy slide value of 87.40% cannot replace the workbook-controlled 143.70% operating-funds IRR.

A displayed XIRR block associated with the broken-date case contains a confirmed 100-times terminal-scale defect and cannot support inference. When the terminal input is corrected, actual-date XIRRs are approximately 77.4929% for total funds and 143.5897% for operating funds. The full-year closure benchmark, by contrast, has valid annual-period IRRs of approximately 67.5717% and 113.3123%, and actual-date XIRRs of approximately 67.5306% and 113.2446%. IRR and XIRR must be labeled rather than treated as interchangeable.

Regime / convention	Total invested funds	Operating invested funds	Interpretive status
Broken-date model-period IRR	77.5514%	143.6988%	Valid when the convention is named
Broken-date corrected XIRR	77.4929%	143.5897%	Valid after correction of the terminal input
Full-year annual-period IRR	67.5717%	113.3123%	Annual-period control

Regime / convention		Total invested funds	Operating invested funds	Interpretive status
Full-year XIRR	actual-date	67.5306%	113.2446%	Dated-cash-flow control

Implications for Structural Peer Ranking

The scenario family strengthens the RANK framework when used according to its analytical design. The broken-date baseline shows how the original capital engine behaves under a historical timing boundary and floating required return. The higher-sales sensitivity tests whether operating improvement raises NOPLAT, ROIC, FCFF, and value without changing that architecture. The full-year fixed-WACC benchmark subjects the improved operating path to a different valuation boundary and provides an exact closure test. The independent dataset tests portability, while the replicated case confirms implementation lineage without adding an economic observation.

- Rank operating quality from ROIC, economic spread, EVA, cash conversion, and reinvestment within one declared regime.
- Do not rank regimes by headline equity value when timing, WACC, g, claim boundary, or scale differs.
- Use the full-year benchmark for exact DCF–EVA closure and the broken-date baseline for temporal-closure diagnostics.
- Treat an extreme ROE or operating-equity return only after confirming the denominator and sign.
- Use workbook arithmetic as computational evidence; screenshots and visual exhibits are explanatory corroboration.

Scenario identity changes the interpretation of rank in three ways. First, it determines which operating path generated the score: baseline or higher sales. Second, it determines whether the hurdle moves with financing and risk assumptions or is held fixed as a control. Third, it determines whether value is measured at a broken date or on a full-year boundary. A rank is therefore valid only within a disclosed analytical regime. The framework ranks firms, years, or decisions after normalization; it does not reward whichever workbook happens to produce the highest per-share value.

This distinction is central to management-quality assessment. Management may legitimately receive credit for improvement between the baseline and higher-sales paths because NOPLAT, ROIC, and FCFF rise within a retained valuation architecture. Management should not receive credit for value caused only by a lower discount rate or a more favorable terminal assumption. Conversely, a more demanding hurdle should not be misdescribed as operating deterioration. The RANK score must attribute performance to the layer in which the change actually entered.

Boundaries of Inference

The scenario family is designed to clarify attribution, but it does not eliminate every limitation. First, the portability case uses a different period, scale, and capital structure; it demonstrates transferability of the architecture rather than statistical generalizability. Second, the baseline-

sales fixed-WACC growth sensitivity changes more than one valuation assumption and therefore cannot identify a unique causal coefficient. Third, exact DCF–EVA equality in the full-year benchmark is a closure property of matched timing and assumptions; it does not by itself establish that the benchmark contains the most realistic forecast. Fourth, replicated workbooks strengthen traceability but do not enlarge the independent sample.

These limitations are not weaknesses concealed by the framework. They define the domain within which each result is valid. Structural credibility follows from declaring the boundary of inference, preserving the causal distinction between operations and valuation assumptions, and refusing to convert sensitivity outputs into empirical claims they were not designed to support.

Resolution of Anticipated Interpretive Questions

Why are 321.71 and 335.08 both present?

They do not value one unchanged regime. Approximately 321.71 is the broken-date EVA/Economic Value display under the baseline operating path and floating annual WACC with a 108/365 first-period stub. DCF under that regime is slightly lower, approximately 320.994–320.996 per share. The 335.0839 result belongs to the full-year, higher-sales, fixed-WACC, $g=7\%$ benchmark in which DCF and EVA close exactly.

Why are there two ROIC sequences?

The 82.13%–62.98% sequence belongs to the baseline-sales broken-date regime. The 93.26%–71.40% sequence belongs to the higher-sales operating path used in both the operating sensitivity and the full-year closure benchmark. They are not two computations of the same operating case.

Why can a stronger operating path produce a lower headline value elsewhere?

Because the comparison may also change the valuation date, WACC regime, terminal growth, or discount-period convention. A lower headline value does not prove weaker operations unless the valuation boundary is held constant. Operating attribution and valuation-regime attribution must be separated.

Are the two fixed-WACC cases versions of one answer?

No. The baseline-sales fixed-WACC sensitivity uses $g=10\%$ and produces approximately 376.25 per share. The full-year closure benchmark uses higher sales and $g=7\%$ and produces 335.0839 per share. Both use fixed WACC, but their operating and terminal assumptions differ materially.

Does the replicated implementation represent an additional scenario?

No. It has a distinct archived file identity, but its governing economics match the higher-sales broken-date sensitivity. It is evidence of provenance and implementation replication, not an additional independent observation.

Why does DCF not exactly equal EVA in the broken-date case but does in the full-year benchmark?

The broken-date architecture contains a stub-period discount convention and path/display differences that generate a temporal-closure diagnostic. The full-year fixed-hurdle architecture uses matched timing and assumptions in its DCF and EVA schedules, producing exact displayed closure. The broken-date gap must be explained, not hidden or forced to zero.

Which IRR should the reader trust?

Use only a measure whose regime, capital boundary, and timing convention are named. For the broken-date baseline, the controlled model-period IRRs are 77.5514% for total funds and 143.6988% for operating funds. Its displayed XIRR block contains a terminal-scale defect and cannot support inference. The full-year benchmark has separately validated annual-period IRR and actual-date XIRR values.

Does the highest valuation automatically deserve the highest RANK?

No. RANK evaluates the quality of the capital engine: operating efficiency, economic spread, EVA, cash conversion, reinvestment discipline, financing quality and valuation coherence. A favorable valuation assumption cannot substitute for superior operating economics.

Conclusion

The computational evidence does not present a set of competing estimates of one unchanged case. It presents a differentiated scenario architecture: an independent portability test, a broken-date baseline, a higher-sales operating sensitivity, a baseline-sales fixed-WACC growth sensitivity, a full-year higher-sales closure benchmark, and a replicated case retained for provenance. These regimes serve different analytical purposes and must be interpreted within the timing, operating, required-return, growth, claim, and scale boundaries that generated them.

Once this hierarchy is stated, the apparent contradictions become explainable. Different ROIC paths follow different operating assumptions; different WACC figures follow annual, fixed, floating, or stub-period conventions; different per-share values follow different timing and terminal assumptions; and different IRRs follow different capital boundaries and date conventions. The closed-loop standard is not that all scenarios must yield one number. It is that each scenario must close internally and that every cross-scenario difference must be attributed to a declared change.

Within the RANK framework, this chapter converts a potential source of reader confusion into a methodological strength. It demonstrates that structural ranking requires the same discipline demanded of the firms being ranked: clear boundaries, traceable causality, controlled assumptions, and explainable results. The scenario architecture therefore complements the book's central thesis rather than standing beside it: reliable rank is not merely a numerical ordering, but an attribution statement whose validity depends on the regime in which the result was produced.