

The Structural Completion of Factor Models

A Mathematical Note, with Worked Examples

Companion to A Risk Manifesto

Richard Bookstaber

August 31, 2026

Introduction

The factor model is the working language of institutional risk: a few systematic factors, and a residual treated as idiosyncratic. It is parsimonious, well understood, and for the broad, stationary co-movement of markets it does its job. But it carries limitations that surface when risk is largest, and there are tasks a risk manager needs that fall outside what it can express at all.

The limitations are familiar in symptom, less so in cause. The residual is assumed diagonal, yet correlated losses keep turning up in it. Factor loadings drift, and the drift is written off as estimation error or a change in regime. The model cannot represent a shock it has no history for. And it carries no internal signal of when its own numbers have stopped deserving trust—coverage that looks complete in sample can sit blind to the exposure that matters.

Beneath these sits a deeper one. A correlation records that two things moved together. It does not say why, and it cannot separate influence from common cause from coincidence (Pearl and Mackenzie 2018). Where shared inputs, supply dependencies, and regulatory linkages create real channels for a shock to travel, that silence is not a technicality. It means a factor resting on association cannot be shown to be wrong; it can only stop working, at which point another is fitted in its place and the zoo grows by one (López de Prado 2023; Harvey, Liu, and Zhu 2016).

There is a plainer objection that precedes the technical ones. The model does not know what the companies in it do. A firm enters as a return series, stripped of economic identity—what it makes, what it needs in order to make it, where it makes it, who buys it, what it cannot substitute. A cobalt miner and a beverage company are distinguishable only by the shape of their recent prices. Cleaner data and longer history do not repair this, because the loss happens at the point of representation: economic entities are entered as statistical objects, and the model then fails when they behave like companies.

This note begins from a different premise: that asset covariance is not a statistical object to be fitted but the image of a *propagation operator* built from observable structure—who

resembles whom, and who depends on whom. From that premise the factor model's behavior follows as consequence rather than assumption, and the consequences fall into two kinds. Some *repair* the model's blind spots: they show the residual is structured and computable, that loading drift is geometry rather than noise, that reliability is measurable, and that the biases buried in a correlation can be attributed and removed. Others *extend* the model to tasks it cannot presently perform: updating correlations for a shock never before seen, de-eventing a history whose episodes will not recur before fitting to it, and running the entire analysis from stated factor exposures with no holdings at all.

Concretely, the note establishes ten points. Each either overcomes a limitation of the factor model or lets it do something it currently cannot:

1. **The residual is information, not noise.** Its entries are fixed by structure and computable in advance, so the residual becomes a quantity to read rather than a term to diversify away. (§2)
2. **Latent factors can be named before they act.** The slow structural modes are genuine factors that standard methods miss: present in the network, and identifiable, before they activate in any return and before they generate any return information at all—where a principal-component fit can find them only after they have already moved. (§2)
3. **The structure fixes how many factors are genuine.** How many factors a market truly supports is not a modeling choice but a determinate property of its structure, computed in advance; it adjudicates whether the enterprise has been disciplined at all—below the count a model under-covers, above it the factor zoo begins. (§2)
4. **The valid domain of a factor model is delimited.** A factor model is faithful on one of three orthogonal mechanisms—the stationary background—and blind to the directional and bottleneck components; the split says exactly how much of a book each covers. (§3)
5. **Correlations can be adjusted for an unseen shock, and history de-evented before fitting.** The correlation matrix has a closed-form response to a named disturbance, so it can be repriced for an event with no precedent; and the events embedded in an estimation window—which will not recur in the form they took—can be projected out, leaving the recurring background a factor model can honestly fit. (§4)
6. **The model knows when to trust itself.** Computable gauges of asymmetry and speed mark when the factor reading is reliable and when it is degrading—and the degradation is itself the early warning. (§4)
7. **Correlations can be de-biased.** A raw correlation is separated into a direct link, a confounder artifact, and a collider artifact, with the direction of dependency distinguishing the last two. (§4)

8. **The blind spot is measurable and falsifiable.** How much of a portfolio's risk a factor span can see is a number, and the invisible remainder can be checked against realized losses. (§5)
9. **The analysis runs from loadings alone.** Given only stated factor exposures, with no holdings, the framework returns an implied portfolio and a bound on the hidden exposure the factor report cannot narrow. (§5)
10. **Loading drift is curvature, not noise.** The beta drift and regime change a desk re-estimates against is the holonomy of a curved manifold—a structural consequence, not estimation error—that no re-specification of the factors removes. (§6)

A single practical fact frames all ten. Although the note is written at the level of theory, nothing in it is merely theoretical: once the two structural matrices S and A are specified—together with a handful of scalars fixed once, not tuned to returns—every quantity that follows, the operators, the modes, the residual, the correlation response, the coverage, is a determinate computation on that data. The exposition is symbolic; the implementation is the same formulas evaluated on a given network.

The proofs behind these statements are not reproduced here. They are established in the mathematical appendix of *A Risk Manifesto* and developed in full in *Crisis Without History* (forthcoming); this note states each result, explains what it means for a factor model, and points to where it is proved. Section 1 summarizes the model; the six sections that follow take the ten points in turn, in the order listed above, and the parenthetical after each point marks the section that settles it.

1 The Model

Two matrices describe the structure of a market. The *similarity* matrix $S \in [0, 1]^{n \times n}$ is symmetric: S_{ij} measures how alike firms i and j are in what they make, use, and are exposed to. The *dependency* matrix $A \in \mathbb{R}_+^{n \times n}$ is directed: A_{ij} is the share of firm j 's critical inputs supplied by i , so a shock at a supplier reaches its customers and not the reverse. Both are read from public information—regulatory filings are the natural example, but supply-chain data, product and geographic classifications, and other standard sources serve as well. Neither is inferred from returns.

From these two matrices the model builds a single operator. The similarity Laplacian $L_S = D_S - S$ (with D_S the degree matrix) diffuses a shock across similar firms; the dependency generator $L_A = I - \eta A^\top$ carries it along supply links, and its inverse, the *propagation map*,

$$P_\eta = (I - \eta A^\top)^{-1} = \sum_{p \geq 0} \eta^p (A^\top)^p, \quad 0 < \eta < 1/\rho(A), \quad (1)$$

sums transmission over paths of every length, weighted by an attenuation η per hop. Combined co-exposure is collected in the *exposure operator* $E = P_\eta(\gamma S)P_\eta^\top$, and the two

channels are weighted into the *propagation operator*

$$M = \alpha L_S + \beta L_A. \quad (2)$$

The scalars $\alpha, \beta, \eta, \gamma$ are the model's only tunable quantities, and they are set once, by argument or a single calibration, not fit to a return series.

Because M mixes a symmetric channel with a directed one it is not symmetric, so its geometry is read through the symmetric part $M_s = \frac{1}{2}(M + M^\top)$, with orthonormal eigenpairs (μ_k, ϕ_k) ordered $\mu_1 \leq \dots \leq \mu_n$. *These eigenpairs are geometric, not statistical.* The eigenvalue μ_k is a dissipation rate along a propagation path—small μ_k marks a slow, far-reaching mode—and ϕ_k is a traversal pattern, a configuration of firms that stays coherent as a shock moves through the network. They belong to the differential geometry of flow on the structural manifold, in the sense that geodesics and propagation do, and they are computed from (S, A) before any return is observed. This must be kept distinct from the other spectrum that appears in this note—the eigenvectors of the return *covariance*, which are the statistical directions of realized co-movement that a factor model extracts. The two are related but not identical; much of what follows lives in the gap between the structural mode a shock *would* travel and the covariance direction that has *recently* moved.

Everything above meets the practitioner's object through one identity. A factor model asserts that the return covariance is low-rank plus diagonal,

$$\Sigma \approx B \Sigma_f B^\top + \Psi, \quad \Psi \text{ diagonal}, \quad (3)$$

with loadings $B \in \mathbb{R}^{n \times k}$, factor covariance Σ_f , and idiosyncratic Ψ . The structural view holds instead that the covariance is *generated*: shocks $u \sim (0, \Sigma_S)$ propagate to $y = P_\eta u$, so

$$\Sigma_y = P_\eta \Sigma_S P_\eta^\top + \sigma^2 I. \quad (4)$$

This is the Structural Covariance Decomposition, proved in the Manifesto appendix. Every result in this note is a statement about the gap between the assertion (3) and the generative identity (4).

Two computable numbers gauge how far the directed operator sits from a normal one, and therefore how far the geometric spectrum can be trusted as a summary:

$$\rho_c = \frac{\| [L_S, L_A] \|}{\| L_S \| \| L_A \|}, \quad \varepsilon = \frac{\| M_a \|_2}{\lambda_{\min}(M_s)}, \quad M_a = \frac{1}{2}(M - M^\top). \quad (5)$$

The condition $\varepsilon < 1$, the *weakly-asymmetric range*, is the standing assumption under which the approximations below carry controlled error; it is checked on the network, not assumed.

2 The Residual, the Latent Factors, and Their Number

A factor fit recovers the leading eigenspace of the return *covariance*—the statistical directions of §1, not the geometric modes. Writing that covariance in its own eigenbasis

$\Sigma_y = \sum_j \lambda_j v_j v_j^\top$ with $\lambda_1 \geq \dots \geq \lambda_n$, and letting $P_k = \sum_{j \leq k} v_j v_j^\top$ be the top- k projector a k -factor model spans, the residual covariance is what remains:

$$R = (I - P_k) \Sigma_y (I - P_k) = \sum_{j > k} \lambda_j v_j v_j^\top. \quad (6)$$

A factor model treats R as diagonal—idiosyncratic, diversifiable by counting names. Under (4) it cannot be. The generated covariance $P_\eta \Sigma_S P_\eta^\top$ has off-diagonal entries along every supply path, its (i, j) term summing $\eta^{p+q} [(A^\top)^p S A^q]_{ij}$ over routes, and no rank- k truncation removes them. The residual co-movement, measured in units of total volatility,

$$\rho_{ij} = \frac{R_{ij}}{\sqrt{[\Sigma_y]_{ii} [\Sigma_y]_{jj}}}, \quad (7)$$

is therefore nonzero entry by entry, and it is fixed by (S, A, η, σ) : one can write these numbers down before drawing a single return, and the realized residual converges to them as the sample grows.

This matters because it inverts the usual reading of the residual. What a factor model discards as noise is in fact the network’s fine structure—the exposure carried by specific supply routes and partial similarities the leading factors are too coarse to represent. The sign is not incidental: firms that co-load a removed common factor have that shared piece subtracted, so what survives in R tends to correlate negatively. And the structure cannot be absorbed by enlarging the model: for any $k < n$ the residual is a sum of the discarded eigenmodes, and it is exactly diagonal only if those modes are themselves coordinate directions, which generically they never are—so adding factors shrinks the residual but never diagonalizes it. The residual is a quantity to be read, and its map is computable in advance. *A Risk Manifesto* establishes that this residual is structurally informative; *Crisis Without History* carries out its term-by-term computation.

The same distinction between the two spectra explains where latent factors come from. Two alignments matter here, and they should not be conflated. Between the full directed operator and its symmetric core the alignment is a theorem: the modes of M sit within $\|M_a\|_2 / \min_{j \neq k} |\mu_k - \mu_j|$ of the modes of M_s (Davis–Kahan; the Manifesto appendix), an $O(\varepsilon)$ correction in the weakly-asymmetric range. Between the covariance directions v_k and the modes ϕ_k the relationship runs through the generative identity itself: Σ_y is computed from (S, A) , so its principal directions are structural objects—and the far-propagating patterns are the directions along which propagated variance accumulates, so the leading v_k track the slow ϕ_k closely on any given network, with the alignment itself a computable quantity rather than an assumption. A principal-component fit is, in effect, estimating structure. But a mode’s structural *presence* and its covariance *weight* are different things. Order a portfolio’s exposure two ways—by realized variance, and by structural persistence:

$$\text{Var}(w^\top r) = w^\top \Sigma_y w = \sum_j \lambda_j (w^\top v_j)^2, \quad I(w) = w^\top M_s^{-1} w = \sum_k \frac{(w^\top \phi_k)^2}{\mu_k}. \quad (8)$$

The first is dominated by directions that have recently moved; the second, weighting by $1/\mu_k$, is dominated by slow modes whether or not they have moved. A slow mode that has not yet been excited carries little realized variance but large persistence: invisible to the first sum, dominant in the second. That is a latent factor—present in (S, A) , waiting for a shock to project onto it, at which point its variance rises and a factor model “discovers” it after the fact. The geometry names it first. This is the practical content of the quiet-before-the-storm portfolio: calm by variance, heavily exposed by persistence.

These two readings—the residual from below, the latent modes from above—meet in a single number, and it is one a factor model has never had: how many genuine factors a market supports. That number is the effective dimension of the *systematic* covariance—the exposure operator, with the idiosyncratic floor removed,

$$E = P_\eta(\gamma S)P_\eta^\top, \quad D_{\text{eff}}(E) = \frac{(\sum_j \lambda_j)^2}{\sum_j \lambda_j^2}, \quad (9)$$

where λ_j are the eigenvalues of E . The formula is the ordinary effective rank a statistician would apply to a sample covariance—but here E is not estimated from returns; it is computed from (S, A) , and the $\sigma^2 I$ noise floor of (4) is gone. Same arithmetic, opposite epistemics: $D_{\text{eff}}(E)$ is a determinate structural quantity, read off in advance, with no history required. (The book’s spectral dimension $D_{\text{spec}} = (\sum_k \mu_k)^2 / \sum_k \mu_k^2$ on the propagation modes is the operator-level twin of this count.)

This is the number that adjudicates whether the enterprise has been disciplined at all. Below $D_{\text{eff}}(E)$ a model under-covers, leaving genuine structure in the residual; above it, further factors recover nothing structural and only fit sample noise—the factor zoo, given at last a stopping rule. And it explains the residual directly. A finite sample of length T detects only those genuine factors that clear the σ -driven noise floor; the rest—real in E , but sub-threshold in the sample—are what land in the residual. The residual entries above, then, are not noise but the genuine factors the sample was too short to resolve, and $D_{\text{eff}}(E)$ counts how many there truly were. Residual, latent factors, and the factor count are one fact seen three ways.

3 The Domain a Factor Model Can Represent

The gap between (3) and (4) is not diffuse; it partitions cleanly into three mechanisms. Built from the geometric spectrum of M_s and the tension between the two channels, three projectors

$$\Pi_{\text{bot}} = \sum_{k: \mu_k \leq \tau} \phi_k \phi_k^\top, \quad \Pi_{\text{ten}} \text{ (the tension block)}, \quad \Pi_{\text{diff}} = I - \Pi_{\text{bot}} - \Pi_{\text{ten}} \quad (10)$$

are mutually orthogonal and sum to the identity: Π_{bot} projects onto the slow, bottleneck modes; Π_{ten} onto the directions where the two channels fail to commute, carrying the

circulating, order-dependent part of a response, and built from the tension operator $[L_S, L_A]$ net of the bottleneck block; Π_{diff} onto the remaining diffusion bulk. Any propagated response $y = Ew$ then splits uniquely,

$$y = \Pi_{\text{diff}}y + \Pi_{\text{ten}}y + \Pi_{\text{bot}}y, \quad F_{\text{diff}} + F_{\text{ten}} + F_{\text{bot}} = 1, \quad (11)$$

with the mechanism fractions $F_{\bullet}(w) = \|\Pi_{\bullet}y\|^2 / \|y\|^2$ measuring how much of a given book's risk each mechanism carries. The orthogonality of the three projectors is proved in the Manifesto appendix.

A factor model is exact on one of these blocks and blind to the other two. On Π_{diff} the two channels commute, the induced terrain is flat, and a single constant loading matrix reproduces the response—this is the stationary “background” co-movement a factor model was built for. On Π_{ten} and Π_{bot} the channels do not commute and the slow modes concentrate; the terrain is curved (§6), and no constant loading matrix holds there. So $F_{\text{diff}}(w)$ is, up to leakage between blocks, the fraction of a portfolio's risk a factor model can represent faithfully, and $F_{\text{ten}}(w) + F_{\text{bot}}(w)$ is the fraction it must misattribute—read after the fact as instability rather than as the directional, path-dependent, chokepoint-bound structure it is.

The reading is a division of labor, and it is what makes the framework complementary to a factor stack rather than a replacement for it. The background belongs to the factors, and can be handed to them with confidence: it is the *recurring* co-movement—the co-movement that will still be there next year because it rests on standing shared exposures rather than on any particular disturbance. The event structure is exactly what they cannot see, and is where the structural reading earns its place. One caveat carries into the next section: the background is stationary only while the terrain evolves slowly on its own clock, a condition that section makes quantitative.

4 The Correlation Matrix: Adjustment, Trust, and De-biasing

Three operations on the correlation matrix follow from the generative identity, and none of them requires a return history of the event in question.

Adjusting for a shock never seen. A named disturbance reweights structural channels, $M(s) = M_0 + V(s)$, where $V(s)$ raises the weights on the affected dependency links—a cobalt disruption lifts the edges along the cobalt supply chain. Its effect on the whole correlation matrix is closed-form. Differentiating (4) and using $\partial(X^{-1}) = -X^{-1}(\partial X)X^{-1}$ on P_{η} ,

$$\frac{\partial P_{\eta}}{\partial s} = \eta P_{\eta} \frac{\partial A^{\top}}{\partial s} P_{\eta}, \quad \frac{\partial \Sigma_y}{\partial s} = \eta P_{\eta} \frac{\partial A^{\top}}{\partial s} P_{\eta} \Sigma_S P_{\eta}^{\top} + \left(\cdot \right)^{\top}. \quad (12)$$

Because $\partial A / \partial s$ touches only the shocked edges, the response is sparse: most pairs do not move, and the ones that do concentrate in Π_{ten} and Π_{bot} , a short and checkable list. The mode-level version is the first-order relation $\dot{\mu}_k = \langle \phi_k, \dot{M}_s \phi_k \rangle$ with excitation coefficients $a_k = \langle \phi_k, V(s) \phi_k \rangle$, and the linear reprice is trustworthy up to a remainder of order

$s^2\|V\|^2/\Delta_{\min}$, where $\Delta_{\min}$ is the spectral gap—small when the shock is mild or the gap wide, large when propagation dominates. The operational point is that a desk can reprice its factor correlations for an event the market has never experienced, because the reprice is computed from the structure the event acts on, not from a history of the event.

De-eventing history before fitting. A realized sample covariance $\hat{\Sigma}$ mixes the recurring background with the particular events that happened to fall inside the estimation window, and the problem with the mixture is not merely that it is unstable. The events are *non-recurring*: a correlation generated by one supply disruption, one funding squeeze, one export ban describes a configuration that held for the months that episode lasted and will not return in that form. A factor model fitted to the mixture therefore carries a fossil—a set of loadings calibrated to a world that no longer exists—and it carries it with the full confidence of a well-estimated parameter. The more eventful the window, the more precisely wrong the resulting factor correlations are, which is the opposite of the usual intuition that a dramatic sample is an informative one. The projectors separate the two: with $\hat{\Sigma}_{\text{bg}} = \Pi_{\text{diff}}\hat{\Sigma}\Pi_{\text{diff}}$ and $\hat{\Sigma}_{\text{ev}} = (I - \Pi_{\text{diff}})\hat{\Sigma}(I - \Pi_{\text{diff}})$,

$$\hat{\Sigma} = \hat{\Sigma}_{\text{bg}} + \hat{\Sigma}_{\text{ev}} + \underbrace{\Pi_{\text{diff}}\hat{\Sigma}(I - \Pi_{\text{diff}}) + (I - \Pi_{\text{diff}})\hat{\Sigma}\Pi_{\text{diff}}}_{\text{cross leakage}}, \quad (13)$$

where the two block terms are orthogonal in the trace inner product and the cross leakage is the $O(\varepsilon)$ -scale coupling between background and event structure—small in the weakly-asymmetric range, and computable. A factor model is correctly specified on $\hat{\Sigma}_{\text{bg}}$, where loadings are constant.

This is the disciplined form of something practitioners already do by hand. Dropping or down-weighting crisis windows is directionally right and structurally crude: it removes *periods*, when the object that needs removing is a *mechanism*. An event's footprint does not respect date boundaries—it begins before the window a desk would excise and decays after it, and meanwhile the recurring co-movement inside that window is discarded along with it. Projecting by mechanism removes the tension and bottleneck content wherever in the sample it falls and keeps the background throughout.

And the events are not thrown away; they are moved to where they can be handled properly. The background goes to the factor model, which fits it well. The event structure goes to the scenario response above, which reprices for the disturbance actually in question rather than the one that happened to land in the sample. De-eventing and shock-repricing are the same split run in opposite directions: strip the specific history out of the estimate, then put the specific scenario back in deliberately, by name.

Knowing when to trust the reading. Both operations are reliable only in the weakly-asymmetric, slowly-varying range, and the framework measures where it stands. Alongside the static gauges ρ_c, ε of §1, the adiabatic parameter

$$\varepsilon(t) = \frac{\|\dot{M}_s(t)\|_F}{\Delta_{\min}(t)^2} \quad (14)$$

compares how fast the structure is moving to the spectral gap it must move through. First-order tracking error is of order $\varepsilon(t)$, and at a gap minimum the probability that two modes exchange identity is $\exp(-\pi/(2\varepsilon))$. Factors are dependable while these gauges are small and degrade predictably as they rise—and because the gauges are computed from the structure itself, their rise is an early warning that the factor reading is about to become unreliable, delivered before the instability shows up in returns.

De-biasing the correlation. A single covariance entry conflates distinct causal routes, and the generative form separates them. Expanding $P_\eta = I + \eta A^\top + \eta^2 (A^\top)^2 + \dots$ in (4) with $\Sigma_S = \gamma S$,

$$\Sigma_y[i, j] = \underbrace{\gamma S_{ij}}_{\text{direct}} + \underbrace{\gamma \eta \sum_m [(A^\top)_{im} S_{mj} + S_{im} A_{mj}]}_{\text{one hop}} + \underbrace{\gamma \sum_k (P_\eta)_{ik} S_{kk} (P_\eta)_{jk}}_{\text{shared ancestor } k} + \dots \quad (15)$$

The last group is a *confounder*: two firms with $S_{ij} = 0$ and no direct link still show $\Sigma_y[i, j] > 0$ when a common ancestor k feeds both, and the correlation appears when that ancestor is disturbed and vanishes when it is calm—a spurious co-movement a factor model would absorb as a spurious factor. Its dual is a *collider*: conditioning on a common descendant k induces a correlation that is absent unconditionally, and it takes the explicit form of a Schur complement, the induced (i, j) term $-\Sigma_y[ik] \Sigma_y[kk]^{-1} \Sigma_y[kj]$ being nonzero exactly when both firms drive k . Direction is what tells the two apart— $i \leftarrow k \rightarrow j$ for the confounder, $i \rightarrow k \leftarrow j$ for the collider—and a symmetric correlation matrix, carrying only $A + A^\top$, cannot separate them. Working in Σ_y instead attributes each entry to its route, labels it, and in the confounder case names the responsible ancestor and anticipates when the correlation will appear or fade.

5 The Reach of a Factor Model, and Analysis from Loadings Alone

How much of a portfolio's risk a factor span can even see is itself a computable quantity. Scanning the unit sphere of possible shock directions and weighting each by the response it would produce, $\sigma_w(\omega) = |w^\top E^{1/2} \omega|$, gives a directional risk density ρ_w ; measuring how much of its mass lies in the cone the factor span F nearly covers, $\Omega_F(\varepsilon) = \{\omega : \|\Pi_F \omega\| \geq 1 - \varepsilon\}$, gives coverage and its complement,

$$\text{Cvg}_w(F) = \int_{\Omega_F(\varepsilon)} \rho_w(\omega) d\omega, \quad \text{Uncov}_w(F) = 1 - \text{Cvg}_w(F). \quad (16)$$

This is not a descriptive statistic but a falsifiable one: if realized losses concentrate in a direction ω^* outside the factor cone, then at least $1 - (1 - \varepsilon)^2$ of that shock's energy was provably invisible to F , whatever the model's in-sample fit. The essential caution is that coverage measures *visibility*, *not safety*: a portfolio can be fully covered and still carry the worst chokepoint concentration, because coverage speaks only to whether the factors span the shock, not to how large the exposure is. The same operator yields the effective number

of independent structural directions a book actually loads: distributing its risk over the principal directions (λ_j, ψ_j) of E ,

$$\pi_j(w) = \frac{\lambda_j (w^\top \psi_j)^2}{w^\top E w}, \quad N_{\text{eff}}(w) = \exp \left(- \sum_j \pi_j \log \pi_j \right), \quad (17)$$

counts how many structural channels carry the book's risk. Its shortfall against the nominal name count is the quantitative form of false diversification—many names, few independent structural bets—and a sharper measure than the number of factors a model happens to carry. Coverage and its falsifiability are established in the Manifesto appendix.

Because none of these objects needs a portfolio until the final step, the analysis can proceed from a factor report alone, with no holdings disclosed. Given stated factor exposures x and a set of reference portfolios $P = [w_1, \dots, w_K]$, the representative portfolio consistent with those exposures solves, in the exposure inner product $\langle u, v \rangle_{S+A} = u^\top E v$,

$$\min_c \|x - Pc\|_{S+A}^2 \implies (P^\top E P) c^* = P^\top E x, \quad \xi = x - Pc^*, \quad (18)$$

a unique solution when E is positive definite and P has full column rank. The implied portfolio $w = Pc^*$ carries the full structural analysis, and the unspanned residual ξ measures what the factor report leaves undetermined: the true portfolios consistent with the same loadings form a family over which structural exposure varies by at least $\|\xi\|_{S+A}$ —a dispersion the factor numbers alone cannot narrow. This is the honest onboarding path: from a set of loadings, a floor, a range, and a precise statement of why the report cannot say more.

6 Loading Drift as Curvature

The deepest of the factor model's limitations is geometric, and it is the one no amount of re-specification repairs. A loading matrix $B(x)$ assigns factor sensitivities at each structural configuration x ; as the structure evolves, B is carried along a path on the manifold $(\mathcal{M}, g)$ whose geometry the propagation operator induces. Transporting B around a closed loop γ does not return it unchanged: the mismatch is the holonomy of the connection, bounded by the curvature enclosed,

$$\|\tau_\gamma B(x) - B(x)\|_F \lesssim \left(\int_{\Sigma_\gamma} \|\mathcal{R}\| dA \right) \|B(x)\|_F, \quad (19)$$

with $\mathcal{R}$ the Riemann curvature and Σ_γ the surface γ bounds. When the two channels are in tension— $[L_S, L_A] \neq 0$ —and the economy's motion engages that tension, the curvature is nonzero on an open set, and the drift is forced. This is the precise sense in which factor-loading drift, the “beta drift” and “regime change” a desk observes and re-estimates against, is not estimation noise but a structural consequence of moving across curved terrain. The characterization of exactly when the terrain is curved—tension together with activation—is proved in the Manifesto appendix and developed in *Crisis Without History*.

The consequence for practice is blunt. Re-specifying the factor model—adding factors, redefining them, re-estimating loadings—is a change of coordinates on the same manifold, and curvature is a property of the manifold, invariant under change of coordinates. A better chart does not flatten curved ground. So no redefinition of the factor set removes the drift; it can only redistribute it, moving the instability from one loading to another while the enclosed curvature, and with it the holonomy, stays exactly the same. A factor model is a flat chart of curved territory, and the drift is the price of pretending otherwise—a price the structural reading at least makes visible and quantifies, rather than charging silently.

Taken together, the ten results move the factor model from a static statistical summary toward a structural instrument, without discarding anything a desk already runs. The same loadings are computed as before, but now read against the network that generates them: the residual becomes a map, the latent exposure becomes visible before it moves, the drift becomes geometry, the blind spot becomes a measured and falsifiable quantity, and the shock with no history becomes computable. What changes is not the arithmetic of the factor model but its status—from a description of what has co-moved to an image, now corrected and extended,

7 Coda: The Return Side

The ten results repair and extend a factor model’s treatment of *risk*. Portfolio design built on this note’s objects, though, is never purely about limiting risk or purely about seeking return—the same fitted quantities support both, and a single scenario response can be read either way: as a bound to respect or as a position to take. What follows develops the return-seeking reading with the same apparatus as the ten points, states precisely what it requires, and closes with the single constraint through which both readings run.

Exposure as the pooling variable. Let $\{e\}$ index a library of past structural episodes—named disruptions, each with its own $(S_e(t), A_e(t))$ evolving through the episode. For firm i in episode e at time t since onset, with s_e the firm at which the disruption originates, define the exposure path

$$x_i^e(t) = [E_e(t)]_{i,s_e} / [E_e(t)]_{s_e,s_e}, \quad (20)$$

the exposure operator of the model identity read as co-exposure to the episode’s source, computed from disclosed structure alone at each t —no return of firm i ’s enters its own exposure. Because E is generated, not fitted, $x_i^e(t)$ is available before firm i ’s return has moved at all, which is what makes exposure, rather than episode identity, the natural pooling variable: two firms in two unrelated episodes with the same x are, structurally, in the same position, whether or not either episode is the “most similar” one to the other by name. The disclosures this requires are not exotic: the SEC’s 10-K forms, for example, and other regulatory filings, along with the supply-chain and segment detail within them, are archived for decades, so reconstructing $(S_e(t), A_e(t))$ at a past point in time is a data-engineering exercise, not a conceptual gap.

The response is a path, not a scalar. Let $r_i^e(t)$ be firm i 's return residual, at time t in episode e , to a conventional factor model fitted independently of the structural framework—the $(I - P_k)$ projection of (6), realized rather than population. Cumulate it into the residual path since onset, $R_i^e(t) = \int_0^t r_i^e(s) ds$, and read off three functionals—the quantitative form of the peak-to-trough and trough-to-recovery columns the book's own episode catalog already carries for each named event, computed here per firm rather than by eye per episode:

$$D_i^e = \min_t R_i^e(t), \quad \tau_i^{e*} = \arg \min_t R_i^e(t), \quad \tau_i^{e,\text{rec}} = \inf\{t > \tau_i^{e*} : R_i^e(t) \geq 0\}. \quad (21)$$

Using the residual rather than the raw return is the return-side form of this note's opening claim (§2), that the residual is structured, not idiosyncratic: whatever a standard factor already prices is removed before the structural variable is asked to explain anything, so a nonzero relation between x and $(D, \tau^*, \tau^{\text{rec}})$ is a direct empirical test of that claim. A factor model enters here and nowhere else: it is the filter that produces $r_i^e(t)$, not a component of the construction that follows— E , M_s , the projectors, and M_s^{-1} are all computed from S and A alone, with no factor model in them, and the fit below would run identically against raw returns if the point were only to locate the response, not to test whether it survives what a conventional factor already explains.

The pooled response, as depth and timing with a band. Collect the panel $\{(x_i^e(t), D_i^e, \tau_i^{e*}, \tau_i^{e,\text{rec}})\}$ over all firms and episodes, and fit three functions of exposure together with their conditional dispersion,

$$D(x), \tau^*(x), \tau^{\text{rec}}(x), \quad q_\alpha(x) \leq \{D, \tau^*, \tau^{\text{rec}}\} \leq q_{1-\alpha}(x) \text{ at confidence } 1 - 2\alpha, \quad (22)$$

where $q_\alpha, q_{1-\alpha}$ are conditional quantiles of the same panel, not a symmetric variance band—depth and recovery time are not symmetric quantities. This band, evaluated at a firm's current exposure and plotted forward in elapsed time, is the risk envelope in its literal, historically calibrated form. The mechanism split of §3 refines (22) rather than replacing it: since $\Pi_{\text{diff}}, \Pi_{\text{ten}}, \Pi_{\text{bot}}$ decompose any $x_i^e(t)$ uniquely, each functional may be fit separately by block, since a factor model's blindness is mechanism-specific (§3) and there is no reason a bottleneck-driven and a diffusion-driven exposure should share one response shape.

Two structural benchmarks, not free-form curves. The spectral and manifold apparatus already in hand supplies a theoretical counterpart for each functional, so the empirical fit is checked against a prediction rather than left to float.

Recovery as travel time. A firm's exposure path is a point moving on $(\mathcal{M}, g)$, the manifold whose curvature governs loading drift. Chapter 10 of *Crisis Without History* gives the branch-level recovery timescale directly, as the reciprocal of the trajectory-averaged decay rate of the slowest engaged mode, $T_{\text{rec}}^{(\alpha)} \sim (\bar{\mu}_{\text{min}}^{(\alpha)})^{-1}$; the firm-level analogue is the structural null

$$\tau^{\text{rec}}(x) \propto \mu_{k(x)}^{-1}, \quad k(x) = \arg \min\{\mu_k : \phi_k \text{ engaged by } x\}, \quad (23)$$

the lifetime of the slowest mode the exposure direction x actually excites, a mode counting as engaged when it carries at least a stated fraction of the firm's largest mode weight. The null has two faces: across firms engaging different slowest modes, recovery times must order by lifetime; and when a single systemic mode engages every connected firm, the prediction is a *common clock*—recovery dates cluster, every tail decays at the same rate $\mu_{\min}$, and the residual spread in dates is carried by each firm's slow-mode share of its own trough, in first-passage form $\tau^{\text{rec}} \approx \tau^* + \mu_{\min}^{-1} \ln(s/\delta)$ with s that share and δ the recovery band.

Depth as persistence-weighted exposure. The trough is not set by exposure alone but by exposure weighted by how long each engaged mode holds its loading: while the disruption is applied, each mode carries amplitude in proportion to $1/\mu_k$, so the cumulative residual rides $\sum_k (\phi_k^\top v) \phi_k / \mu_k$ and the fit is written as theory plus a data-revealed remainder rather than a bare nonparametric curve,

$$D(x) = D_0(x; \theta) + g(x), \quad D_0(x; \theta) \propto -[M_s^{-1}v]_i, \quad (24)$$

the same $1/\mu$ persistence weighting that defines integrated exposure $I(w) = w^\top M^{-1}w$ in the filter bound below: the object that caps a portfolio's risk is, column by column, the object that benchmarks a firm's depth. The remainder g is nonparametric—the return-side name for exactly the gap Ch. 11 of the book proves must exist, the Geometric Obstruction Theorem, now isolated as the specific quantity left over once the structural benchmark is subtracted, rather than asserted whole: the Coda and the Theorem are the same claim, read from the return side and the risk side respectively.

The band widens where tracking degrades. §4 already gives the mode-exchange probability $\exp(-\pi/(2\varepsilon))$ at a gap minimum. The same $\varepsilon(t)$ should govern the width of the fitted band: $q_{1-\alpha}(x) - q_\alpha(x)$ is predicted to widen as $\varepsilon(t) \rightarrow \varepsilon_0$, since that is precisely where mode identity, and with it the exposure direction x itself, becomes unreliable. A band that fails to widen there is evidence against the construction, not merely an inconvenience.

The band as geodesic divergence. Elapsed time gives a second, dynamical route to the same width. Two firms with nearly identical exposure at onset sit on neighboring geodesics of $(\mathcal{M}, g)$, and their separation obeys the Jacobi equation that already governs scenario branching in Ch. 10 of *Crisis Without History*, at a rate set by the sectional curvature

$$K \sim - \frac{\|[L_S, L_A]\|^2}{\|\partial_S M\|^2 \|\partial_A M\|^2} \quad (25)$$

(Ch. 10, App.). Where $K \approx 0$ —diffusion exposure, channels not in conflict—two initially indistinguishable paths separate linearly and the band widens slowly. Where K is strongly negative—tension and bottleneck exposure, channels pulling against each other—separation is exponential: firms that looked identically exposed at onset can end an episode at opposite ends of the outcome distribution. The band is not a fixed-width tube around a mean path; its growth rate over elapsed time is itself the curvature already governing loading drift in §6, read here as dispersion instead of bias.

What identifies the fit. Two conditions, stated rather than assumed away. First, writing $\zeta_i^e(t)$ for firm i 's deviation from the fitted conditionals of (22), exposure must be exogenous to that deviation conditional on mechanism and elapsed time,

$$\zeta_i^e(t) \perp x_i^e(t) \mid \text{mechanism}, t, \quad (26)$$

which fails if x is computed with any knowledge, direct or retrospective, of which firms turned out to move—the panel-construction discipline this condition makes checkable rather than merely urged. Second, because every firm in episode e shares that episode's common shock, the $\zeta_i^e(t)$ are not independent across i within e ; standard errors on (22) must be clustered at the episode level, so the effective sample size for inference is the number of episodes, not the number of pooled firm-episode-time rows. A panel of forty episodes across a thousand firms is still, for inference, a panel of forty.

Execution. A live, unfinished scenario supplies its own exposure $x_i(t)$ by the identical computation of §1 and §4, with no return history of the scenario required—and the bound that limits risk on it is, read the other way, the bound that sizes an opportunity. At the mode level, the Structural Filter Theorem's cap (Ch. 10, App.),

$$\frac{C_{\text{bot}}(w)}{\mu_{\min}} + \frac{1 - C_{\text{bot}}(w)}{\mu_{\tau^+}} \leq \frac{I^*}{\|w\|^2} \implies I(w) = w^\top M^{-1} w \leq I^* + R(\varepsilon) \|w\|^2, \quad (27)$$

constrains w regardless of whether I^* is held fixed while risk is minimized or held fixed while $w^\top D(x)$ —the book's fitted response, with signed weights carrying the long or short direction—is maximized over the same feasible set: the cap does not change; only which side of it the optimization pushes against does. At the mechanism level, Ch. 10's coarse equalization $\|\Pi_{\text{diff}} w\| \approx \|\Pi_{\text{ten}} w\| \approx \|\Pi_{\text{bot}} w\|$ has the same dual: relaxed toward whichever block the fitted $D(x)$ favors, with the worst case held inside the empirical band $q_\alpha(x)$ rather than the theoretical I^* , it becomes a tilt instead of a parity condition. Four combinations follow—cap or maximize, at the mode level or the mechanism level, switched by the same $\varepsilon(t)$ clock Ch. 10 already uses to move between its two policies—and working them into an explicit portfolio rule is exactly the exercise Chapter 10 of *Crisis Without History* carries out in full; what matters here is that the same constraint does both jobs.

The full construction, its identification requirements, and its out-of-sample validation are developed in *Crisis Without History*, whose Ch. 10 appendix proves the Structural Filter Theorem invoked above and whose worked example in Appendix A.11 below carries the construction onto a concrete, reproducible panel.

Appendix A. Worked Example: The DRC Critical-Minerals Network

Every result in the body is a determinate computation once S and A are fixed. This appendix carries the ten points onto a single concrete network—the ten-firm critical-minerals example of *Crisis Without History*, Ch. 2—so that each abstract statement lands as a number. Four of the points are transcribed from the book’s own worked results (Ch. 2, Ch. 3, Ch. 12), cited to their source and re-verified against the matrices below to the printed precision. The remaining six are computed here directly from (S, A) and the stated parameters, with every convention pinned in the text and the internal checks reported—closed forms against finite differences, projector orthogonality to machine precision, fractions summing exactly to one. Every number on these pages is reproducible from the matrices printed in A.0.

A.0 The network

Ten firms span the cobalt and tantalum supply chains out of the Democratic Republic of Congo, with one deliberate control (Coca-Cola, unconnected). The directed dependency matrix A (A_{ij} = share of column firm j ’s critical inputs from row firm i) and the symmetric similarity matrix S are

	Gl	Ya	Um	CA	Cp	Te	Me	Lo	St	Co
$A =$	Glencore	0	0	.40	0	0	0	0	0	0
	Yageo	0	0	.15	0	0	0	.05	0	0
	Umicore	0	0	0	.35	.25	0	0	0	0
	CATL	0	0	0	0	0	.18	.10	.06	0
	Carpenter	0	0	0	0	0	0	.08	.12	0
	Tesla	0	0	0	0	0	0	0	0	0
	Medtronic	0	0	0	0	0	0	0	0	0
	Lockheed	0	0	0	0	0	0	0	0	0
	Stryker	0	0	0	0	0	0	0	0	0
	Coca-Cola	0	0	0	0	0	0	0	0	0

	Gl	Ya	Um	CA	Cp	Te	Me	Lo	St	Co
$S =$	Glencore	1	.65	.20	.10	.15	.05	.10	.05	0
	Yageo	.65	1	.15	.08	.10	.05	.12	.05	0
	Umicore	.20	.15	1	.25	.30	.10	.15	.10	0
	CATL	.10	.08	.25	1	.15	.40	.20	.25	0
	Carpenter	.15	.10	.30	.15	1	.10	.20	.25	0
	Tesla	.05	.05	.10	.40	.10	1	.35	.10	0
	Medtronic	.05	.05	.10	.20	.15	1	.20	.70	0
	Lockheed	.10	.12	.15	.25	.25	.35	.20	1	.20
	Stryker	.05	.05	.10	.15	.30	.10	.70	.20	1
	Coca-Cola	0	0	0	0	0	0	0	0	1

with parameters $\alpha = 0.6$, $\beta = 0.4$, $\gamma = 1$, $\eta = 0.9$. From these, $M = \alpha L_S + \beta L_A$ and every quantity below is determined. (Firm order: Glencore, Yageo, Umicore, CATL, Carpenter, Tesla, Medtronic, Lockheed, Stryker, Coca-Cola.)

A.1 The residual (point 1, §2) — from Ch. 12

Generating returns from the structural model ($r_t = P_\eta u_t + \sigma \epsilon_t$, $u_t \sim \mathcal{N}(0, S)$, $\sigma = 0.3$, $T = 250$, seed 42) and fitting three principal-component factors explains 58.6% of sample variance and leaves a residual that standard practice would call idiosyncratic. Computed from (S, A, η, σ) alone, the six largest residual correlations and two declared- zero controls are

Residual correlation	Estimated	Realized
Carpenter–Medtronic	−0.215	−0.207
Umicore–Yageo	−0.189	−0.221
Lockheed–CATL	−0.188	−0.149
Yageo–Carpenter	−0.168	−0.136
Carpenter–Stryker	−0.109	−0.071
Medtronic–Stryker	−0.104	−0.058
Coca-Cola–Tesla (<i>control</i>)	0.000	+0.109
Umicore–CATL (<i>control</i>)	+0.003	+0.018

Every estimated entry appears in the sample with the same sign and comparable magnitude; the controls, estimated at zero, carry only sampling noise. The residual a factor model discards is, entry by entry, a structural quantity.

One reproducibility note. The Estimated column is deterministic and verifies exactly against (S, A, η, σ) under the total-volatility normalization of §2. The Realized column is tied to the generating script’s random-number stream: an independent regeneration under the printed recipe (`numpy default_rng(42)`) yields the same signs and comparable magnitudes throughout—−0.188, −0.224, −0.140, −0.156, −0.152, −0.040, +0.187, +0.026, with three components explaining 58.4%—but not the identical draws; the original stream should be archived alongside the table.

A.2 Latent factors (point 2, §2) — from Ch. 3

The symmetric core M_s has ten canonical modes:

Mode	μ_k	Character
1	0.327	systemic: all nine connected firms; the slow, latent mode Coca-Cola alone, $\mu = \beta$ exactly (isolated)
2	0.400	
3–9	0.851–1.635	contrast modes (upstream/downstream, sector splits)
10	1.822	fastest, med-tech contrast

The slowest mode, $\mu_1 = 0.327$, is systemic and far-reaching; through the persistence weight $1/\mu_1$ it dominates integrated exposure $I(w) = \sum_k (w^\top \phi_k)^2 / \mu_k$ even when it contributes little to recent variance—a latent factor named by the geometry before a principal-component

fit would find it. Coca-Cola’s isolated mode sits at exactly β , a spectral fingerprint of a disconnected firm. (The full directed operator also carries a complex pair at $1.601 \pm 0.013i$, absent from any symmetric summary.)

A.3 When the reading can be trusted (point 6, §4) — from Ch. 2

The reliability gauges are computed directly on the network: the stability condition $\eta \lambda_{\max}(\frac{1}{2}(A + A^\top)) = 0.28 \leq 1$ holds, so $\lambda_{\min}(M_s) = 0.327 > 0$; the asymmetry ratio is $\varepsilon = 0.34 < 1$ and the normalized tension $\rho_c = 0.09$. The network sits inside the weakly-asymmetric range, so the spectral estimates above carry controlled error—the framework certifies its own reliability here, and would flag the network if it did not.

A.4 Curvature, the source of drift (point 10, §6) — from Ch. 2

Splitting M into symmetric and antisymmetric parts gives $\|M_{\text{sym}}\|_F = 4.10$ against $\|M_{\text{anti}}\|_F = 0.17$: the directed content is about 0.2% of the operator’s weight, the exact part a symmetric correlation summary would erase. The tension between the channels is nonzero and measured, $\|[L_S, L_A]\|_F = 0.43$, so the terrain is curved and loading drift on this network is a real, quantifiable phenomenon rather than estimation noise. (The drift itself is a holonomy—an integral of curvature along a path of structural change—so a full worked figure requires specifying that path; see the remaining computations.)

A.5 The genuine-factor count (point 3, §2) — computed

The exposure operator’s spectrum, computed from the matrices of A.0, is

$$\lambda(E) = (3.904, 1.794, 1.375, 1.000, 0.948, 0.846, 0.530, 0.380, 0.325, 0.276),$$

giving the structural count

$$D_{\text{eff}}(E) = \frac{(\sum_j \lambda_j)^2}{\sum_j \lambda_j^2} = 5.49, \quad D_{\text{eff}}(\Sigma_y) = 5.87.$$

The count is about six, and the spectrum shows the six by name: the fourth eigenvalue is *exactly* 1.000 with eigenvector $e_{\text{Coca-Cola}}$ —the isolated control, contributing nothing off-diagonal—so the six leading directions are five connected systematic patterns plus the control, and the spectrum’s second-largest relative gap falls precisely after mode six ($\lambda_6/\lambda_7 = 1.51$). The operator-level twin $D_{\text{spec}}(M_s) = 8.54$ is notably different, which is why the count is defined on E : the factor count lives in the generated covariance, not in the propagation operator’s flatter spectrum.

Against this stand the three detected factors of A.1: the sample’s three principal components carry 58.6% of variance (population top-three share 59.8%). The gap is then visible entry by entry. Reconstructing the residual from the two omitted *connected* systematic

modes alone (modes five and six) reproduces the table of A.1: Carpenter–Medtronic -0.222 against the full -0.215 ; Umicore–Yageo -0.145 against -0.189 ; Lockheed–CATL -0.247 against -0.188 ; Yageo–Carpenter -0.229 against -0.168 ; Carpenter–Stryker -0.078 against -0.109 ; only Medtronic–Stryker ($+0.048$ against -0.104) is carried instead by the fast med-tech tail. The two sub-threshold genuine modes carry the residual’s entire off-diagonal structure (their contribution exceeds it, at 123%, the small tail partially cancelling). Three fitted, about six genuine, and the residual is the gap between them—the three-versus-six question answered on the page. (For the alignment claim of §2: the leading covariance direction and the slowest structural mode align at $|\langle v_1, \phi_1 \rangle| = 0.985$ on this network.)

A.6 The mechanism split (point 4, §3) — computed

Conventions, stated in full: the bottleneck projector collects the modes with $\mu_k \leq \tau = 0.6$ (the systemic mode and the isolated control); the tension projector takes the left singular directions of $(I - \Pi_{\text{bot}})[L_S, L_A]$ with $\sigma_i \geq 0.2 \sigma_{\text{max}}$, which keeps four (singular values 0.283, 0.171, 0.103, 0.059; the next is 0.032); diffusion is the remainder (adopted convention—a gap-detection rule for τ and $\kappa = 0.25$ have also been used elsewhere and give the same rank on this network; not yet reconciled across scripts). The three projectors are mutually orthogonal to machine precision ($\sim 10^{-16}$) and the fractions sum to one exactly. For a Glencore shock the response splits

$$F_{\text{diff}} = 0.075, \quad F_{\text{ten}} = 0.299, \quad F_{\text{bot}} = 0.626 :$$

over 92% of that response is event structure—circulating or chokepoint-bound—that a constant-loading factor model misfiles. An equal-weight basket across the EV chain (Umicore, CATL, Tesla) is more extreme still: 0.046/0.120/0.835.

A.7 The correlation Jacobian, and history cleaned (point 5, §4) — computed

Take the cobalt scenario as a fractional cut of the Glencore→Umicore link ($\partial A_{\text{GL,Um}}/\partial s = -0.40$). The closed-form Jacobian of §4 agrees with a central finite difference to 1.6×10^{-10} , and its row norms rank the repricing: Umicore 0.685, Glencore 0.387, Yageo 0.251, CATL 0.247, Carpenter 0.202, Lockheed 0.075, Tesla 0.057, Stryker and Medtronic 0.040, Coca-Cola exactly 0. Only 10% of the Jacobian’s mass lies in the diffusion block—the repricing concentrates in the event structure, as §4 asserts. Cleaning runs on the same sample as A.1: splitting $\hat{\Sigma}$ by the projectors gives Frobenius-mass shares of 31.8% (background), 65.1% (event), and 3.1% cross leakage—the $O(\epsilon)$ coupling of §4, small and now measured—with the two block terms trace-orthogonal to machine precision.

A.8 De-biasing a confounded pair (point 7, §4) — computed

The CATL–Carpenter entry attributes by route: of the total 0.431, direct similarity contributes 0.150 (35%), Umicore-mediated terms 0.235 (54%), and deeper Glencore–Yageo

ancestry 0.047 (11%). The responsible ancestor is identified, not guessed: the shared-ancestor products $(P_\eta)_{ik}(P_\eta)_{jk}$ are 0.0709 for Umicore against 0.0092 for Glencore and 0.0013 for Yageo. A correlation of 0.431 between firms with 0.15 similarity is thus labeled: two-thirds artifact of common ancestry, concentrated at Umicore, and predicted to appear when that node is disturbed and fade when it is calm.

A.9 Coverage and the effective bet count (point 8, §5) — computed

Take the two-factor span $F = \{\text{market, EV chain}\}$. The tolerance-free reading is the direction-energy share $\|\Pi_F E^{1/2} w\|^2 / \|E^{1/2} w\|^2$: an equal-weight ten-name book scores 0.975 and the EV basket 0.972—their risk directions are visible to the span—while a medical book (Medtronic–Stryker) scores 0.408: 59% of its structural risk direction is invisible to a market-plus-EV factor model, before any return is drawn. (The cone integral of §5 gives the same ordering, 0.483/0.482/0.362 at $\varepsilon = 0.5$; in ten dimensions a small- ε cone is geometrically tiny—0.13% of the sphere at $\varepsilon = 0.1$ —so the tolerance must be read against the dimension.) The effective bet count of §5 delivers the sharpest single number in this appendix: the equal-weight ten-name portfolio has

$$N_{\text{eff}} = 1.25$$

against a nominal count of ten—ten names, one-and-a-quarter independent structural bets, nearly everything riding the systemic mode. The EV basket scores 2.12 (nominal three) and the medical pair 3.65 (nominal two): N_{eff} counts the structural channels a book loads, not the lines on its statement.

A.10 From loadings alone (point 9, §5) — computed

Suppose the only disclosure is a stated exposure of 30% broad market and 70% EV chain, and the reference set is three sector baskets: Materials (Glencore, Yageo, Umicore, Carpenter), EV (CATL, Tesla), and Medical (Medtronic, Stryker), equal-weighted. The normal equations of §5 give

$$c^* = (0.396, 0.578, 0.051),$$

an implied portfolio that overweights CATL and Tesla at 0.289 each with 0.099 across the materials names—and an unspanned residual $\|\tilde{\zeta}\|_{S+A} = 0.147$ against $\|x\|_{S+A} = 0.776$: 19% of the stated exposure is simply not determined by the loadings, the dispersion floor of §5 made concrete. The structural analysis then runs on the implied book exactly as if the holdings had been disclosed.

A.11 The return side: a synthetic episode panel (Coda) — computed

The Coda’s construction is exercised here on a panel simulated from the same network and parameters as A.0–A.10, with the generating process printed so the exercise is reproducible,

in the convention of Ch. 12. This is a synthetic panel, not a claim about the real return history of the named firms; it tests whether the pooled fit recovers what the structural model itself generated.

Generating process. A disruption at source firm s drives a severity path that ramps linearly to $s_{\max}$ over $T_1 = 6$ periods and decays exponentially over $T_2 = 1$ —faster than the slowest mode’s lifetime $1/\mu_1 = 3.06$, so that once the forcing resolves, the free tail belongs to the network, not to the disturbance. Each geometric mode k of M_s is forced by that path and relaxes at its own rate μ_k , $\dot{h}_k = -\mu_k h_k + \text{sev}(t)$ —the homogeneous part is not a modeling choice but the manuscript’s own relaxation dynamics $\dot{x} = -M_s x$ (App.), decomposed mode by mode; only the exogenous forcing $\text{sev}(t)$, standing in for the disruption’s evolution, is added here. The firm-level cumulative residual path is the resulting superposition onto the shock direction $v = E_{*,s}/E_{s,s}$,

$$R_i(t) = -\sum_k (\phi_k^\top v) \phi_k[i] h_k(t),$$

with i.i.d. Gaussian increments added before computing (21) to obtain a realized path from the noiseless predicted one. One operational convention: a relaxing path recrosses zero only asymptotically, so recovery is recorded when R_i re-enters ten percent of its own trough depth rather than at the strict zero of (21). No return data enters v or h_k ; only the noise added on top and the severity path forcing the modes are not read directly off the manuscript.

One episode, firm by firm. A cobalt-chain disruption at Glencore, $s_{\max} = 0.8$:

Firm	Mechanism	x_i	D pred.	D real.	τ^{rec} pred.	τ^{rec} real.
Glencore	diffusion	1.000	−1.057	−1.090	13.2	15.2
Yageo	diffusion	0.650	−0.860	−0.733	13.8	11.2
Umicore	tension	0.648	−0.823	−0.903	14.5	17.0
CATL	tension	0.304	−0.513	−0.511	15.5	30.0
Carpenter	diffusion	0.296	−0.520	−0.719	15.5	42.8
Tesla	diffusion	0.099	−0.354	−0.549	16.2	14.5
Medtronic	tension	0.077	−0.327	−0.239	16.5	10.5
Lockheed	diffusion	0.167	−0.413	−0.431	16.0	55.2
Stryker	diffusion	0.082	−0.332	−0.247	16.5	10.2
Coca-Cola (<i>control</i>)	isolated	0.000	−0.000	−0.029	—	—

Depth tracks the persistence-weighted benchmark of (24) almost exactly: across the nine engaged firms, the correlation between predicted depth and $-[M_s^{-1}v]_i$ is 0.998 (0.92 against the noisy realized depths). The isolated control’s noise-floor depth and undefined recovery are exactly what a firm with no path and no similarity to the source should show. Recovery dates cluster (13.2–16.5) despite an order-of-magnitude range in x —the common-clock face of (23), since on this network the systemic mode is engaged by every connected firm

(each carries it at no less than 76% of its largest mode weight, against a 50% engagement threshold). This clustering is a property of this ten-firm network’s single dominant slow mode, not a general law: a larger network with genuine near-partitions would carry several slow modes rather than one, and (23) then predicts several clocks—one recovery timescale per structurally distinct cluster—not a single shared date.

The structural null. The common clock is verified directly. Fitting each firm’s predicted tail decay over $t \in [15, 40]$ gives rates 0.327 ± 0.0002 across all nine engaged firms— $\mu_1 = 0.327$ recovered to three decimals, firm by firm: the recovery clock is the network’s slowest mode, not a property of any firm. The residual spread in dates then follows the first-passage form of (23): each firm’s recovery date correlates with its slow-mode share of trough at $\rho = 0.99$, and the analytic date $\tau^* + \mu_1^{-1} \ln(s/0.1)$ reproduces the predicted dates at $\rho = 0.99$. Timing is a property of the network; magnitude, through $-[M_s^{-1}v]_i$, is a property of the firm’s position in it.

Pooled panel. Forty synthetic episodes, each with its source drawn from the three upstream firms (Glencore, Yageo, Umicore) and its $s_{\max} \sim \text{Unif}(0.3, 1.0)$, pooled into 320 firm-episode observations:

Quantity	Value
Corr. (depth predicted, depth realized)	0.77
Corr. (depth benchmark $-s_{\max}[M_s^{-1}v]$, depth predicted)	0.99
Corr. (depth benchmark, depth realized)	0.76
Corr. (recovery predicted, recovery realized)	0.13
Share of pairs with a defined predicted recovery	1.00
Share with a defined <i>realized</i> recovery	0.79

Reading the panel. Depth survives noise well, and so does its structural benchmark: the persistence-weighted exposure $-s_{\max}[M_s^{-1}v]$, computed with no reference to any path, explains the realized depths essentially as well as the full noiseless simulation does (0.76 against 0.77)—the benchmark is not an approximation to the dynamics, it is the dynamics’ own summary. Recovery *timing* does not survive noise: predicted and realized recovery dates correlate at only 0.13, even though the noiseless clock is exact to three decimals. This is not a failure of the construction; it is the same asymmetry the book states as a design principle: the framework is precise about position and coarse about timing, and this panel shows why empirically as well as by argument—a firm’s *depth* is an instantaneous function of where it sits structurally, while its *recovery date* is a first-passage time, and first-passage times are far more sensitive to accumulated noise than level statistics are. A risk system built on this construction should trust the depth and band estimates more than the point recovery-date estimate, and should report the latter as the wide, mechanism-dependent range this panel shows it to be, not as a date.

Two scope notes on this panel. First, the correlations above describe this synthetic panel; they are not the clustered-inference estimates (26) requires of a real one, since with forty

episodes sharing sources among only three firms the panel does not have forty independent clusters to infer from—a real panel, with genuinely distinct episodes, would report standard errors clustered at the episode level rather than the plain correlations used here to check self-consistency. Second, this panel holds (S, A) fixed within each episode, so it exercises the depth and recovery-clock predictions but not the band-width prediction of §6: with no curvature evolving along the path, there is no geodesic divergence for a fixed network to generate, and that check is left for a panel in which the network itself moves during the episode.

The conventions above (the bottleneck threshold, the relative tension threshold, the scenario normalization, the cone tolerance, and the severity path $T_1=6, T_2=1$, noise scale, ten-percent recovery band, and fifty-percent mode-engagement threshold of A.11) are the only free choices made in this appendix, and each is stated where it is used; everything else follows from the matrices of A.0 by the formulas of the body.

References

- Harvey, Campbell R., Yan Liu, and Heqing Zhu. "...and the Cross-Section of Expected Returns." *Review of Financial Studies* 29, no. 1 (2016): 5–68.
- López de Prado, Marcos. *Causal Factor Investing*. Cambridge: Cambridge University Press, 2023.
- Pearl, Judea, and Dana Mackenzie. *The Book of Why*. New York: Basic Books, 2018.