

ESCRC Calculator

Quick Start Guide

Developed by SCIRISK · Updated September 2026 — v5.4.2

What these results are. Model-implied estimates under a stated, p-disclosed regional-dependence assumption, for decision support — screening and ranking suppliers and comparing scenarios — **not capital setting**. Parameter-uncertainty bands show that input uncertainty exists; band width is currently a judgment-based $\pm 30\%$ range on recovery time and impact coefficient, not calibrated evidence — see the FAQ below.

Two severity formulas. By default, each supplier's loss scales with its share of your COGS, pricing pro-rata revenue exposure — how much of your spend flows through a supplier — and **not** criticality, meaning how much of your output stops when that supplier stops. This can significantly understate loss at concentrated, low-spend suppliers. Declaring a supplier's criticality switches it onto two parameters you set separately: **criticality** (the share of your output that halts) and **unrecoverability** (how much of that you never make up). A cheap, low-spend part can still carry a high criticality.

Shipped presets use the criticality formula. Every preset supplier carries a criticality and an unrecoverability set by SCIRISK, so a preset is a worked example of that formula rather than an assessment of the named firm. **Your own networks stay on the pro-rata formula** until you supply a criticality yourself, and reproduce the same figures every time until you do. Check **formula_version** on any export: *ic-prorata* is the pro-rata formula, *kappa-beta-v1* the criticality formula. If a supplier's criticality cannot be researched the run is blocked rather than given a default, because a guessed criticality is indistinguishable from a researched one afterwards.

What is and isn't established. The criticality model is not independently validated: it has never been checked against a sealed set of disruptions it has not already seen, and unrecoverability cannot be checked against public reporting at all. Loss sizes are anchored to documented disruption episodes, but disruption *rates* — how often events happen — are not, because no dataset of site-years at risk exists to calibrate them against. Rates are bounded by documented history rather than fitted to it.

Capital setting is out of scope, and not merely pending. The modeled loss distribution is bounded by construction — a supplier cannot stop more than completely, or for longer than the horizon — so it cannot reproduce the extreme tail that large business-interruption claims show is real. That is a property of how the model is built, not a calibration gap further work closes. Treat every figure as decision support for screening and ranking.

A quick guide to running your first supply chain risk analysis.

1. Sign in or create an account

When you open the calculator, you'll be asked to sign in. If this is your first time, click "Create one" below the form, enter your email and a password (8+ characters), and you're in — no separate confirmation step. Your account is what your AI credits and saved networks are tied to; buy a plan once and save a network once, and both follow you on any device.

Sign-in protection. After several failed sign-in attempts, further tries are briefly rate-limited (per account and per IP) to protect your account; wait a few minutes and try again, or use password reset.

2. Start your network

You have three ways to get suppliers into the calculator:

- **Option A — Use a preset.** Apple, Toyota, TSMC, Pfizer, and Novartis come pre-loaded with hand-calibrated supplier networks. The HP (2011 Thailand hindcast) preset is a real pre-event supplier network you can load to see how the model's flags compare against what actually happened in the 2011 Thailand floods — an illustrative hindcast, not a calibration or a claim of predictive accuracy.
- **Option B — Add your own focal firm.** Click "+ Add Focal Firm", type a company name, and the AI proposes a realistic multi-tier supplier list (up to 15 suppliers across Tier 1, 2, and 3) to review and load. Uses AI credits — see "Plans & AI Credits."
- **Option C — Load a saved network.** Click the Save / Load badge (top-right) and load a previously saved network back exactly as you left it.

You can also click "+ Add Supplier" at any time to add a supplier manually — always free.

3. Fill in your focal firm numbers

Before running a simulation, set annual revenue and COGS for your focal firm (these convert supplier annual spend into focal-firm revenue exposure), and each supplier's tier, annual spend, and risk parameters (disruption rates, recovery times, impact coefficient).

Don't want to estimate every parameter by hand? Select suppliers and click "Research N suppliers" — the AI calibrates disruption rates, recovery times, and impact coefficients in a single batch, with a cited basis for every value. Every researched supplier is matched to an industry × geography calibration bucket whose disruption-rate bounds come from documented history; the AI must stay inside those bounds, the calculator independently enforces them, and the same validation runs at every entry point — manual edits, AI research, network load, and server-side on save — rejecting out-of-bounds values loudly rather than silently clamping them.

Where annual spend figures come from. When the AI proposes a network it does not invent dollar figures — it estimates each supplier's size as a share of spend: for a Tier-1 supplier, its share of your total COGS; for a Tier-2 or Tier-3 supplier, its share of the parent supplier's input cost, chained back through the parent links into your focal-firm dollars. Those shares are industry-structure estimates, not your actual procurement figures, and — unlike the risk parameters — they are not covered by the citation and verification tools. Presets declare their spend directly, and you can always type a supplier's share (percent of COGS, or percent of its parent) or edit the value yourself. Because your VaR and Expected Shortfall scale directly with these numbers, treat AI-proposed spend as a starting scaffold and replace it with your real annual spend figures wherever you have them.

4. Run the simulation

Click **Run Simulation**. The calculator runs a Monte Carlo simulation (100,000 replications by default, adjustable up to a 500,000 cap) and returns:

- Portfolio VaR and Expected Shortfall — model-implied estimates under a stated, p-disclosed regional-dependence assumption, for screening and ranking, not capital setting. They understate exposure concentrated in small-spend sole-source suppliers (see the known limitation above).
- A VaR confidence interval next to VaR — Monte Carlo sampling noise at your run size, not input uncertainty (see "bands" below, and the FAQ). It is an asymmetric order-statistic interval read straight off the sorted loss sample, so the upper and lower halves need not be equal.
- Per-supplier Expected Loss, standalone VaR, Tail Intensity, and RoESCR.
- A loss distribution chart, color-coded by how each outcome compares to your VaR threshold.

Regional dependence (p). The p button sets how strongly major events in the same geography move together (0 = independent, up to 0.95). Raising p can make a fixed-confidence VaR go *down* even as risk rises — that's expected, not a bug (see the FAQ) — so compare ES, not VaR, across different p. The bands panel also shows an ES readout at p = 0, 0.35, and 0.6 so you can see the sensitivity directly.

Parameter-uncertainty bands. The "bands" button re-simulates across each parameter's calibrated or disclosed range and reports a 90% band around VaR and ES — how much the answer would move if your input estimates were somewhat different. Recovery time and impact coefficient use a disclosed, judgment-based ±30% range (not calibrated evidence); disruption-rate ranges come from the calibration buckets. This is a different question from the VaR confidence interval above — see "What's the difference..." in the FAQ.

Tail Intensity. Standalone VaR ranks suppliers by dollar risk, so huge suppliers dominate even when smaller ones are far more hazard-exposed. Tail Intensity surfaces the "small but dangerous" ones: **the red flag goes to the top 3 by tail loss per dollar of spend**, which is a different question from the number in the column. A flagged supplier can therefore show fewer days than an unflagged one — the days say how costly a bad year is, the flag says how little you spend to be exposed to it. **Read the unit carefully: the "days" are loss-equivalent days — the worst-year loss expressed as equivalent days of the revenue flow it was priced against — not a forecast of how long a disruption lasts.** On the pro-rata formula that flow is the supplier's own, and losses are scaled by the impact coefficient (inventory buffers and alternative sources absorb most of a disrupted day), so a 35-day disruption at an impact coefficient of 0.14 reads as ≈ 4.9 loss-equivalent days. On the criticality formula it is the share of your output that supplier gates, so the figure measures how long a bad year costs you *per day of gated output* — driven by recovery time and unrecoverability. **Criticality itself cancels out of this number**, by design: how concentrated a supplier is, is what the criticality column already tells you, and Tail Intensity is there to add the hazard dimension rather than repeat it. Use it to rank, never to read off outage durations.

Reproducibility & export.

- Every result carries a reproducibility stamp (engine version + run seed) shown beside the results. Same inputs + same seed + same engine version reproduce the figures bit-for-bit.

- Pin a seed to reproduce an earlier run exactly, or let each run generate a fresh one.
- Export the full run record as JSON or CSV — focal firm, every supplier parameter, simulation settings, seed, engine and priors versions, portfolio VaR/CI/ES and per-supplier results — so any displayed figure can be verified outside the browser.
- **Raw samples.** A third download alongside JSON/CSV exports one row per Monte Carlo iteration — portfolio loss plus each top supplier's own loss — so you can recompute VaR/ES independently, or feed the exact sample into your own model.
- Each simulation run is also logged to your account's server-side audit trail.

5. Where the numbers come from

Every number is only as good as the parameters behind it. Each parameter carries one of five evidence tiers, weakest to strongest: default (not researched) → model estimate → cited (AI value with a named source) → verified citation (checked and supports the value) → validated database (from an authoritative external dataset). Where relevant the Sources panel offers FDA inspection and drug-shortage data and the INFORM Risk country multiplier; applied values always show what they replaced and cannot be applied twice.

Calibration bucket. Each researched supplier is matched to an industry × geography bucket by its product/name and country; the matching rule guards against false-positive matches (e.g. a display-panel supplier reading as pharma). If more than one bucket could apply, the supplier is flagged rather than silently guessed at, and you can set both the industry and the geography bucket by hand on the Parameters tab's **Bucket** columns.

Geography coverage. Geography spans regional buckets including South America, Africa, the Middle East, Oceania, and Russia/CIS. Each bucket carries a plain calibration status: **anchored** (bounds traced to named, dated events), **thin** (a placeholder awaiting stronger calibration), or **catch-all** (no regional match). A country string the calculator can't place is given a visible flag on the geography control rather than being routed silently, so you can correct it with the geography Bucket control.

Verifying citations. The “Verify sources” buttons re-check each researched citation and return a verdict: green check (verified), amber (unverifiable), red (contradicted); an unverifiable or contradicted parameter demotes that supplier to low data quality (applied server-side). You choose the verification engine:

- **Training mode (default)** — the model checks each citation against its own training knowledge, with no web search. Fast and inexpensive at 1 credit per supplier. It cannot confirm anything published or changed after the model's training cutoff, and returns “unverifiable” whenever it can't recall a source with real confidence.
- **Web search mode** — a live, web-search-enabled check that actually looks each citation up online. Slower (real searches, ~30–60s per batch) and priced higher to reflect the real search cost, at 3 credits per supplier. Use it to confirm recent or published sources that training mode can't reach. Batches are capped at 3 suppliers in this mode.

Database-sourced values (FDA, INFORM) are verified by construction — never billed and never demoted. Verification is billed per supplier actually sent (server-counted), so verifying two suppliers in web search mode costs 6 credits.

Evidence Coverage. One number summarizing the share of purchase-volume exposure sitting on cited-or-better evidence versus defaults, with verified parameters weighing more than merely-cited ones. It measures coverage, not correctness: a network can score 100% and still contain a wrong number — use the score to see where evidence is thin, not as a substitute for reviewing citations.

6. Save, compare, export

- **Save your network** (Save / Load badge): stores suppliers, focal-firm numbers, and settings to your account, available on any device. Up to 20 named networks; Load restores exactly, including all provenance and evidence tiers. Results aren't saved — only the network; re-running after loading is free.
- **Compare scenarios** (Compare badge): capture a result, tweak the network, run again, capture again. The table shows VaR, Expected Shortfall, total Expected Loss, and Evidence Coverage with deltas versus your baseline. Evidence Coverage is frozen at capture, and each captured scenario carries its engine version and seed, so any scenario is exactly reproducible via Restore.
- **Export for stakeholders:** export a report with full numbers and methodology notes, plus the JSON/CSV run record from Section 4 for independent verification.

7. Stress-test with Scenarios and the What-if builder

Scenarios (Network tab). Four ready-made stress templates — Hormuz Strait closure, Taiwan Strait disruption, Suez Canal blockage, Panama Canal drought — apply an elevated disruption rate and recovery time to the suppliers they actually match (by exact country, never a loose substring), plus a regional-correlation bump where relevant. Applying one automatically captures your current network as a “Baseline” and opens the Compare panel with the shocked result beside it, so you always see before-and-after, not just an after.

What-if tab — custom shocks. Build your own shock: scope it to a country, industry bucket, tier, or a hand-picked list of suppliers, then set a tariff, a risk overlay, or both. A tariff's direct duty is shown as its own annual-cost figure, separate from VaR and ES — a certain cost and a change in tail risk are different quantities, and merging them into one number would be misleading, not simplifying. The risk overlay (re-sourcing disruption, retaliation risk) needs a one-line rationale before it can move off “no change,” since no published data maps a tariff rate to a disruption-rate multiplier — that number is your judgment, recorded on the run. The comparison view also shows which supplier actually drives your tail risk, so a shock that barely moves VaR is explained rather than left looking broken.

Both paths use the same pinned-seed Compare machinery, so the VaR/ES delta you see is attributable to the shock, not to random sampling noise.

8. Explain my results & the board report

From any completed run, “Explain my results” gives you an executive view and a one-page, print-ready board report — both computed directly from your run's own numbers, no AI required and no credit cost. An optional AI narrative adds plain-language commentary on top, built only from the figures already on screen; switching it off changes nothing about the numbers themselves.

Every version of this view leads with a **provenance band** — **PRESET** (unmodified shipped defaults), **MODIFIED** (AI-researched or hand-edited, with a count), or **SCENARIO** (a stress template or custom shock is applied) — so a report never gets read as your firm's current measured position when it's actually a hypothetical. Supplier shares shown here are always shares of standalone risk, never of portfolio VaR, which isn't simply additive once regional dependence is switched on.

Plans & AI credits

Two things use AI, drawing from the same shared credit balance: Add Focal Firm (one lookup uses up to 15 credits), Batch AI Research (one credit per supplier selected), and Verify sources (1 credit per supplier in training mode, 3 in web search mode; database-sourced values are exempt). A counter in the top-right shows your balance. Manually adding suppliers, editing parameters, applying FDA/INFORM data, saving networks, capturing scenarios, and running simulations are always free.

Frequently asked

Q: Our biggest supply risk is a tiny sole-source part. Why doesn't the model see it?

A: It depends on whether that supplier carries a criticality. Check the scope line above your results, or **formula_version** on any export — the answer is different for the two cases.

If formula_version is ic-prorata (no supplier declared a criticality — your own networks, until you supply one): the limitation is live. Each supplier's loss scales with that supplier's share of your COGS, so the model prices **pro-rata revenue exposure** — how much of your spend flows through a supplier — and not **criticality**, meaning how much of your output stops when that supplier stops. A small, sole-source component and a commodity of the same annual spend produce the same tail, which can significantly understate the small supplier's real exposure. Use VaR and ES for what they were scoped for — screening and ranking suppliers, and comparing scenarios — and treat the figure for any supplier you know to be sole-sourced as a floor rather than an estimate.

If formula_version is kappa-beta-v1 (or kappa-beta-v1+timing — every shipped preset, and your own network once you supply a criticality): it does see it. Severity is driven by the declared criticality — the share of your output that halts — and not by share of spend, so a small sole-source part at criticality 0.7 produces a far larger tail than a same-spend commodity at 0.05, and the pro-rata understatement above does not apply to that supplier. What it depends on instead is the criticality and unrecoverability supplied: they are inputs, not model outputs, a criticality wrong by half moves the loss by roughly half, and unrecoverability could not be validated against public reporting at all.

Q: Tail Intensity shows a number of “days” that looks nothing like how long a supplier was actually down. Is it wrong?

A: It isn't downtime, so it should not match one. Tail Intensity is the supplier's worst-year loss at your chosen confidence level, divided by its daily revenue exposure: loss expressed in **equivalent days of revenue flow**, a ranking device rather

than a duration. It is also an unconditional annual screening figure — a “normal worst year,” not “given a Thailand-scale flood hits every supplier in the region at once.” That correlated, event-conditional scenario is closer to what raising p (regional dependence) models — see the next question.

What it measures depends on which formula the run used. On the pro-rata formula the impact coefficient scales the loss down — a 35-day disruption at 0.14 reads as ≈ 4.9 loss-equivalent days, far below the calendar outage. On the criticality formula it is days of the output that supplier gates, so it isolates recovery time and unrecoverability; the supplier's criticality cancels out, because that is what the criticality column itself already tells you. Two suppliers with the same recovery profile will therefore show the same Tail Intensity even when one gates seven times as much of your output — read the two columns together, not either alone.

Q: Why did raising p make VaR go DOWN? Is that a bug?

A: No — it's a real, provable property of correlation. Regional dependence makes suppliers in the same geography more likely to have their bad *and* good years together, pushing probability mass to both ends of the loss distribution — more quiet years and more co-disaster years. A fixed-confidence VaR can fall if the extra quiet years outweigh the fatter tail at that exact percentile; that's a known property of VaR itself, not a defect in this model. **Expected Shortfall never has this problem** — it is mathematically guaranteed not to decrease as p rises. Use ES, not VaR, when comparing runs at different p .

Q: What's the difference between the 90% parameter band, the confidence level, and the VaR confidence interval?

A: Three independent numbers that happen to share a screen. **Confidence level** (e.g. 95%) picks which loss quantile counts as VaR — a risk-tolerance setting. The **VaR confidence interval** measures Monte Carlo sampling noise — how much the reported VaR would move across random seeds at your sample size; it shrinks toward zero as N grows and says nothing about whether your inputs are right. The **90% parameter band** measures input uncertainty — how much VaR/ES would change if your disruption-rate, recovery-time, or impact-coefficient estimates were somewhat different; it does not shrink with N , and only tightens with better evidence.

Q: Can I use the calculator on my phone or tablet?

A: Yes — the whole app is optimized for phone and tablet screens: panels stack into a single column, tap targets meet comfortable minimums, and inputs don't trigger iOS zoom-on-focus. The status badges (AI credits, saved networks, scenarios, provenance, evidence coverage) and the run-record stamp (engine version, seed, export) collapse behind a small toggle so they don't crowd a small screen. The laptop/desktop layout is unchanged, and every displayed number is identical across devices.

Q: What's the difference between a Scenario and the What-if custom shock builder?

A: Scenarios are four ready-made templates modeling named, live geopolitical events — pick one and apply it in one click. The What-if builder is for shocks the templates don't cover: you choose exactly who is affected and by how much, and it's the only place tariffs are modeled, with the duty cost kept separate from risk. Reach for a Scenario first; reach for What-if when you need a shock templates don't cover, or you need the tariff-cost accounting.

Q: Why is a tariff's cost shown separately instead of inside VaR?

A: A tariff you actually pay is a certain cost, known once the rate is set — not a probability of loss. Adding a certain cost into a risk figure like VaR would shift the number by exactly that amount without saying anything new about risk, and on realistic tariff scenarios that constant can dwarf the genuine tail-risk figure it would be hiding. Keeping the two separate — an annual cost line, and a risk delta from the (optional) disruption overlay — is what lets you read each for what it actually is.

A note on AI-assisted data and scope

Anything labeled “AI” is a starting point, not a verified fact: calibrated priors bound what the AI can claim, verification checks its citations, and validated databases override where authoritative data exists — but review the citation shown for each value. Treat every result as a model-implied estimate under a stated, p -disclosed regional-dependence assumption, for decision support and not capital setting, and not a substitute for expert judgment. Where no criticality is declared, the model prices a supplier's pro-rata revenue exposure rather than its criticality and understates concentrated sole-source risk; where criticality is declared — which includes every shipped preset — it prices criticality instead, and the level then rests on the criticality and unrecoverability supplied, by you or by SCIRISK on a preset. On either path the model prices availability loss only — not input-price or pass-through risk — its simulated tail is bounded against the tail estimated from disclosed disruption losses, and its event frequencies are literature-anchored rather than derived from an exposure base, so the ordering of suppliers is more dependable than the level. Parameter-uncertainty bands show that input uncertainty exists, not a calibrated confidence statement: band width on recovery time and impact coefficient is currently a disclosed,

judgment-based $\pm 30\%$ range, not calibrated evidence. The VaR confidence interval and the band estimator have both been independently confirmed by SCIRISK's model validator; see the Methodology Specification for the full validation record, and the FAQ above for what each interval does and doesn't tell you.

Questions or stuck on something? Reach out to the SCIRISK team — we're happy to help.