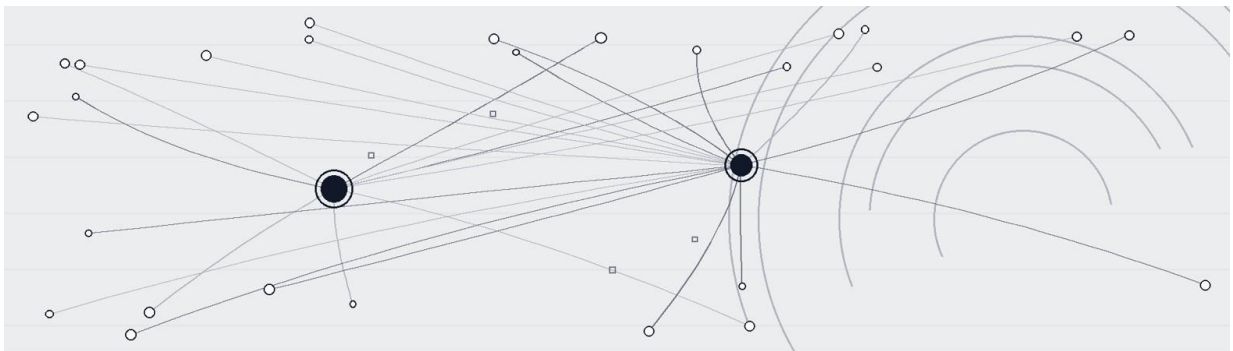


AUGUST 2026
APPLIED QUANTUM

THE APPLIED QUANTUM CCF CONCENTRATION SENSITIVITY MODEL

CONCEPT NOTE



Optional impairment-propagation analysis consuming measured framework outputs

Version 1.0-RC – August 2026

Steve Vaile and Marin Ivezic, Applied Quantum

Disclosure: Applied Quantum provides post-quantum migration advisory services. This work is published openly under CC BY 4.0 and is computable without engaging Applied Quantum.

COPYRIGHT AND LICENSE

© 2026 Steve Vaile and Marin Ivezic / Applied Quantum. This work is licensed under Creative Commons Attribution 4.0 International (CC BY 4.0). You are free to share and adapt this material for any purpose, including commercial use, provided you give appropriate credit to Steve Vaile, Marin Ivezic and Applied Quantum, link to the licence, and indicate any changes made.

Licence details: <https://creativecommons.org/licenses/by/4.0/>

The reference implementation and canonical test vectors publish separately under Apache-2.0 on the Applied Quantum GitHub. Neither licence reads onto the other.

DISCLAIMER

This document is provided as professional guidance based on the authors' and Applied Quantum's practitioner experience and publicly available standards, regulations and industry research. It does not constitute legal, regulatory, financial or compliance advice. Organisations should consult qualified legal counsel, auditors and regulatory authorities for guidance specific to their jurisdiction, sector and circumstances.

The regulatory instruments, certification records, standards and vendor capabilities referenced in this document reflect the state of the landscape as of August 2026, and that landscape moves. Readers should verify current status against primary sources (regulators and supervisory authorities, NIST, EMVCo and PCI SSC, the CycloneDX specification, public certification listings, vendor documentation) before making decisions.

References to specific vendors, products or tools are for illustrative purposes and do not constitute endorsement.

This is a v1.0 release candidate: v1.0 final publishes in November 2026 after a pilot cycle. Open items are public on the roadmap at CCFramework.org, and corrections are published as dated entries on its errata page.

ABOUT THIS DOCUMENT

Document type	Concept note
Parent document	The Cryptographic Concentration Framework – Universal – v1.0-RC (August 2026)
Intended audience	Quantitative risk and network-analysis teams evaluating impairment propagation
Assumed knowledge	The Universal Framework's outputs; comfort with network models
Scope	Optional impairment-propagation analysis consuming measured framework outputs. The framework depends on none of it, and the note's open questions are stated and worked publicly.

Status	Version 1.0-RC – release candidate; v1.0 final November 2026 after pilot
---------------	--

VERSION HISTORY

Version	Date	Changes
1.0-RC	August 2026	First public release, as a release candidate; drafted through four external review rounds and a family-wide line-edit pass.
1.0-RC	20 August 2026	Co-author review edits accepted (Steve Vaile, 19–20 August).
1.0-RC	7 September 2026	Heading aligned to the concept-note labelling ruled at RC; changelog added; docx rendering corrected.
1.0-RC	8 September 2026	Prose revised throughout; the trajectory dependence of a threshold recorded as open question 4. No change to the model’s inputs, outputs or method.

HOW TO USE THIS DOCUMENT

This concept note is optional and separate. It consumes measured framework outputs and proposes an impairment-propagation analysis on top of them. The framework depends on none of it.

ACCOMPANYING RESOURCES

The framework family publishes at CCFramework.org: the Universal Framework, the Payments and Financial Services extensions, the payments whitepaper, the CBOM Conformance Statement, the Technical Companion and the instruments, together with the public roadmap, the errata and corrections page, the provenance survey and the frontier-census methodology.

Every canonical figure in the family comes from the reference implementation, with ten canonical test vectors that fix each one exactly. Both are open under Apache-2.0 on the Applied Quantum GitHub, and any implementation that reproduces the fixtures is conformant.

The Applied Quantum CBOM Profile, the input data contract, publishes at CBOMProfile.org. The Applied Quantum PQC Migration Framework, whose sector taxonomy the extensions follow, publishes at PQCFramework.org. Practitioner background and adjacent analysis sits on PostQuantum.com, as further reading rather than a normative source.

TABLE OF CONTENTS

1. What this model is for.....	5
2. What it is not.....	6
3. What the model assumes	6
4. Inputs.....	7
5. Outputs.....	9
Tolerance threshold.....	9
Propagation ordering	9
Coverage sensitivity	9
6. Method	10
6.1 Evidence requirements.....	11
6.2 Experiment records	12
6.3 Uncertainty and sensitivity analysis.....	13
6.4 Validation and safeguards.....	13
7. Governance.....	14
8. Publication.....	15
Resolved and open at v1.0-RC	15

This document specifies structure and inputs. It is a specification only once the runnable model, formal equations, parameter distributions, validation cases and synthetic dataset publish with it; until then, it asserts structure, not results.

An open model that takes measured CCF outputs and returns the concentration threshold at which an institution's impact tolerance breaks.

Companion to the [CCF Universal Framework v1.0-RC](#). No CCF metric, determination or output requires this model, and an institution can complete a CCF assessment without running it.

1. WHAT THIS MODEL IS FOR

A CCF assessment establishes which cryptographic dependencies support a service, how concentrated those dependencies are, how far a shared failure domain may extend, and whether the resulting service consequences would exceed an approved tolerance. It does not calculate the subsequent response of a network of institutions. A payment processor's impairment may delay its clients, require a counterparty to reroute traffic, consume operational capacity elsewhere, or prevent a settlement instruction from reaching its destination. Those effects depend on relationships and operating conditions that a concentration index does not contain.

Consider a service scoring 8,400 at layer 4. It is effectively single-source, and its failure-impact determination records that it fails its impact tolerance. The chief risk officer's next questions are how many counterparties fail with that layer, how many more are impaired, and at what concentration the tolerance breaks. CCF does not answer those network questions.

The sensitivity model propagates a single-dependency failure through an interbank or payment network and returns propagation ordering, sensitivity and thresholds. Its intended use is a controlled comparison: given a specified dependency failure, a documented network and a set of operating assumptions, how do service impairment and recovery change when the dependency distribution or a relevant mitigation changes? This can support scenario design, prioritisation of evidence collection and examination of alternative architectures. It does not establish a universal concentration limit for financial institutions.

This distinction matters because identical CCI values can describe different networks. Two institutions may each have two equally weighted implementation families, giving a layer-2 CCI of 5,000. In one case the families are independent; in another, both depend on the same upstream code. Their indices match, but their reach differs. Even with identical reach, one institution may have a tested substitute and the other none. A

sensitivity model must consume those underlying relationships and service facts rather than reconstruct them from the headline index.

2. WHAT IT IS NOT

Loss estimates. The model produces no monetary output at any scale. An earlier version of the analysis applied an assumed loss rate to published GDP and reported dollar ranges. Assessment of that analysis centred on the loss rate and left the structure aside. That output has been removed and will not return.

Forecasting. There is no Q-Day to observe, so the model cannot be back-tested against such an event. It is a structural sensitivity tool, not a forecast.

Optional use. An institution can compute CCF, act on it and never run this model. Completing a CCF assessment, satisfying a CCF input conformance level, or using the Board Report does not require it. The Universal Framework's failure-impact determination remains the applicable method for assessing the consequences of the measured dependency domains. The model may provide evidence for that determination; it must not replace an operational fact with an unlabelled simulation assumption.

The output is operational and structural. It does not estimate the arrival date of a cryptographically relevant quantum computer, rank vendors' security quality, or assign a probability of compromise from an evidence tier.

3. WHAT THE MODEL ASSUMES

The framework supplies measured inputs. An institution that has completed a CCF assessment has measured concentration per layer and coverage per service from its own inventory. For network structure, the model assumes only the topology. This is one stated, adjustable assumption, as required by the model risk policy governing the tool. Propagation dynamics and unobserved edge weights are separately parameterised under Section 6.

The experiment records which inputs are observed, estimated, elicited from service owners, or chosen solely for sensitivity testing. The following distinctions preserve the boundary between measured CCF inputs and model assumptions.

Network structure and weights. A connection can represent a message route, a processing dependency, a settlement obligation or another relationship. Those meanings are not interchangeable. The model must state what each edge means, its direction, and the unit of its weight. A binary connection alone does not reveal traffic volume, available capacity, financial exposure or the consequences of losing the connection.

Initial failure and affected membership. A named failure scenario determines which dependency is impaired and what the impairment means. A coding defect might expose integrity without stopping processing; a module withdrawal may remove capacity immediately. For uncertain ancestry or design membership, the model should run distinct feasible membership scenarios. It must not treat all non-excludable memberships as established, or assign them equal probabilities merely because a Monte Carlo routine requires a distribution.

Transmission to services. The relationship between affected cryptographic assets and service impairment must be explicit. Losing 20% of an asset population need not remove 20% of throughput. The affected assets may include the only settlement authorisation path, or they may support capacity with a tested substitute. The failure-impact determinations provide the service-specific starting point for this mapping.

Response and recovery. Rerouting, isolation, suspension of a counterparty, manual processing and emergency key replacement depend on operating policies and available capacity. Their activation times, limits and dependencies require evidence or labelled assumptions. A model without a time dimension cannot produce a meaningful estimate of whether a two-hour tolerance is exceeded.

Joint state. Population migration percentages and marginal dependency shares do not identify which participants share the same exposure. Where joint membership is unavailable, the experiment should use feasible bounds or explicit alternative assignments. Independence is one possible assumption to test, not a default established by the absence of data.

These assumptions should be adjustable separately. Otherwise a result attributed to diversification may actually arise from a changed routing policy or an additional recovery resource introduced at the same time.

4. INPUTS

From the CCF assessment, measured:

Input	Source
CCI per layer per service	Universal Framework C.3
Largest share per layer	Universal Framework C.4
Flagged-only index per layer	Universal Framework C.2
Effective Coverage per service	Universal Framework C.6

Input	Source
Counterparty counts by class	Universal Framework C.6
Impact tolerance per service	Institution's operational resilience framework
Discovery coverage and conformance level	CCF CBOM Conformance Statement §9

Parameterised by the institution:

Input	Default	Notes
Network topology	Stylised core-periphery	Replaceable with observed exposures where available
Counterparty dependency map	Derived from <code>appliedquantum:path:counterpartyRef</code> on the path records	Which services depend on which counterparties
Dependency identity register	The assessment workbook's resolved dependency list per layer per service	The failure that propagates must be identifiable; aggregate scores cannot say which nodes share it
Dependency reach matrix	Derived from the workbook: dependency by service, share carried	Two networks with identical indices can have radically different propagation
Impairment threshold	Service-defined	The point at which a service is treated as impaired
Recovery profile	None	Optional, for tolerance-breach timing

Assumed, and labelled as such in every output: the topology where observed exposures are unavailable, and the propagation dynamics.

5. OUTPUTS

Tolerance threshold

The primary output is the tolerance threshold for each layer of each service: the concentration level at which a single-dependency failure at that layer produces impairment exceeding the impact tolerance. Each threshold is reported as a CCI value and as a largest share, matching the layer-specific failure-impact determination in Universal C.5.

This output turns the static concentration score into an institution-specific target. An institution scoring 8,400 against its computed tolerance threshold knows that it fails and by how much it must move. The threshold is an output of the model run on that institution's data, not a constant of the framework.

Propagation ordering

The secondary output ranks which layer's failure reaches furthest. It identifies where assurance effort has the greatest effect, the question the earlier blast-radius-by-depth analysis addressed qualitatively.

The model may compare initiating failures to identify which scenarios affect the greatest service population or create the longest disruption. The ranking is conditional on the chosen scenario severity and response assumptions. A root-withdrawal scenario and a retroactive key-generation compromise are not directly comparable merely because both produce an affected-asset percentage. Reports should separate confidentiality, integrity and availability consequences where their measures differ.

Coverage sensitivity

The tertiary output shows how impairment changes as Effective Coverage rises while concentration is held constant. Raising coverage across the institution's own estate frequently leaves impairment materially unchanged until its counterparties raise theirs.

A coverage experiment must therefore identify the cryptographic function, target definition and paths changed. Raising a single service-level percentage without identifying the changed population is insufficient to determine a network effect. Historical confidentiality exposure also requires separate treatment: protecting future connections does not repair previously captured traffic.

Every output carries the discovery coverage and conformance level of the input data, the topology source, and the statement that no output is a loss estimate. The report identifies the initiating domain, the affected membership scenario and the controlling assumptions beside the result.

6. METHOD

The model uses two established techniques to distinguish hard default from distress.

Eisenberg – Noe clearing. The clearing model counts hard defaults. That count is expected to be smaller than the distress count because netting, central clearing and prefunded collateral made direct default cascades rarer after 2008. Hard defaults are reported beside distress, not instead of it.

The clearing framework concerns payments in a network of financial obligations. Its use requires operational impairment to be translated into constraints on available resources or payments. A failed cryptographic service must not simply be labelled an insolvent institution, and an availability interruption must not automatically become a loss on every financial obligation.

DebtRank. DebtRank counts distress: institutions impaired enough to stop transacting without failing. A cryptographic failure impairs through quarantine and withdrawal, so runs lead with distress. This is a modelling choice whose validation cases remain open.

Its use requires a defensible mapping from cryptographic impairment to the state variable propagated by the chosen formulation. Different DebtRank variants and update rules may behave differently; the implementation must identify the precise formulation rather than cite the model family alone. Which network-model family ships first, and the validation cases for the distress-versus-default ordering, remain open.

Monte Carlo analysis. The model runs Monte Carlo analysis over the parameterised uncertainty and reports distributions rather than point values.

Effective Coverage and mitigating variables. Effective Coverage is a post-quantum migration-state variable, not a universal mitigant. It enters only scenarios whose failure mode migration addresses: confidentiality loss and forgery at layer 1, and downgrade at layer 5. A classical shared-code defect, a firmware recall, a trust-anchor compromise or key material already generated by a defective design is not mitigated by coverage. Before a run, each scenario declares its mitigating variables and whether each reduces probability, reach, consequence or recovery time.

Edge weights. Eisenberg – Noe and DebtRank require relative exposure between nodes; connectivity alone does not provide it. Use observed exposures, such as settlement values or credit lines, where they exist. Otherwise draw the weight matrix from a stated distribution conditional on node class and topology, and sweep that distribution through the Monte Carlo analysis. Label the weighting assumption in every output alongside the topology. Outputs describe distributions over those weightings, not properties of one assumed network. Fixing weights inside a stylised topology makes two assumptions while stating only one.

Migration-process limitation. Migration is treated as a state per asset, not as a process. The model therefore understates residual and implementation risk during migration.

6.1 Evidence requirements

The assessment's published heat map is a useful starting point, but it is not a sufficient model input. The proposed model requires the underlying dependency register, reach relationships and service mapping. It also requires a distinction between observed facts, assessment determinations and scenario assumptions.

Input	Role in the experiment	Source and limitation
Service and asset population	Defines what can be affected and the denominator for each output	CCF scope register and canonical identities; retain incomplete-population qualifications
Executing-unit assignments	Identifies the dependencies used in each layer	Assessment supplement; a service maximum cannot substitute for these assignments
Reach graph and membership states	Identifies shared ancestors, designs and other failure domains	Evidenced and possible memberships remain distinct; an upper bound is not an occurrence probability
Criticality and supported functions	Connects affected assets to service consequences	Service-specific determinations, including the five CCF tests
Effective Coverage and its supporting path state	Records whether the functions addressed by the scenario meet the approved target	Require own-state, counterparty, hop and population evidence; do not infer missing relationships from a percentage
Service tolerances and consequence thresholds	Defines the conditions that constitute an unacceptable outcome	Approved institutional statements; an absent threshold is not a permissive threshold
Operational substitution and recovery	Describes the response after the dependency fails	Tested or estimated capacity, activation conditions and

Input	Role in the experiment	Source and limitation
		recovery times, with evidence and uncertainty
Counterparty relationships	Connects the institution's services to other participants	Path records and the assessment supplement; a legal-entity identifier is not itself a payment-flow weight

The CBOM Profile supplies inventory attributes and identifiers. It does not supply the complete counterparty network, a payment-obligation matrix, a liquidity model or a behavioural response function. These would have to be obtained separately or declared as synthetic inputs.

A model manifest should identify the exact assessment, framework, Profile and implementation revisions used. It should preserve the CCF and EC conformance claims, discovery coverage, evidence limitations and the date of each operational estimate. Updating a network assumption must not silently update the recorded inventory.

6.2 Experiment records

An experiment would begin with a named scenario and a frozen baseline. The baseline identifies the service population, dependency assignments, admissible reach memberships, operational conditions and recovery assumptions. The experiment then changes one defined intervention or a documented bundle of interventions. Examples include replacing a shared implementation on selected critical functions, moving generation to a demonstrably independent design, narrowing an ancestry bound through new evidence, or changing a negotiated path so that it meets the service's target state.

The model should distinguish an **estate intervention** from an **evidence intervention**. Replacing an implementation changes the estate. Learning that two existing implementations are independent changes what is known. Both can change a modelled result, but only the first is a deployment change. This follows the same discipline used for movement reporting in Universal C.8.

A run should retain a trace from the initiating dependency through the affected assets and functions to each reported service outcome. This trace is necessary for challenge. A service shown as impaired because a named root was withdrawn should have a recorded path to that root and a documented reason that its alternatives cannot sustain the required function.

The experiment should also preserve the distinction between actual recorded facts and the state assumed during the scenario. Simulating the failure of a sound hybrid combiner does not mean the deployed combiner is known to be defective. Simulating a common ancestry at the upper bound does not establish that ancestry.

6.3 Uncertainty and sensitivity analysis

The model should separate uncertainty about facts from variability introduced for an experiment. Missing lineage, uncertain recovery time and an assumed counterparty response are different limitations. Combining them into one unlabelled output distribution would make the result difficult to interpret and impossible to attribute to a practical improvement.

For bounded membership, deterministic scenario analysis should be available before probabilistic sampling. The analyst can examine evidenced membership, a justified upper-bound assignment and intermediate feasible assignments. These are alternative structural scenarios, not probability quantiles.

Monte Carlo analysis uses distributions with a documented basis or explicitly selected to explore sensitivity. The run must retain the parameter distributions, dependence assumptions, random seed, sample count and convergence checks. A percentile obtained from an assumed distribution is a percentile of that model experiment, not an empirically established probability of service failure.

Sensitivity should be reported at the level where an owner can act. A result dominated by uncertainty in one HSM's generation provenance supports an evidence request. A result dominated by an untested four-hour recovery assumption supports an operational exercise. Those are different actions from replacing a supplier, even where each reduces the reported uncertainty interval.

6.4 Validation and safeguards

Validation should begin with simple, fully specified networks for which the expected result is known. A disconnected service must not become impaired without an explicit common dependency. Removing an unused dependency must not affect an unrelated service. An asset must not be counted twice because it reaches the same ancestor through two paths. A stronger intervention should not be described as beneficial where it removes the only working recovery option.

The validation plan should cover input transformation, reference integrity, reach membership, service transmission, time evolution, scenario controls and output interpretation. Tests should include shared ancestors, independent families, unresolved boundaries, alternative trust paths, capacity-limited fallback, sound and unverified hybrid constructions, incomplete populations, and the difference between marginal and joint population coverage.

Historical cryptographic incidents can support tests of particular mechanisms, such as inherited implementation defects or issued-key replacement. They do not validate an entire future network scenario or its quantum timing. Conversely, the absence of an observed quantum cryptanalytic event does not prevent verification of arithmetic, dependency propagation, historical operational behaviours or implementation correctness.

Independent reviewers should be able to reproduce a run from the versioned inputs and manifest, identify the assumptions driving its conclusions, and explain when the output should not be used. Material disagreements about a transmission rule should remain visible in the model documentation rather than being resolved through an undisclosed parameter choice.

7. GOVERNANCE

The sensitivity computation is a model. Where SR 11-7, PRA SS1/23 or equivalent requirements apply, an institution using an output in a decision documents the assumptions, states the limitations and obtains independent validation.

The model ships with a documentation pack containing an assumption register, a limitations statement, input data quality requirements tied to conformance level, and a validation test plan. The [Model Documentation Pack](#) provides the documentation structure, with the sensitivity model's equations, dynamics, calibration and experiment design recorded alongside it.

Before a runnable release is presented as suitable for decision support, the release should include the mathematical and computational specification, input and output contracts, an explicitly synthetic example dataset, a reference implementation, deterministic validation cases, sensitivity tests, known limitations and reproducibility instructions. The release should identify which propagation method is implemented and which alternatives remain research options.

The framework does not depend on this model. CCF's own exposure remains with the assessment computation governed under Universal G.6, where classification is the institution's determination. Classification of that computation does not change the separate model status of the sensitivity tool.

The contribution of the proposed model would be a more explicit account of how a measured dependency can become a network-level operational problem. Its value depends on maintaining the boundary between what the institution has measured, what it has determined from service evidence, and what the experiment assumes. CCF remains usable without crossing that boundary into simulation.

8. PUBLICATION

This document and any datasets are licensed under Creative Commons Attribution 4.0. The runnable model's code is licensed under Apache-2.0 and aligns with the CCF Universal version baseline. A clearly labelled synthetic example dataset ships with the runnable model.

RESOLVED AND OPEN AT V1.0-RC

Resolved in prior rounds: the model consumes assessment-workbook outputs and never computes framework metrics itself (one computation engine, one validation surface); edge weights are parameterised and swept, never silently fixed.

1. Resolved: a clearly labelled synthetic example dataset ships with the runnable model. A normative default topology would be adopted unexamined, and no example at all means most institutions never run it – the label is the control.
2. Resolved: the tolerance threshold reports per layer of each service, matching the layer-specific failure-impact determination in Universal C.5. Section 5 carries the change.
3. Open, carried publicly: which network-model family ships first in the runnable release, and the validation cases for the distress-versus-default ordering. Decided with the runnable release, on pilot-cycle input, and tracked at ccframework.org/roadmap.
4. Open, added 8 September 2026: a threshold for a non-linear measure depends on the path by which concentration changes, so the runnable release states whether thresholds are reported for a defined intervention trajectory or as a single value per layer, and under which ordering assumption. Decided with the runnable release.