



METHOD STUDY

QAOA Concentration Study

Which quantum method concentrates on the certified optimum - and what decides it. Training objective, ansatz, start state and penalty weights measured over 12,397 exact runs on 20 synthetic portfolios and 4 real demo books, every run scored against the proven optimum. A method study behind the Portfolio Optimization, Credit Allocation and Collateral Allocation RQPs.

August 2026



qubit-lab.ch

qubit-lab.ch

contact@qubit-lab.ch

© 2026 qubit-lab.ch · All rights reserved · qubit-lab.ch/legal

Contents

The study on one page	3
1. Question and frame	4
1.1 Why concentration	4
1.2 The four choices	4
1.3 What "escape" and "modal" mean	4
2. Design	5
2.1 Arms and budgets	5
2.2 Engine, referee and matching	5
2.3 Pre-registration and amendments	5
3. Instances	6
3.1 Synthetic	6
3.2 Real demo books	6
4. Phase 1 - the landscape	7
4.1 The penalty ratio dominates	7
4.2 Two things the matched budget showed	7
5. Phase 2 - the start state	8
5.1 A right seed is perfect	8
5.2 A wrong seed locks	8
5.3 The mixer is not the culprit - a null result	9
5.4 Depth pays off under CVaR, not under expected energy	9
5.5 The heuristic as it is, and the anti-seed	9
5.6 The adaptive ansatz from a wrong seed	9
6. Phase 3 - the real books	10
6.1 What replicates	10
6.2 What gets harder	10
6.3 Budget and cap unlock the real misses - partly	10
6.4 From nothing, at scale - the trap	11
7. Synthesis - a decision tree, and the settings it maps to	12
8. Limitations	13
9. Reproducibility	13
Appendix A - The pre-registered design (DESIGN.md, 15 August 2026)	14
A.1 Question	14
A.2 Phase 1 - lean core (locked 15 August)	14

A.3 Phase 2 - initialisation deep-dive (locked 15 August, run 16 August) 14

A.4 Phase 3 - extensions, "only if the core warrants" 14

A.5 Analysis plan (fixed with the design) 15

A.6 Deviations from the plan, disclosed 15

Disclaimer 15

The study on one page

With a certified referee in the loop, every reasonable QAOA configuration eventually finds the optimum on portfolio instances of prototype size. A 0% gap alone therefore does not rank methods. What differs between methods is **concentration**: how much probability the trained circuit puts on the proven optimum, whether that optimum is the single most likely readout, and how many measurement shots it would take to see it. This study measures concentration across four choices a user of the RQPs actually makes - training objective (expected energy or CVaR), ansatz (fixed depth or adaptive), start state (uniform or a classical seed) and the penalty weights - on 20 synthetic dense Markowitz instances at 12 qubits and 4 real demo books at 7, 14 and 16 qubits. 12,397 arms, exact statevector simulation, matched evaluation budgets, pre-registered design.

Six findings survived, in the order they matter for practice:

- 1. Tune the problem before the circuit.** The ratio of budget penalty to risk weight is the strongest single driver of concentration - Spearman -0.91 with $\log P(\text{optimum})$ under CVaR training, -0.76 under expected energy. Energy gap and covariance rank predict almost nothing.
- 2. A classical seed is a prior, not a result.** Started from the classical heuristic's answer, the circuit puts a right seed on top 100% of the time; started from a wrong seed it stays there - about half the runs escape at one asset off, one in six at two, one in twenty at three on synthetic instances ($n = 240$ per distance), and 0 of 96 standard runs from the heuristic's real wrong answers, at any budget. Swapping the matched mixer for a plain X mixer changes nothing: the lock lives in the starting state.
- 3. The seed was wrong on every real book.** The product's own heuristic missed the certified optimum on all four demo books (2-3 assets off, 0.65-1.60% gap), and 8,000 classical starts did not fix the 16-qubit one.
- 4. The adaptive ansatz escapes what the seed cannot - given depth and budget.** From controlled wrong seeds it escapes 83 / 75 / 50% of runs at 1 / 2 / 3 assets off on the real books (100 / 97 / 88% synthetic); the real 16-qubit miss 3 of 3 - but only at cap 12 and about 10,000 evaluations. Starved of depth or budget it fails; it never wins on cost.
- 5. From nothing, at scale, nothing concentrates - and there is a trap.** A deep CVaR circuit from a uniform start concentrates on its own at 7 qubits; at 14-16 qubits it does not, at these budgets. At 16 qubits the optimum was still the "most likely state" on 5 of 6 runs - at 5 times uniform, with the ten most likely states holding 0.09% together. A flag, not a result.
- 6. CVaR training beats expected-energy training** on every axis measured here, and depth pays off under CVaR ($P(\text{optimum})$ 0.08% to 0.30% from $p = 1$ to 6, uniform start) where it does not under expected energy.

The practical answer is a decision tree, and every branch is a setting of the products: trust the seed - classical seed start; doubt it - adaptive ansatz, deep and well-fed; know nothing - soften the penalty, CVaR at depth 4-6, and read the probability mass, not the flag.

qubit-lab.ch, Daniel Hug. Runs 15-16 August 2026; document version 1.0, 16 August 2026. Every number is re-derived from the committed raw results shipped with the study package; where a memo figure was superseded on re-derivation, the re-derived value is used and noted.

1. Question and frame

1.1 Why concentration

The RQPs run every quantum candidate against a certified classical referee - an exact solver (SCIP branch-and-bound) or exhaustive enumeration - so the honest question is never "did the quantum method find the optimum" (with a referee it always can be checked) but "how hard does the trained circuit point at it". A circuit that assigns the optimum 10% of its probability mass at 16 qubits (65,536 states) has done something; a circuit that assigns it 0.008% - five times a blind coin flip - has not, even if that 0.008% happens to be the tallest of 65,536 grains.

The product reports print, for every run: the probability of the optimum, its multiple of the uniform baseline $1/2^n$, the mass of the ten most likely states, whether the optimum is the modal (most likely) state, its rank, and the number of shots needed to observe it once with 99% confidence. This study asks which method choices move those numbers, and by how much.

1.2 The four choices

Choice	Levels measured	Where it lives in the products
Training objective	expected energy; CVaR (alpha = 0.10, the mean of the best 10% of the sampled energies)	<code>qaoa_cvar_mode</code> off / objective
Ansatz	standard fixed depth p in {1, 2, 3, 4, 6}; adaptive (ADAPT-QAOA style, mixer chosen per layer from an operator pool, natural stop) with cap 6 or 12	<code>qaoa_ansatz</code> standard / adaptive, <code>qaoa_adapt_max_layers</code>
Start state	uniform superposition; classical seed start (Egger et al. 2021: RY rotation toward a seed bitstring with blur epsilon, matched mixer)	<code>qaoa_warm_start_mode</code> none / classical, <code>qaoa_warm_start_epsilon</code>
Problem shaping	penalty ratio <code>lambda_budget</code> / <code>lambda_variance</code> in {0.5 ... 8}, covariance rank	<code>lambda_budget</code> , <code>lambda_variance</code>

1.3 What "escape" and "modal" mean

Throughout: **modal = optimum** means the certified optimum is the single most likely state of the trained circuit.

P(optimum) is its probability. **Escape** is used for seeded runs: the trained circuit's most likely state is the certified optimum rather than the (wrong) seed it started from. **Shots to 99%** is $\ln(0.01) / \ln(1 - P(\text{optimum}))$.

2. Design

2.1 Arms and budgets

Phase	Arms	What varies	Budget per arm
1 - landscape	1,425 (of 1,440; 15 expected-energy adaptive uniform cells failed on a gradient underflow and were excluded)	objective x ansatz {p 1, 2, 3, 6; adaptive cap 6, 12} x start {uniform, seed from the heuristic as is} x 20 synthetic instances x 3 optimizer seeds	3,000 objective evaluations, matched; single training run
2 - initialisation	10,200	objective x ansatz {p 1, 3, 4, 6; adaptive cap 6} x start {uniform; right seed at epsilon 0.10 / 0.25 / 0.40; wrong seed 1 / 2 / 3 assets off; the heuristic as is; an anti-seed ≥ 6 assets off} x mixer {matched, plain X} x 20 instances x 3 seeds	3,000 evaluations; product-style training: 3 restarts x 1,000
3a - real books	720	the same arms on 4 real instances built through the Portfolio Optimization RQP's own QUBO pipeline, seeded from its own heuristic	3,000 evaluations, 3 x 1,000
3b - budget and cap	52	adaptive cap {6, 12} x budget {3,000 / 10,000 / 20,000} and standard p = 4 as control, on the real wrong seeds of the 14- and 16-qubit books	as stated

Total: **12,397 arms**. Every arm reports P(optimum), x-uniform, top-10 mass, modal flag, rank of the optimum, shots to 99%, layers used, evaluations spent, and - for adaptive arms - the mixer sequence.

2.2 Engine, referee and matching

All arms run on an exact statevector engine (numpy) with the same conventions as the product core: MSB-first bit order, cost layer $e^{-i \gamma H_C}$, standard mixer $e^{-i \beta X}$ per qubit or the Egger matched mixer $e^{-i \beta (\cos \theta Z + \sin \theta X)}$ per qubit, adaptive operator pool of the global X, single-qubit X and Y, and two-qubit XX / YY / YZ strings on the 40 strongest QUBO couplings (plus the matched mixer when a seed is used). Angles are optimized with COBYLA (initial step 0.3), as in the product core. Screening gradients are the analytic commutator for expected energy and a symmetric finite difference for CVaR (which has no closed form); CVaR screening costs about 290 evaluations per grown layer, which the budgeting reserves in advance.

The referee is exhaustive enumeration of all 2^n states for every instance; the optimum is proven, not estimated. Budgets are matched in objective evaluations, screening included, so the adaptive ansatz pays for its operator selection out of the same budget the standard ansatz spends on optimization.

2.3 Pre-registration and amendments

The design was written before the first full run and is shipped with the study package. Three amendments were recorded during the study: the evaluation budget was raised from 600 to 3,000 after a smoke test showed CVaR-adaptive obtaining a single layer at 600; the per-layer budgeting was fixed to reserve future screening cost; and Phase 2 replaced the single training run by the product-style 3 x 1,000 after Phase 1's "depth barely matters" turned out to be a training-budget artefact (section 5.4).

3. Instances

3.1 Synthetic

Twenty dense Markowitz portfolio QUBOs at 12 qubits, generated on a grid of covariance rank (2, 4, 8, 12 factors) x penalty ratio $\lambda_{\text{budget}} / \lambda_{\text{variance}}$ (0.5 to 8), with a soft budget constraint. The soft budget makes every instance fully dense, so the "density" axis of the design is really covariance rank. On 7 of the 20 the demo-budget classical heuristic finds the optimum; on 13 it lands 2-4 assets away.

3.2 Real demo books

Book	Qubits	States	Heuristic seed vs certified optimum	Certified gap
Portfolio demo 7	7	128	3 assets off	0.65%
Portfolio adaptive demo 14	14	16,384	3 assets off	1.60%
Portfolio demo 16 (workbook classical budget)	16	65,536	2 assets off	1.40%
Portfolio demo 16 (core-default classical budget: 8,000 random starts, 40 local searches, 200,000 neighbour evaluations)	16	65,536	2 assets off - the same wrong seed	1.40%

The real instances were built through the product's own QUBO builder in CVaR band mode and seeded from the product's own classical heuristic - the seed a real user's warm start would inherit. That the heuristic misses on every book, and that more classical budget does not move the 16-qubit miss, is itself finding 3.

4. Phase 1 - the landscape

4.1 The penalty ratio dominates

Across the 20 synthetic instances, the single largest driver of concentration is not the ansatz or the objective but how hard the budget penalty dominates the risk term. On the reference cell - standard ansatz, $p = 3$, uniform start, 60 arms per objective (20 instances $\times$ 3 optimizer seeds) - $\log P(\text{optimum})$ correlates with the penalty ratio at Spearman **-0.91** under CVaR training and **-0.76** under expected energy. The instance's energy-gap ratio and its covariance rank predict essentially nothing on the same cell ($|\text{rho}| \leq 0.08$). The finding is robust to the aggregation: taking each instance's median $P(\text{optimum})$ over all Phase 1 arms of an objective gives -0.78 for both objectives.

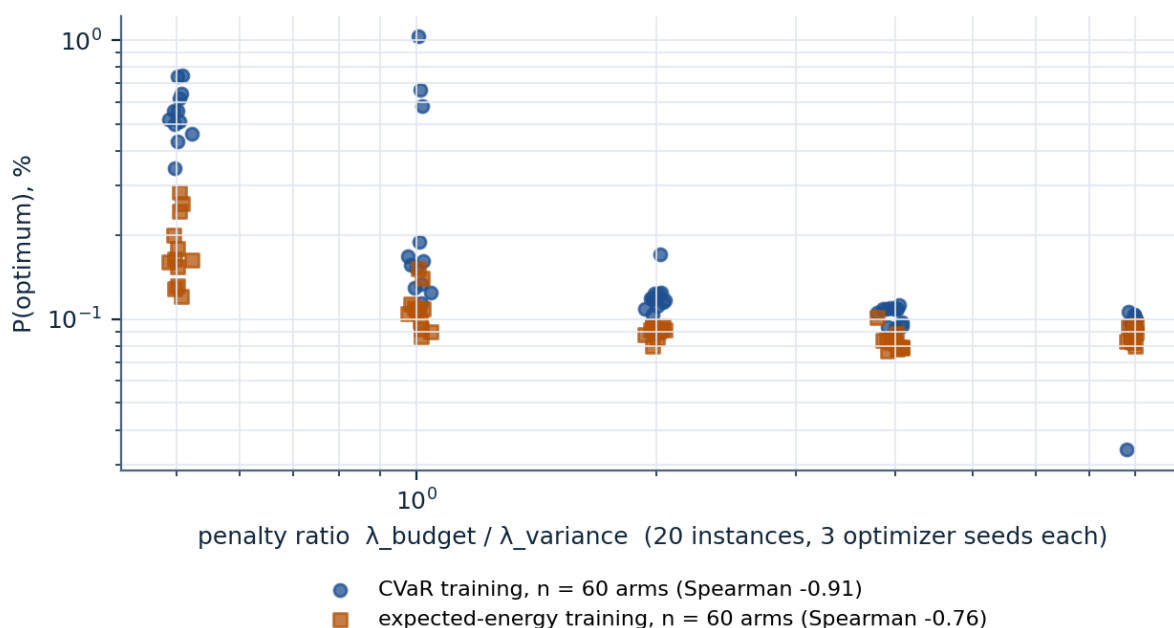


Figure 1 - $P(\text{optimum})$ against the penalty ratio on the reference cell (standard ansatz, $p = 3$, uniform start; 20 synthetic instances $\times$ 3 optimizer seeds per objective). Points are jittered horizontally for legibility. The softest penalty that still holds the budget concentrates best under both objectives.

Practical consequence: use the softest **lambda_budget** that still holds the budget constraint (the report's budget-gap and QUBO-breakdown panels show whether it does), then tune the circuit. This is the deflating opener - and the first thing to turn.

4.2 Two things the matched budget showed

At a matched 3,000-evaluation budget with a single training run, the adaptive ansatz starves: CVaR screening costs about 290 evaluations per layer, so CVaR-adaptive from a uniform start reached exactly the 64-state symmetry plateau ($P = 1/64 = 1.56\%$) on 59 of 60 cells - still 11 times the best standard uniform arm (0.14%), while expected-energy adaptive from a uniform start was the weakest method of all. The lesson survived into every later phase: the adaptive ansatz cannot win on efficiency, and it should never be judged at a starved budget.

The second observation - that depth barely mattered for the standard ansatz ($p = 1$ to 6: 0.07% to 0.10%) - did not survive: it was an artefact of the single training run, corrected in Phase 2 (section 5.4).

5. Phase 2 - the start state

5.1 A right seed is perfect

Seeded from the correct optimum (blur epsilon 0.25, matched mixer), every method puts the optimum on top on 100% of runs. The blur trades concentration for seed-dependence (standard ansatz, matched mixer, $n = 240$ arms per row and objective):

Epsilon	Median P(optimum), CVaR	Median P(optimum), expected energy	Note
0.10	29.5%	64.6%	roughly 3x the mass of 0.25
0.25	10.1%	9.5%	the measured middle ground - the product default
0.40	4.7%	0.8%	about half; the standard ansatz keeps the modal lock at every epsilon, the expected-energy adaptive ansatz loses it on 45% of runs at 0.40

5.2 A wrong seed locks

Seeded from a portfolio 1, 2 or 3 assets away from the optimum, the standard warm-started circuit escapes to the optimum on 46% / 17% / 5% of runs (CVaR, matched mixer, depths pooled; $n = 240$ per distance), with median P(optimum) collapsing from 8.9% to 2.3% to 0.7%. Its escape radius is about one asset. On the real books (Phase 3) the corresponding figures are 23% / 17% / 17% for controlled seeds ($n = 48$ per distance) - and 0 of 96 for the heuristic's actual wrong answers (4 books x 3 seeds x 4 depths x 2 objectives).

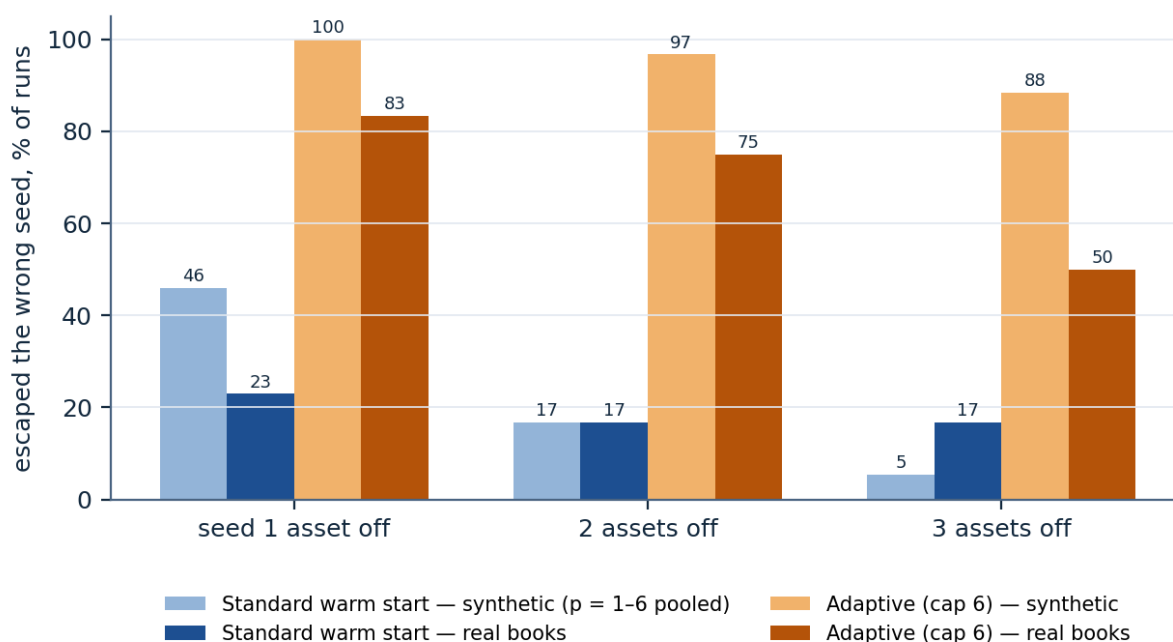


Figure 2 - Share of runs in which the trained circuit's most likely state is the certified optimum rather than the wrong seed it started from, by seed distance (CVaR training, matched mixer). Standard warm start pooled over depths (synthetic $n = 240$, real $n = 48$ per distance); adaptive ansatz at cap 6 and 3,000 evaluations (synthetic $n = 60$, real $n = 12$ per distance).

5.3 The mixer is not the culprit - a null result

The Egger construction matches the mixer to the seed state; a natural hypothesis was that the matched mixer causes the lock. It does not. With the plain X mixer on the same seed state the escape rates are 56% / 16% / 8% (n = 240 per distance) - a few points either way - while right-seed concentration drops. The lock lives in the starting state. The shipped matched mixer stays.

5.4 Depth pays off under CVaR, not under expected energy

With product-style training (3 restarts x 1,000 evaluations) from a uniform start, CVaR concentration rises steadily with depth; expected-energy training stays flat (n = 60 per depth and objective):

Depth p	P(optimum), CVaR (median)	Optimum modal, CVaR	P(optimum), expected energy
1	0.078%	25%	0.067%
3	0.124%	37%	0.095%
4	0.215%	45%	0.097%
6	0.295%	50%	0.104%

This is why the products' uniform-start guidance is CVaR at depth 4-6.

5.5 The heuristic as it is, and the anti-seed

Seeded from the heuristic's answer as it stands (right on 7 instances, wrong on 13), the standard warm start put the optimum on top on 35% of runs - exactly the right-seed fraction, i.e. it escaped none of the wrong ones - and the adaptive ansatz on 65%. Seeded deliberately far away (≥ 6 assets off), the warm start was worse than starting from nothing: a warm start is only as good as its seed.

5.6 The adaptive ansatz from a wrong seed

The one method that escapes wrong seeds is the CVaR-trained adaptive ansatz: 100% / 97% / 88% at 1 / 2 / 3 assets off on synthetic instances (cap 6, n = 60 per distance), with median P(optimum) 10% at every distance, mixer irrelevant. Expected-energy adaptive escapes 18% / 2% / 0% - the recommended pairing is CVaR.

6. Phase 3 - the real books

6.1 What replicates

Right seed to 100% modal optimum, every method (median $P(\text{optimum})$ about 10% for standard $p \geq 3$ and adaptive; 4.5-7.8% at $p = 1$). Escape from controlled wrong seeds: standard 23% / 17% / 17% ($n = 48$ per distance), adaptive **83% / 75% / 50%** ($n = 12$ per distance) at 1 / 2 / 3 assets off - the same shape as synthetic, adaptive a little weaker at distance 2-3 on real QUBOs. A wrong seed is a confident wrong answer on real books as on synthetic ones.

6.2 What gets harder

- The heuristic misses on every book (table 3.2), and the 16-qubit miss persists at the core-default classical budget: the wrong seed is not a demo-budget artefact; the local search sits in a genuine basin.
- On those real seeds, at 3,000 evaluations, only the 7-qubit case escapes: standard 0 of 96; CVaR-trained adaptive 3 of 3 on the 7-qubit book and 0 of 9 on the 14- and 16-qubit ones (expected-energy adaptive 0 of 12).
- Best cell per instance is always a right-seeded standard circuit ($P(\text{optimum})$ 10-63%). With a correct seed the standard warm start is unbeatable at this budget; the adaptive ansatz's whole value is what happens when the seed is wrong.

6.3 Budget and cap unlock the real misses - partly

Phase 3b asked whether "at this budget" was the true limiter. It was, for one of the two:

Real wrong seed	Method	3,000 evals	10,000 evals	20,000 evals
16-qubit book, 2 assets off	adaptive cap 12	P 0.26%, no escape	P 13.7%, escaped 3 of 3, 10 layers	3 of 3
16-qubit book	adaptive cap 6	0.6%, none	0.6%, none	0.6%, none
16-qubit book	standard $p = 4$	0.7%, none	0.7%, none	0.7%, none
14-qubit book, 3 assets off	adaptive cap 12	P 0.28%, none	P 10.0% (12 layers), no flip	10.0%, none
14-qubit book	adaptive cap 6 / standard $p = 4$	none	none	none

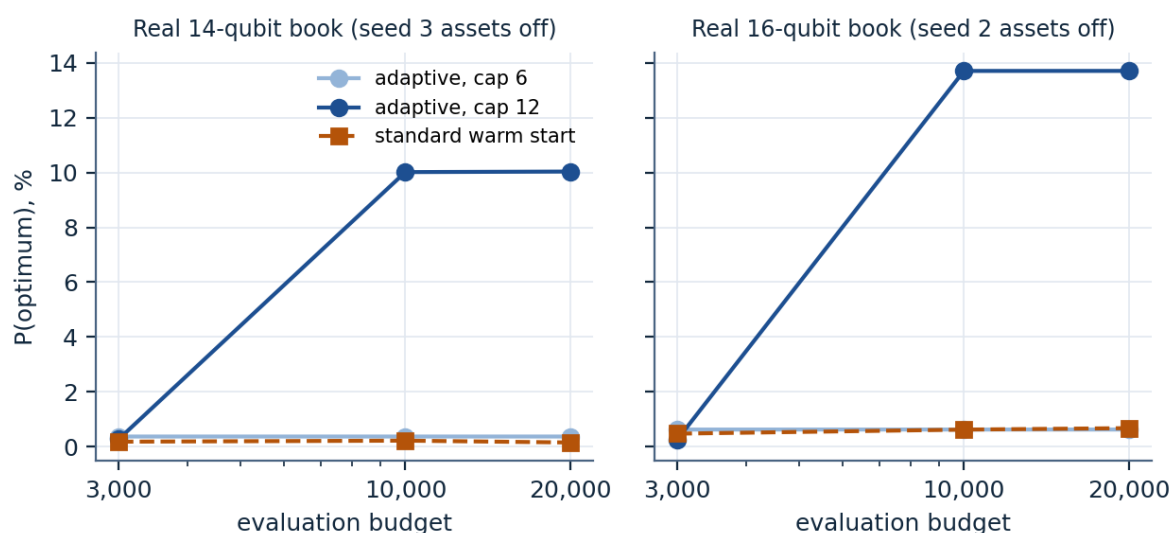


Figure 3 - $P(\text{optimum})$ on the two real wrong seeds against evaluation budget: adaptive at cap 6 and 12, standard warm start as control. Depth without evaluations, or evaluations without depth, do nothing; both together escape the 16-qubit miss and lift the 14-qubit one from 0.3% to 10%.

Mechanism: the adaptive ansatz needs depth to build the escape route and evaluations to afford that depth's screening. The product's default cap of 12 is vindicated; the manual now recommends 200-300 iterations per layer on real instances. The standard warm start did not escape at any budget tested: its lock is not a budget artefact.

6.4 From nothing, at scale - the trap

From a uniform start the 7-qubit book behaves like the synthetic instances: CVaR $p = 6$ puts 4% on the optimum and makes it modal on 2 of 3 runs, top-10 mass 33%. At 14 qubits $P(\text{optimum})$ is below 0.1% and the optimum is never modal. At 16 qubits the trap appears: CVaR $p = 6$ shows the optimum as the modal state on 5 of 6 runs - at 0.008-0.010%, five to seven times uniform, with the ten most likely states holding 0.07-0.10% together, about 50,000 shots to observe it once. A flat readout in which the optimum happens to be the tallest of 65,536 grains. The "most likely = optimum" flag without mass is not a result - which is exactly why the reports print probability, x-uniform, top-10 mass and shots-to-99% next to it.

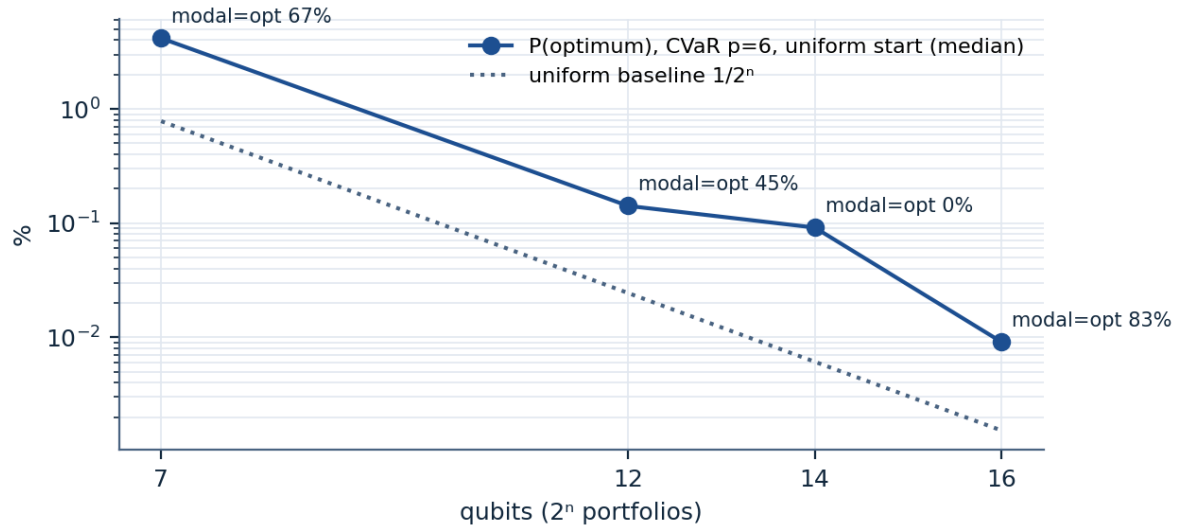


Figure 4 - Median $P(\text{optimum})$ from a uniform start (CVaR, $p = 6$) against register size, with the uniform baseline $1/2^n$ and the share of runs on which the optimum was the modal state. At 16 qubits the modal flag is true on 83% of runs while the mass is a few times uniform.

7. Synthesis - a decision tree, and the settings it maps to

What you know going in	What the study says	Product settings
You trust the classical seed	Classical seed start: optimum on top 100% of runs, ~10% mass at epsilon 0.25 (roughly 3x at 0.10). Read it against the certified gap - a wrong seed comes back just as confident.	qaoa_warm_start_mode classical, qaoa_warm_start_epsilon 0.25, qaoa_cvar_mode objective
You doubt the seed	Adaptive ansatz, deep and well-fed: escapes controlled misses 83 / 75 / 50%, the real 16-qubit miss 3 of 3 - only at cap 12 and ~10,000 evaluations. Never wins on cost; wins on not inheriting the seed.	qaoa_ansatz adaptive, qaoa_adapt_max_layers 12, qaoa_maxiter 200-300 (per layer), qaoa_cvar_mode objective
You know nothing	Soften the penalty first, then CVaR at depth 4-6 from a uniform start. Below about 10 qubits it concentrates on its own; at 14-16 it does not at these budgets. Read the mass and shots-to-99%, not the flag.	softest lambda_budget that holds, qaoa_cvar_mode objective, qaoa_p 4-6

Two null results belong next to the tree and are printed with it: the mixer swap changed nothing, and the adaptive ansatz never won at matched budgets. All four switches ship in the Portfolio Optimization, Credit Allocation and Collateral Allocation RQPs, and every report prints the concentration next to the certified gap. The referee does not care which method wins - that is the point.

8. Limitations

- **Exact simulation only.** No hardware, no shot noise in training. Finite-shot sampling would make the adaptive ansatz's operator screening worse (a winner's-curse effect over 300 noisy gradient estimates), which is why the products restrict the adaptive method to the exact-probability regime and replay finished circuits on hardware.
- **Scale.** 7-16 qubits; the synthetic grid at 12. The "from nothing" finding is stated for these budgets and sizes, not as a scaling law.
- **Instances.** Dense Markowitz QUBOs with a soft budget; four demo books. Other constraint structures (one-hot assignment, slack ladders) were not part of this study.
- **Budgets.** Matched at 3,000 evaluations for the main phases; the budget ladder to 20,000 was run only for the two real misses.
- **Untested.** Wrong seed $\times$ epsilon; alpha other than 0.10; adaptive with the expected-energy objective beyond Phase 1 (where it was the weakest method).

9. Reproducibility

The study package - the pre-registered design with its amendments, the engine and instance generator, the run scripts, all 12,397 raw results (one JSON record per arm with every metric named above) and the three analysis memos from which this document's numbers are taken - is available from qubit-lab.ch on request. Phases 1, 2 and 3b are self-contained (Python, numpy, scipy); Phase 3a builds the real instances through the Portfolio Optimization RQP backend, which is not part of the package, but every real-book result is included in the raw data. Every number in this document can be re-derived from those files with a few lines of pandas; the memos state which filters and aggregations produced each one.

Reference for the classical seed start: D. J. Egger, J. Marecek, S. Woerner, "Warm-starting quantum optimization", Quantum 5, 479 (2021). Reference for the adaptive ansatz: L. Zhu et al., "Adaptive quantum approximate optimization algorithm for solving combinatorial problems on a quantum computer", Physical Review Research 4, 033029 (2022) - adapted here for the CVaR objective.

Appendix A - The pre-registered design (DESIGN.md, 15 August 2026)

The design was written and committed before the first full run ("pre-registration in spirit"). It fixed the question, the factor space, the budget-matching rule, the metrics and the analysis plan. It did **not** register directional hypotheses: the six findings of this document are measurements, not confirmed predictions, and are labelled as such. What follows condenses the design lock as committed; amendments carry their dates.

A.1 Question

For a fixed circuit-evaluation budget, how much probability mass do different QAOA method choices place on the certified optimum of dense finance QUBOs - and does the answer depend on the instance? Concentration is the honest quantum-value metric in this framing: every method reaches a 0% certified gap eventually; what distinguishes them is $P(\text{optimum})$, the multiple of uniform, whether the modal state is the optimum, and the shot budget that concentration implies. The study measures those across the full method space instead of the two instances of the 14 August 2026 spike.

A.2 Phase 1 - lean core (locked 15 August)

Factor	Levels
Objective	expected energy; CVaR ($\alpha = 0.10$)
Ansatz	standard p in {1, 2, 3, 6}; adaptive cap in {6, 12}
Initialisation	uniform; warm (Egger, $\epsilon = 0.25$, matched mixer) seeded from the classical heuristic's answer as it stands, right or wrong per instance
Instance	20 synthetic dense Markowitz QUBOs at $n = 12$ spanning density x penalty ratio x gap ratio
Seed	3 optimizer seeds

Budget matching: every arm gets the same total objective-evaluation budget $B = 3,000$ (standard: spread over restarts of 200; adaptive: per-layer budget plus screening evaluations, counted - CVaR screening 2 x pool per layer, analytic screening for expected energy counted as one state preparation per layer). If an adaptive arm's cap x per-layer cost would exceed B , the per-layer budget is reduced to fit and actual evaluations are recorded. **Amendment before any run:** B raised from 600 to 3,000 after the smoke test showed CVaR-adaptive obtaining a single layer at 600.

Metrics per arm (exact statevector, no shot noise): $P(\text{optimum})$, x-uniform, top-10 mass, modal = optimum, rank of the optimum, best-exported gap (top-100 pool), shots to 99% = $\ln(0.01) / \ln(1 - P(\text{optimum}))$, layers used, evaluations used, wall time. Instance descriptors: n , edge density, coupling spectrum statistics, penalty ratio, gap ratio $(E_2 - E_1) / |E_1|$, heuristic gap and Hamming distance of the heuristic seed from the optimum. Cells: $2 \times 6 \times 2 \times 20 \times 3 = 1,440$ arms.

A.3 Phase 2 - initialisation deep-dive (locked 15 August, run 16 August)

Initialisation levels: uniform; right seed at ϵ in {0.10, 0.25, 0.40}; wrong seed at Hamming distance 1, 2, 3 from the optimum ($\epsilon = 0.25$); the heuristic as it stands; anti-seed (a random bitstring at Hamming distance $\geq n/2$, recorded). Mixer split in scope: warm state x {matched Egger mixer, plain X mixer}. Extra metric: escape rate (the modal state leaves the seed and lands on the optimum). **Amendments before any Phase 2 run (review of the Phase 1 results):** the depth axis was corrected - Phase 1's ~200-evaluation restarts under-trained deep circuits and hid the depth effect, so Phase 2 standard arms use fixed 3 restarts x $B/3$, matching how the product trains, with p in {1, 3, 4, 6}; adaptive keeps cap 6 (cap 12 was budget-starved at $B = 3,000$); an underflow guard was added to the adaptive natural stop (the cause of the 15 failed Phase 1 cells) and to shots-to-99%; qubit count stays 12 for exactness and breadth. Cells: $20 \times 2 \times 5 \times (1 + 8 \times 2) \times 3 = 10,200$.

A.4 Phase 3 - extensions, "only if the core warrants"

Listed: α sweep {0.05, 0.25}; shot-noise training at 1k / 10k shots; $n = 14$ replication; the four real workbook instances as anchors.

A.5 Analysis plan (fixed with the design)

Main effects and two-way interactions on $\log P(\text{optimum})$ by ordinary least squares with instance fixed effects; a depth-matched fairness table (adaptive cap 6 against standard $p = 6$); a per-instance winner map against the descriptors; seeds reported as spread (min / median / max), never averaged away in tables; negative results reported; memo and figures committed under the study folder.

A.6 Deviations from the plan, disclosed

- **The OLS model was not run.** The memos and this document use per-factor tables, medians with seed spread, escape rates with sample sizes, and Spearman correlations on named cells instead. The main-effects picture is therefore descriptive; no interaction terms are estimated.
- **The depth-matched fairness comparison** is reported through the matched-budget observations of section 4.2 and the budget ladder of section 6.3 rather than as a single table.
- **The per-instance winner map** was produced in Phase 1 (best cell per instance: CVaR-adaptive from a warm start on 11 instances, expected-energy adaptive from a warm start on 8, CVaR-adaptive from a uniform start on 1; the standard ansatz never won at the matched single-run budget) and is superseded by the Phase 2/3 findings; it is kept in the memos.
- **Phase 3 was re-scoped** after Phase 2: the alpha sweep and shot-noise training were not run; the $n = 14$ spot check became the four real demo books at 7, 14 and 16 qubits (Phase 3a) plus the budget-and-cap ladder on the two real misses (Phase 3b), which was not in the original list.
- **Counts corrected on re-derivation:** the Phase 3a memo's "0 of 60" for standard warm starts from the real wrong seeds is 0 of 96 (4 books x 3 seeds x 4 depths x 2 objectives); the memo is annotated, this document uses the re-derived value.

Disclaimer

This document reports a method study on demo data with a rapid quantum prototyping tool by qubit-lab.ch. Its figures are unedited outputs of the described runs, intended for technical review, education and structured discussion. They are not financial advice, not a performance claim for other problems or sizes, and not a hardware result. qubit-lab.ch is not affiliated with the authors of the cited references.

Legal notice. © 2026 qubit-lab.ch. All rights reserved. This document and the software it describes are proprietary to qubit-lab.ch and provided for authorized use only; reproduction or redistribution without written permission is prohibited. Outputs of the RQP tools are intended for technical review, experimentation, demonstration, and structured discussion - they are not financial advice. Full legal terms: qubit-lab.ch/legal.