



METHOD STUDY

The Advantage Island of Variational Classifiers

The companion map to the quantum-kernel study - where the trainable quantum model (VQC) wins, what a fair attempt costs, and which method a given problem deserves. 65 pre-specified training arms across three label families, register sizes 4-16 and the product's own configuration knobs, every arm scored against ten classical variants and the quantum kernel method head-to-head. A method study behind the QML RQP.

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 2

1. Management summary 3

2. Question and frame 3

 2.1 What the first study left open 3

 2.2 Two pieces of theory, both with fine print 3

 2.3 Best-case constructions, and one deliberate difference from our own sample 4

3. Design 4

 3.1 The fair recipe and the families 4

 3.2 The referee, and the head-to-head 4

 3.3 Staged grids under a registered ceiling 4

 3.4 Instruments 4

 3.5 Registration - provable this time 5

4. The territory map 5

5. The predicted wall does not appear in the tested regime 6

6. The knobs, and the price of a fair run 6

7. Instruments for the variational side 7

8. Hypothesis scorecard 7

9. Limitations 8

10. Reproducibility 8

11. Outlook 8

Appendix A - The registered design, amendments and deviations 8

 A.1 Registration - mechanism, scope and limits 8

 A.2 Deviations, disclosed 9

References 9

Disclaimer 9

The study on one page

The first advantage-island study mapped the quantum kernel classifier (QSVM). This companion asks the questions that map left open: does the trainable quantum model - the variational quantum classifier (VQC) - own territory of its own, when is it reasonable at all, what does a fair attempt cost, and what can be screened before spending? The stakes are concrete: the canonical trainability literature warns of vanishing gradients at scale, and the kernel-correspondence result suggests a support-vector machine on quantum-state distances might replace variational models outright. Neither folklore survives contact with the referee unmodified.

Throughout this study, "advantage" means a statistically significant predictive advantage against the specified classical referee set under the stated sample and estimation conditions. It does not mean computational speed-up, hardware advantage, or advantage over all possible classical algorithms.

The design: 65 pre-specified fair-recipe training arms (of a registered ceiling of 150) across three label families - VQC-native "teacher" labels, the kernel study's kernel-native labels, and its classical mixture dial - at 4-16 qubits, plus the product's configuration knobs on three registered cells and cheap pre-run instruments calibrated on every dataset. Every arm is scored against ten classical variants including an order-matched periodic twin and a tuned RBF kernel SVM, and against the QSVM head-to-head on identical rows. The design, hypotheses and analysis plan were locked, publicly pushed and tagged **before the first run** (git tag [vqc-island-design-v1](#)); every deviation is disclosed in Appendix A.

Five findings survived:

1. **Territory is size-gated, and the crossover is monotone.** At 4 qubits the kernel method wins everywhere - including a perfect score on the VQC's own native labels. But the paired VQC/QSVM ratio on VQC-native labels rises monotonically, $0.98 \rightarrow 1.02 \rightarrow 1.10 \rightarrow 1.14 \rightarrow 1.18$ across 4-16 qubits, while kernel-native territory stays the QSVM's at every size. The kernel correspondence holds; predictive dominance does not: at small registers the QSVM also wins on the VQC's native labels, but its held-out advantage disappears as registers grow - an inductive-bias story, not a contradiction of the correspondence [4, 5].
2. **No trainability wall observed through 16 qubits.** Initial gradient norms stay flat-to-rising through 16 qubits (0.055-0.092; the local-cost, shallow-circuit fine print of the barren-plateau literature, measured), and on native structure the fair-recipe VQC posts its best score AT 16 qubits (0.93). The registered 14-16-qubit degradation prediction is refuted on native territory.
3. **512-shot inference noise is negligible in the tested regime.** All 16 persisted models re-evaluated with sampled expectation values at 512 shots move by median -0.0002 macro-F1. The VQC's measurement cost at inference - a handful of expectation values per row - is structurally more forgiving than the QSVM's per-kernel-entry estimation.
4. **A fair run costs less than feared - if the stopping rule is sane.** With patience scaled proportionally, 150-iteration runs stay quantum-ahead on matched structure; the registered prediction of verdict flips between 150 and 600 iterations is refuted. Fairness is full circuit depth plus proportionate early stopping, not brute-force budget. Depth is the load-bearing knob, the near-identity warm start recovers half of what missing depth costs, and no knob buys territory the data does not grant.
5. **Cheap screening works - and one of our own bets lost.** A 25-iteration probe run predicts the full-run verdict at AUC 0.76 (registered threshold met). Our registered anti-folklore bet - that initial gradient variance would NOT predict outcomes - is refuted: it predicts at AUC 0.79. The kernel study's gram-tail instrument does not transfer to VQC outcomes (AUC 0.54), confirming the shipped kernel-only scope of the island panel.

Scope, stated once and meant everywhere: exact-statevector training (training under shots is registered future work); synthetic native-label families - a best-case benchmark envelope per family, not a real-world claim; one ansatz family (hardware-efficient, linear entanglement, local cost); 65 trainings against a disclosed ceiling, with single-seed knob probes disclosed as such.

1. Management summary

Written for decision-makers; the measured basis for every sentence follows in sections 2-8. The five questions below are the questions this study was commissioned to answer.

Q1 - Two quantum models: do we need both? Yes - but not at every size. Below roughly eight qubits the kernel method dominated everything we measured, including data built for the trainable model. Above that, structure decides: on its native structure the trainable model pulls ahead - by 18% at sixteen qubits - while kernel-native structure stays with the kernel method. Which method a dataset deserves is not a matter of taste; it is measurable, and the crossover is now mapped.

Q2 - When is the trainable model reasonable at all? Within this study's regime: whenever the data's structure matches its architecture and the run is configured fairly. The feared trainability wall did not appear through sixteen qubits - gradients stayed healthy, and the model's best result came at the largest size tested. The walls that do exist are structural (wrong data) and configurational (unfair runs), both detectable.

Q3 - Which knobs matter, and what does a fair attempt cost? Circuit depth matters most; a near-identity warm start recovers about half of what missing depth costs; extra restarts help on hard cells; one popular knob (periodic head refitting) mildly hurts. And the price of fairness is lower than our own pilot suggested: with a sane stopping rule, a quarter of the full training budget already delivers the same verdicts on matched data. No knob, at any price, won on data whose structure belonged to someone else.

Q4 - Can we trust a variational verdict? Only if the run was fair - and fairness is now checkable, not aspirational: full depth, proportionate stopping, matched encoding. The product flags budget-limited runs automatically, and this study is the calibration behind that flag.

Q5 - What can we screen before spending? Two cheap measurements predict full-run outcomes: a 25-iteration probe training, and - to our own registered surprise - the variance of the model's initial gradients. Both now qualify for the assessment's variational instrument set. The kernel study's instrument does not transfer, which is why the product keeps its kernel panel scoped to kernel methods.

The practical deliverable, as in the first study, is a map rather than a slogan: a measurable procedure for deciding, before training money is spent, which quantum method - if either - a dataset should receive, and what a fair attempt will cost. The deeper lesson of the knob experiments belongs here too: do not begin by choosing a quantum algorithm - begin by asking which representation makes the data learnable; the method choice follows.

2. Question and frame

2.1 What the first study left open

The companion study measured where the quantum kernel classifier can win and shipped the instruments that place a dataset on that map. It deliberately excluded the product's second quantum model. A pilot arm there produced this study's founding lesson: a thirteen-cell "the VQC never wins" verdict was overturned overnight by a pre-specified verification protocol - the binding constraint was training budget, not physics. A fair run is part of the method; this study is that lesson, industrialized.

2.2 Two pieces of theory, both with fine print

Trainability. The barren-plateau results [1, 3] warn that gradients of random parametrized circuits vanish exponentially with register size - the canonical reason to doubt variational models at scale. The fine print [2]: local observables in sufficiently shallow circuits can avoid the exponential regime under the stated conditions. Our stack - expectation values of single-qubit observables, a shallow hardware-efficient ansatz - lives in exactly that exception, and the pilot saw no initial-gradient decay at 4-16 qubits. This study instruments the question rather than assuming either answer: initial gradients, stall behavior and budget use are measured on every arm.

Kernel correspondence. Schuld's result [4] shows that many supervised variational models are mathematically kernel methods, and that kernel-based training carries guarantees variational training lacks - suggesting the QSVM might

simply replace the VQC. But the correspondence is a statement about model classes, not about finite training budgets on finite data. Whether the kernel method actually poaches the variational model's territory is an empirical question; it is registered here as an explicitly two-sided test, published either way.

2.3 Best-case constructions, and one deliberate difference from our own sample

The VQC-native family labels points by a fixed random teacher instance of the exact architecture the student trains - the same construction as the product's certified sample, with one deliberate difference: the sample's generator screens teacher seeds against the classical bar; this study does NOT screen, because screening would select for VQC-favourable instances and bias the territory test. The teacher seed was fixed a priori. Kernel-native and mixture families are the first study's generators, unchanged. All native-label results are best-case envelopes for their family, not real-world claims.

3. Design

3.1 The fair recipe and the families

Every training arm uses the fair recipe established by the pilot's verification: hardware-efficient ansatz at full depth (4 repetitions), local expectation-value readout with a linear head, Adam, 600 iterations with early stopping whose patience is proportionate to the budget, and the feature map matched to the label family (the teacher's own angle map on VQC-native cells; the kernel study's zz-style map on kernel-native and mixture cells). Cross-map arms appear only as registered knob probes. Exact statevector throughout; training under sampled measurements is registered out of scope.

3.2 The referee, and the head-to-head

The classical referee is the first study's, unchanged: ten variants - logistic regression, random forest, gradient boosting and the order-matched periodic twin in parity and best-reference modes, a tuned RBF kernel SVM, and a neural network - each with its own validation-tuned threshold. Verdicts follow the same absolute rule (whole 95% paired-bootstrap confidence interval above 1.0), with the selection-aware panel-max bootstrap reported alongside (agreement 35 of 36 Stage-A arms for both quantum models). In addition, every cell runs the QSVM with the family's own feature map on identical splits, and the paired VQC-versus-QSVM ratio on identical test rows is the study's primary territory readout.

3.3 Staged grids under a registered ceiling

A fair variational training costs minutes to hours, not seconds - so this study registered a hard ceiling of 150 full trainings and staged its grids: Stage A, the territory map (36 arms: four label settings x three sizes x three splits); Stage B, the wall probe (6 arms at 14-16 qubits); Stage C, the knobs (23 variant arms on three registered cells, single split, disclosed); Stage D, instruments and the inference-shot pass (probe runs and re-evaluations, not counted). Total spent: 65 counted trainings, 24.0 training-hours; per-arm wall-clock is recorded in the raw results.

3.4 Instruments

Per unique dataset: a 25-iteration probe training; initial-gradient statistics over 20 random initializations; the kernel study's configured-map gram statistics; classical headroom from the stored referee results. Scored by AUC against variational island membership (win fraction $\geq 2/3$), axis orientations fixed in the registered design - including the anti-folklore bet that initial-gradient variance would fail. Sixteen datasets, seven members: small-N calibration, disclosed as such.

3.5 Registration - provable this time

The complete design - grids, ceiling, hypotheses S2-H1..H7, analysis plan, instrument orientations - was committed and pushed to the public repository with the annotated tag **vqc-island-design-v1** (commit **8c976d4**) before any study run executed, so the chronology is publicly verifiable. Skeptical readers will find the mechanism's exact scope and limits in Appendix A.1. Amendments were recorded before their runs and are disclosed in Appendix A; all numbers in this document are re-derived from committed raw results by committed scripts.

4. The territory map

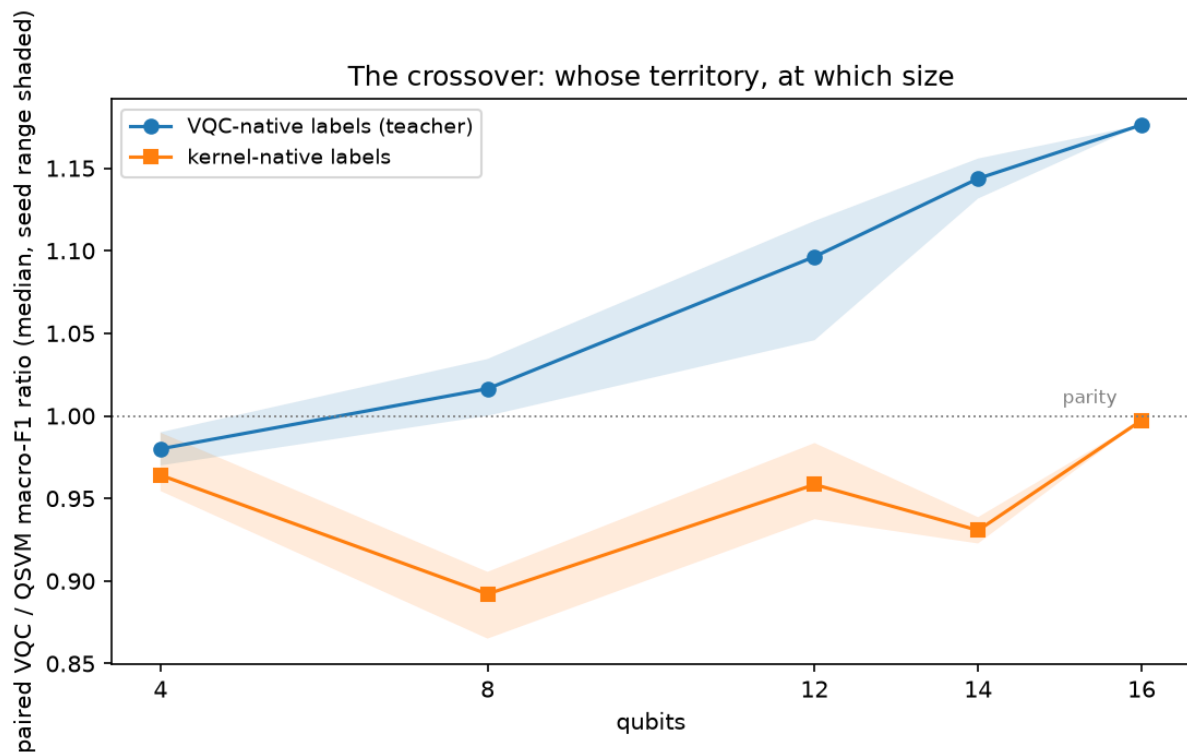


Figure 1 - The crossover. Median paired VQC/QSVM macro-F1 ratio on identical test rows (seed range shaded), by register size and label family. On VQC-native labels the ratio rises monotonically from 0.98 at 4 qubits to 1.18 at 16; on kernel-native labels it stays at or below parity at every size, touching 0.997 only at 16 qubits where the kernel study measured the QSVM's own erosion.

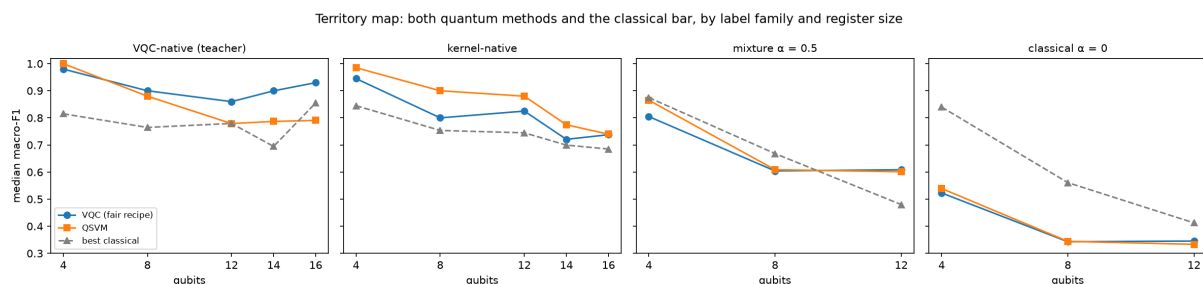


Figure 2 - The full territory grid: both quantum methods and the best classical variant, by family and size. VQC-native territory changes hands with size; kernel-native territory does not; the classical mixture stops both quantum methods identically.

At 4 qubits the kernel method wins every family it can win - including a perfect 1.00 on all three splits of the VQC's own teacher labels, where the trained VQC reaches 0.97-0.99. By 8 qubits the two methods are statistically inseparable at home (paired 1.00/1.02/1.03); from 12 qubits the VQC leads its territory outright, and the lead widens to 1.18 at 16. The theory-consistent reading: Schuld's correspondence [4] concerns model classes and training-set fit and does not expire with size; what reverses is held-out predictive dominance - above the crossover the restricted variational model generalizes better on its native structure, the explicit-versus-implicit distinction of Jerbi et al. [5] made empirical. On kernel-native labels the QSVM wins 83% of arms across all sizes (the VQC 42%, mostly at the extremes) and the paired

ratio never favours the VQC. On the classical mixture, the wall is model-independent: at $\alpha = 0$, zero whole-CI wins for either quantum method in 18 arms - the same referee discipline, the same result, from two entirely different quantum models.

The management reading of Figure 1 is the study's central deliverable: below the crossover, one quantum method suffices; above it, the data's structure chooses the method - and both the crossover and the structure are measurable in advance.

5. The predicted wall does not appear in the tested regime

The registered wall hypothesis expected the fair-recipe VQC to degrade at 14-16 qubits. It did not - not on its own territory. Its macro-F1 at the registered sizes reads 0.97 / 0.88 / 0.81-0.86 / 0.89 / **0.93** across 4-16 qubits (the 12-qubit dip is one seed; the other two splits sit at 0.86): the best native-territory score arrives at the largest size tested. Initial gradient norms are flat-to-rising across the whole range (medians 0.055-0.092) - no barren-plateau signature, exactly as the local-cost fine print [2] permits - and the stall instrumentation shows the 16-qubit run using 90% of its budget productively rather than stalling early. The refuted half of the hypothesis is printed in the scorecard; the confirmed half (no gradient decay) is what the folklore needed testing against.

The measurement-cost story splits cleanly from the kernel study's. There, estimating every Gram entry is the wall, arriving about two register sizes after the per-entry heuristic predicts. Here, inference needs only a handful of expectation values per row: re-evaluating all sixteen persisted models with binomially sampled measurements at 512 shots moves test macro-F1 by a median of -0.0002 (extremes -0.045 / +0.057) with thresholds re-tuned per shot level - and, in a sensitivity check added on external review, by a median of -0.004 (extremes -0.025 / +0.027) with the exact-trained threshold **held fixed**. The robustness is a property of the trained classifier, not of a recalibration step. Within this study's scope - inference only - 512-shot noise is negligible for the variational model. Training under sampled measurements is the registered open flank.

6. The knobs, and the price of a fair run

Single-cell sensitivity tests (not population estimates): depth and structural match dominate; budget beyond 150 scaled iterations buys little

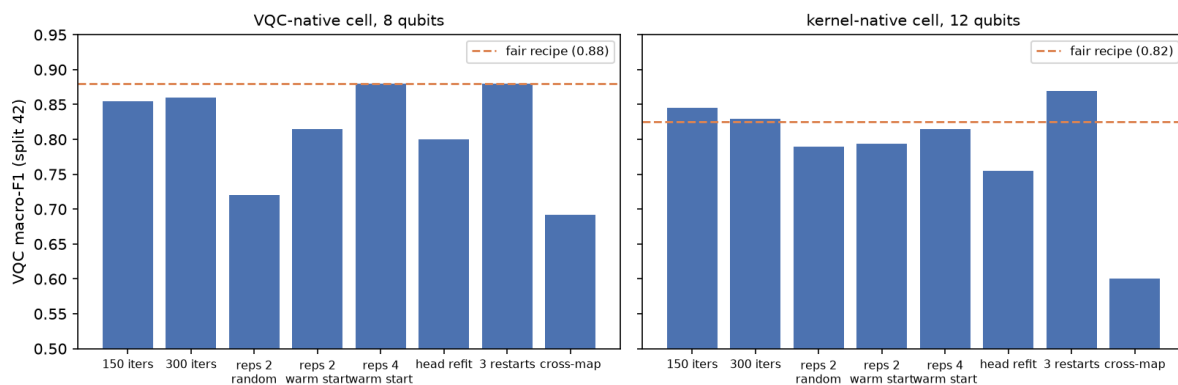


Figure 3 - Single-cell sensitivity tests, not population estimates: the registered knob probes on the two structure-matched cells (single split, disclosed). Depth dominates; the near-identity warm start recovers about half of what missing depth costs; the cross-map arm crashes; a quarter of the iteration budget already matches verdicts; head refitting mildly hurts; restarts are the kernel cell's best knob.

Three lessons, each priced:

- **Depth is the load-bearing knob.** Halving the ansatz repetitions costs the teacher cell $0.88 \rightarrow 0.72$. The identity-block warm start - the variational analog of a warm-started QAOA - claws back to 0.81 at reduced depth and adds nothing at full depth: it is depth insurance, not a substitute.
- **Structural match dominates everything.** Swapping in the other family's feature map crashes both matched cells ($0.88 \rightarrow 0.69$; $0.82 \rightarrow 0.60$), and on the mixture cell every knob combination - depth, warm start, refit, budget, all of it - stays classical-ahead. Knobs tune within territory; they cannot buy territory.

- **Fairness is cheaper than our own pilot implied.** With patience scaled proportionally, 150-iteration runs keep the quantum-ahead verdict on both matched cells (0.855 / 0.845 against fair-recipe 0.88 / 0.825); the registered prediction of verdict flips between 150 and 600 iterations is refuted. The pilot's dramatic budget rescues were artifacts of the old unscaled stopping rule - the product fix that followed (proportionate patience, budget-limited flag) turns out to be most of what "fair" means. Restarts (three, best-by-validation) are the kernel cell's best knob (0.825 → 0.870); periodic exact head refitting mildly hurts on both cells and stays off.

7. Instruments for the variational side

The sixteen study datasets, each with a known outcome, calibrate the assessment's missing variational axes - with one registered bet lost in public:

- **The probe run qualifies, with wide error bars.** A 25-iteration truncated training predicts the 600-iteration verdict at AUC 0.76 (bootstrap 95% interval 0.47-0.97): the point estimate meets the registered threshold, but at sixteen datasets the interval includes chance. Cost: about a minute per dataset at small sizes.
- **Initial-gradient variance carries signal - against our registered bet.** We registered the anti-folklore position that init-gradient statistics would fail as predictors (AUC < 0.60). Measured: the highest observed AUC in this calibration, 0.79 (bootstrap interval 0.55-1.00, above chance though wide) - datasets on which random initializations produce varied gradients are the ones the trained model goes on to win. The bet is printed as lost; both instruments qualify for provisional inclusion in the variational assessment, to be re-ranked as the benchmark grows.
- **The kernel study's instrument does not transfer** (gram tail AUC 0.54), and classical headroom carries no variational signal on these families (0.38). Both confirm decisions already shipped: the island panel stays kernel-scoped, and headroom remains a gate against classical-solved data rather than a quantum-opportunity signal.

Calibration caveat, stated plainly: sixteen datasets, seven members - directional evidence for the radar's variational sector, to be recalibrated as the benchmark grows, not a finished calibration.

8. Hypothesis scorecard

Pre-registered outcomes, in the order registered. As in the companion study: refuted entries are results, not failures.

Hypothesis	Outcome
S2-H1 - VQC wins its native family at 4-8 qubits; QSVM-poaching left explicitly two-sided	Confirmed, and the open question resolved. VQC win fraction 1.0 at 4-8; the QSVM also wins there (poaching real at small sizes) and cedes monotonically - paired 0.98 → 1.18 by 16 qubits (Section 4).
S2-H2 - kernel territory stays the QSVM's at fair budget	Confirmed. QSVM 83% wins vs VQC 42%; paired ratio never above parity for the VQC (Section 4).
S2-H3 - classical wall is model-independent	Confirmed. 0 of 18 quantum wins at alpha = 0 (Section 4).
S2-H4 - degradation at 14-16 qubits via stall, not gradient decay	Gradients: confirmed (no decay through 16 qubits). Degradation: refuted on native territory - best score at 16 qubits (Section 5).
S2-H5 - depth, warm start, match, refit-null, budget flips	Depth and match confirmed; warm start confirmed as depth insurance; refit null leans mildly harmful; the 150-vs-600 verdict-flip prediction refuted (Section 6).
S2-H6 - probe run predicts (≥ 0.70); gradient variance fails (< 0.60); gram does not transfer	Probe point estimate meets threshold (0.76, wide interval); gradient-variance bet REFUTED (0.79, interval above chance); gram non-transfer confirmed (0.54) (Section 7).
S2-H7 - verdicts stable at 4,096 inference shots, margin-dependent at 512	Confirmed, stronger than registered: stable at both levels, median delta -0.0002 at 512 (Section 5).

9. Limitations

- **Best-case envelopes per family**; no real-world claim follows. The mixture family's classical wall is the honest counterweight, and our real-data benchmarks remain classical territory.
- **One ansatz family** (hardware-efficient, linear entanglement, local cost) - the regime the barren-plateau fine print protects. Deep, global-cost or highly expressive architectures are outside scope, and the no-wall finding must not be quoted beyond this regime.
- **Exact-statevector training**. The inference-shot result does not license claims about training under sampled measurements - registered future work.
- **Ceiling-bounded replicates**: three splits per Stage-A cell, one to two at 14-16 qubits, single-split knob probes, one teacher seed per size. Seed spread is shown wherever a headline number is quoted; the 12-qubit teacher dip is reported, not smoothed.
- **Small-N instrument calibration** (16 datasets); AUCs are directional.
- **Numerical note**: overflow warnings from the simulation stack appeared during cross-map training runs; results were consistent across cells and are reported as measured (Appendix A).

10. Reproducibility

The complete study - registered design (tag [vqc-island-design-v1](#), commit [8c976d4](#), pushed before the first run), engine, runners, raw per-arm results including per-arm wall-clock and persisted model parameters, analysis scripts and generated tables - is committed under [studies/vqc-island-2026-09/](#) in the qubit-lab repository. Every number in this document traces to [tables_*.generated.json](#). The referee, training path and feature maps are the shipping product's own code; the fair recipe is the product's certified sample configuration.

11. Outlook

Three follow-ups are queued: placing real financial datasets on both studies' maps (the search procedure the management summaries promise, run in earnest); training under sampled measurements, the variational model's untested flank; and the product integration of the variational instrument axes - the probe run and the initial-gradient-variance signal - alongside the already-shipped kernel panel, completing the two-sector radar the assessment roadmap calls for.

Appendix A - The registered design, amendments and deviations

A.1 Registration - mechanism, scope and limits

Design, hypotheses, grids, ceiling, instrument orientations and analysis plan locked 29 August 2026; committed and pushed with annotated tag [vqc-island-design-v1](#) (commit [8c976d4](#)) prior to the first study run. One design note was added after the storyline lock but before any run: the teacher generator's explicit no-screening rule (Section 2.3).

For precision: this is public pre-specification with verifiable push order - immutable repository history demonstrates that the tagged design precedes every result - and NOT an institutional preregistration; no third-party registry took custody of the design before execution, and an archival DOI deposit made after the fact could attest the tag but not add chronology. (The first study in this series was pre-specified in writing but its public push postdated its first grid; this study closes the chronology gap.) The next study will deposit its design with an independent registry before its first run.

A.2 Deviations, disclosed

1. **Single-seed knob probes.** Stage C runs one split (42) per variant, as registered; knob conclusions are cell-level evidence, not distributions.
2. **Initial-gradient statistics** use one-iteration trainings to read the first gradient norm (20 initializations per dataset) - an implementation choice within the registered "no training" intent, disclosed because one optimizer step technically occurs.
3. **Numerical overflow warnings** (simulation stack, **overflow encountered in square**) appeared during cross-map training arms. Runs completed; results are plausible and consistent across both affected cells; no arm was re-run or excluded.
4. **A seed-level dip** at the 12-qubit teacher cell (0.81 on split 42 vs 0.86 on splits 7 and 11) is reported as spread; the crossover figure shades the full seed range.
5. **The pilot** (13 default-recipe trainings and the verification that overturned them) predates this study's registration and is inherited as motivation, not evidence; its data lives with study 1.

References

1. J. R. McClean, S. Boixo, V. N. Smelyanskiy, R. Babbush, H. Neven, Barren plateaus in quantum neural network training landscapes, Nature Communications 9, 4812 (2018); arXiv:1803.11173.
2. M. Cerezo, A. Sone, T. Volkoff, L. Cincio, P. J. Coles, Cost function dependent barren plateaus in shallow parametrized quantum circuits, Nature Communications 12, 1791 (2021); arXiv:2001.00550.
3. Z. Holmes, K. Sharma, M. Cerezo, P. J. Coles, Connecting ansatz expressibility to gradient magnitudes and barren plateaus, PRX Quantum 3, 010313 (2022); arXiv:2101.02138.
4. M. Schuld, Supervised quantum machine learning models are kernel methods, arXiv:2101.11020.
5. S. Jerbi, L. J. Fiderer, H. Poulsen Nautrup, J. M. Kübler, H. J. Briegel, V. Dunjko, Quantum machine learning beyond kernel methods, Nature Communications 14, 517 (2023); arXiv:2110.13162 - the explicit-versus-implicit model distinction behind Section 4's reading of the crossover.
6. S. Thanasilp, S. Wang, M. Cerezo, Z. Holmes, Exponential concentration in quantum kernel methods, Nature Communications 15 (2024); arXiv:2208.11060 - the kernel-side wall referenced in Sections 4-5.
7. J. Mancilla, T. Tagliani, The Fourier Wall, arXiv:2607.15815 - the periodic-probe methodology of the referee's twin follows this work (independent implementation inspired by the paper; no affiliation).
8. qubit-lab.ch, The Advantage Island of Quantum Kernels (August 2026) - the companion study; map, referee and mixture generator are shared.

Disclaimer

This document reports a method study on synthetic benchmark 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. The reported winning regions are best-case benchmark envelopes measured on labels constructed to favor the respective quantum method; they are not financial advice, not a performance claim for real-world datasets, 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.