

Where Can a Quantum Kernel Help?

A Readout-Locality No-Go that Strengthens with System Size

Author

June 8, 2026

Abstract

Two obstructions to quantum kernel advantage have been established separately: local (projected) readouts can be *dequantized*—reproduced to machine precision by a classical kernel—while global (fidelity) readouts *concentrate* exponentially and stop training. We show these are not two separate phenomena but the two ends of a single axis, *readout locality* k , and that for a broad encoding family the interval between them contains no usable advantage. Fixing the encoding and sliding k from single-qubit Pauli observables (which dequantize) to the full fidelity kernel, we find on eight tabular datasets that no locality beats a classical RBF, while feature concentration rises monotonically toward the global end. We then make the central prediction operational: when a genuinely high-weight signal is *planted* so that only a high-locality readout can represent it, the resulting advantage *decays toward zero as the qubit count grows*—the only readout that can host the signal concentrates it away, with the planted-feature variance and the kernel concentration both collapsing geometrically in q . The advantage window does not merely fail to open; it narrows as system size grows. The result is robust across six encoding families—including three with no closed form, and a CNOT-based map—so the no-go is not an artifact of the commuting structure that makes dequantization provable; on the same encoding the concentration also obstructs gradient training of a variational model (a barren plateau), so it is not specific to kernels. Throughout, a single pre-training quantity, the centered kernel-target alignment gap $\Delta = \text{KTA}(\text{quantum}) - \text{KTA}(\text{RBF})$, predicts where on the axis advantage can appear—including on a kernel with no closed form—turning the obstruction into a measurable test that any future encoding must pass. Our claims are scoped to classical tabular data.

1 Introduction

Quantum kernel methods embed inputs x into a state $|\psi(x)\rangle$ and classify with a support vector machine on a kernel built from those states. Two failure modes are by now well documented. The *fidelity* kernel $|\langle\psi(x)|\psi(x')\rangle|^2$ concentrates exponentially with qubit count and becomes uninformative [4]. The *projected* quantum kernel (PQK), introduced to mitigate this by reading out only local reduced observables [1], is for common encodings classically simulable—its kernel can be reproduced by an explicit classical feature map.

These two results are usually told as separate stories with separate proof techniques. Our first observation is that they are the two endpoints of a single axis. A readout is defined by which observables one measures; order them by *locality* k , the maximum Pauli weight retained. At $k = 1$ one recovers the local PQK (which dequantizes); at $k = q$, summing all Pauli weights recovers the fidelity kernel (which concentrates). The natural question is then not “does this kernel win?” but “*where on the locality axis, if anywhere, can advantage appear?*”—and whether the interval between the two known limits is empty.

We answer this for the standard depth-one IQP encoding family and its non-commuting extension. Our contributions:

1. **One axis, two limits (Section 3).** We make readout locality an explicit continuum from the dequantizing local POK to the concentrating fidelity kernel, anchored at $k = 1$ by an exact closed form and at $k = q$ by the fidelity identity.
2. **An empty advantage window on real data (Section 4).** On eight tabular datasets, no locality k yields advantage over a classical RBF tuned identically; concentration climbs $\sim 10^2 \times$ toward the global readout.
3. **The obstruction strengthens with system size (Section 5).** Planting a high-weight signal that only a high-locality readout can represent, the achievable advantage *decays toward zero as q grows*, with planted-feature variance and kernel concentration collapsing geometrically. This converts a fixed-size null into an asymptotic statement and identifies the mechanism.
4. **Robustness across encoding families (Sections 6–7).** The empty window persists across six families—no entanglement, IQP at depths one to three, a non-commuting re-uploading map, and a CNOT-based hardware-efficient map—including three with no closed form. The no-go is not an artifact of commutativity or of the closed form.
5. **A measurable test, validated without a closed form (Section 8).** The pre-training alignment gap Δ predicts where advantage can appear (Spearman $\rho \geq 0.8$), *including* on a kernel with no closed form, so the obstruction is a property any candidate encoding can be tested against rather than a feature of one circuit.
6. **The obstruction extends to trained models (Section 9).** On the same encoding, a variational model’s gradient variance vanishes geometrically in q —a barren plateau—at the kernel-concentration rate, and the all-to-all encoding leaves no trainable local escape.

The dequantization anchor is specific to the commuting depth-one map; the locality axis, the concentration limit, the scaling argument, and the Δ screen are not.

2 Background and Related Work

Quantum kernels and the power of data. Huang et al. [1] formalize when a quantum model can outperform a classical one, introduce the geometric difference $g(K_C \| K_Q)$ as a necessary-but-not-sufficient condition for advantage, and propose the POK as a robust alternative to the fidelity kernel. On conventional data, engineered relabelings are needed to elicit separation.

Encoding and Fourier expressivity. Schuld, Sweke and Meyer [2] show a data-encoding circuit realizes a truncated Fourier series whose accessible frequencies are fixed by the encoding; depth-one diagonal embeddings give a fixed, low-order trigonometric basis. This is why the dequantized local readout is a specific classical inductive bias rather than an arbitrary one.

Inductive bias and concentration. Kübler, Buchholz and Schölkopf [3] argue quantum kernels generalize well only when the target aligns with their fixed spectrum. Thanasilp et al. [4] establish exponential concentration of fidelity (and sufficiently global) kernels. These are, respectively, the “why no signal” and “why no training” statements that our two limits make operational on one axis.

Benchmarking. Bowles, Ahmed and Schuld [5] document across many models and datasets that quantum models seldom beat classical baselines, and emphasize fair comparison. We adopt their discipline—one shared split and a single symmetric tuning protocol for every kernel—and add a structural account of *why*, together with a pre-training screen.

3 The Readout-Locality Axis

Fix an encoding $x \mapsto |\psi(x)\rangle$ on q qubits. For a Pauli string P let $\text{wt}(P)$ be its weight (number of non-identity factors). Define the locality- k feature map

$$\phi_k(x) = (\langle \psi(x) | P | \psi(x) \rangle : 1 \leq \text{wt}(P) \leq k), \quad K_k(x, x') = \langle \phi_k(x), \phi_k(x') \rangle. \quad (1)$$

At $k = 1$, ϕ_1 collects the single-qubit Pauli expectations: this is the projected quantum kernel [1]. At $k = q$, ϕ_q collects all Pauli expectations, and a short calculation gives $K_q(x, x') = 2^q |\langle \psi(x) | \psi(x') \rangle|^2 - 1$: the fidelity kernel, up to a fixed affine map. Sliding k therefore interpolates from the local PQK to the global fidelity kernel.

The two ends are pinned. For the depth-one IQP map $|\psi(x)\rangle$ (a layer of Hadamards followed by diagonal phases $\sum_i x_i Z_i + \sum_{(i,j) \in E} x_i x_j Z_i Z_j$), the $k = 1$ features admit the exact closed form

$$\langle X_k \rangle = \cos(x_k) \prod_{j:(k,j) \in E} \cos(x_k x_j), \quad \langle Y_k \rangle = -\sin(x_k) \prod_{j:(k,j) \in E} \cos(x_k x_j), \quad \langle Z_k \rangle = 0, \quad (2)$$

verified to $\sim 10^{-15}$ for $q = 2, \dots, 8$. Hence K_1 is exactly a classical RBF over an explicit trigonometric feature map: the local readout *dequantizes*, and entanglement enters only as fixed cross-feature cosines. At the other end, the fidelity of the same depth-one state is $\frac{1}{2^q} \sum_z e^{i\theta(x'-x)z}$, an exponential Ising-type sum that does *not* factor (the cross terms couple it)—the regime behind IQP sampling hardness. Thus classical-simulation cost genuinely rises along the axis even as the closed form trivializes the $k = 1$ end. The product of $|E_k|$ sub-unit cosines in (2) also shows directly that high-weight features shrink as q grows—the concentration that we make quantitative in Section 5.

Protocol. Every kernel is “an RBF over a feature map,” so we compare them through one symmetric routine: a single stratified outer split per condition, nested cross-validation selecting the RBF bandwidth on the training portion only over a shared grid, fixed C , identical inner splits. The classical baseline is the same routine with $\phi = \text{id}$ on the q -dimensional (PCA) representation that every kernel sees. The fidelity endpoint has no bandwidth and is tuned over C by the analogous nested protocol. We report per-fold paired t -tests; because k -fold accuracies are not independent these p -values are indicative and, if anything, optimistic, so we lean on effect sizes and their trends.

4 An Empty Advantage Window on Real Data

We sweep $k = 1, \dots, 4$ and the fidelity endpoint at a fixed $q = 6$ (top-six PCA components) on eight genuine-signal tabular datasets. Figure 1 summarizes the result: on every dataset, the accuracy advantage over the classical RBF is ≤ 0 at *every* locality, including the moderate- k readouts where one might hope advantage could hide. Deficits on the higher-signal datasets are large (breast-cancer -0.29 , sonar -0.28 , qsar -0.16 , ionosphere -0.15 , diabetes -0.12). Across all 39 (dataset $\times$ readout) cells, only two reach nominal significance, both on the lowest-signal datasets and at the false-positive rate expected at $\alpha = 0.05$. Simultaneously, feature concentration (off-diagonal

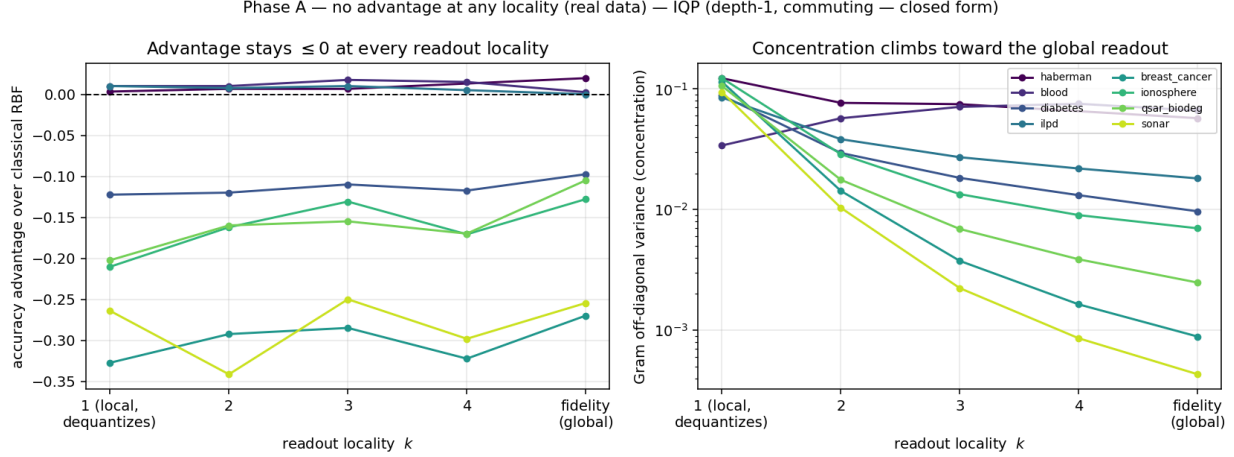


Figure 1: **An empty advantage window on real data (depth-one IQP).** Left: accuracy advantage over the identically tuned classical RBF stays ≤ 0 at every readout locality, on all eight datasets. Right: concentration climbs monotonically toward the global (fidelity) readout. No locality yields advantage.

variance of the unit-diagonal Gram) falls by roughly two orders of magnitude from $k = 1$ to the fidelity endpoint—the concentration limit made visible. No locality yields advantage.

5 The Obstruction Strengthens with System Size

The Phase-A null could be read as “these datasets lack the structure the kernel encodes.” The decisive test is therefore to *plant* that structure and ask whether the kernel can use it. We generate labels from a single high-weight observable $\langle W \rangle_x$ of the encoded state: a weight-3 Pauli string (representable only once $k \geq 3$) and the global weight- q parity (representable only by the fidelity readout). For each we sweep q and measure, through the same fair protocol, the advantage of the readout that *can* represent the signal over the classical RBF. Note the change of axis from Section 4: here the readout locality is held *fixed* at the minimum that represents the planted signal ($k = 3$ for the weight-3 string, the fidelity readout for the global parity), and it is the qubit count q —the horizontal axis of Figure 2, not the locality k of Figure 1—that varies.

Figure 2 shows the central phenomenon. The weight-3 signal yields a real, significant advantage at small q that *decays out of significance as q grows* (+0.138 at $q=4$, $p < 0.01$, to +0.059 at $q=8$, not significant, $p = 0.13$; 10 seeds); the global signal’s alignment gap Δ decays toward zero likewise (+0.023 to -0.001). The mechanism is explicit in the third panel: the planted feature’s variance $\text{Var}\langle W \rangle$ and the kernel’s concentration both collapse geometrically in q (straight lines on a log axis), exactly as the product of sub-unit cosines in (2) predicts. The only readout that can represent high-weight signal concentrates it away, and faster the more global the signal. The advantage window does not merely fail to open at a fixed size—it *narrows as system size grows*, which is the regime in which any asymptotic quantum advantage would have to appear.

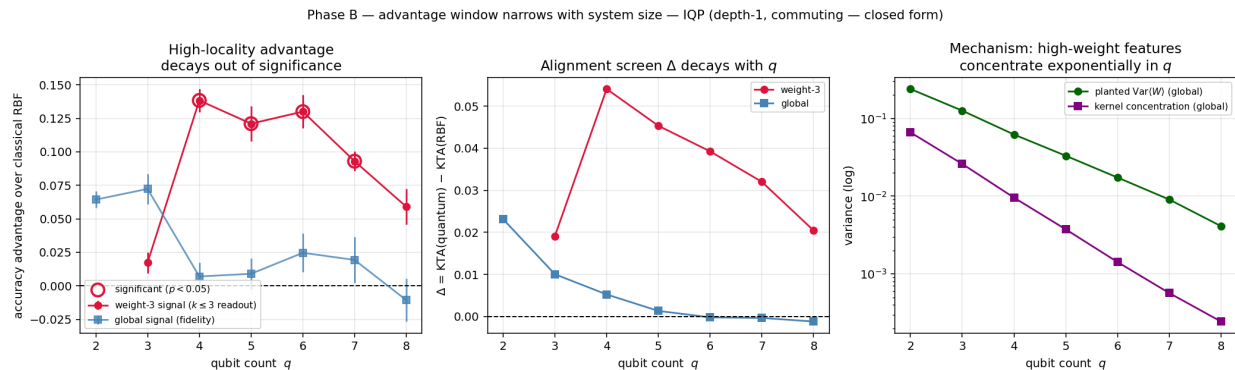


Figure 2: **The obstruction strengthens with system size (depth-one IQP).** Signal planted at the high-locality (global) limit. Left: weight-3 accuracy advantage decays out of significance as q grows. Middle: the alignment screen Δ decays for both planted signals. Right: planted-feature variance and kernel concentration collapse geometrically in q —the mechanism.

6 Robustness to a Non-Commuting Encoding

The dequantization anchor (2) relies on the commuting, diagonal structure of the IQP map. The natural way to seek a usable window is to break that structure. We therefore repeat the entire study with a non-commuting, data-re-uploading encoding—alternating the diagonal IQP block with a layer of data-dependent R_Y rotations—which has *no* known closed form, so its local readout is not provably classical. If any encoding within reach could open the window, this is a natural candidate.

It does not. Figure 3 overlays the two encodings: on real data, the best-over-locality advantage remains ≤ 0 for the non-commuting map on every dataset, and the planted-signal advantage still decays with q . Breaking the closed form removes our *proof* that the local end is classical, but it does not produce advantage at any locality on real data, nor does it stop the high-locality advantage from concentrating away as system size grows. The obstruction is unchanged.

7 Across Encoding Families

The non-commuting map is one alternative; to test whether the empty window is specific to any structural feature, we repeat the real-data study (Section 4) across six encoding families spanning the relevant axes: *angle* (single-qubit R_Y , no entanglement); IQP at *depth one, two, and three*; the non-commuting *IQP+ R_Y* re-uploading map; and a *hardware-efficient* map whose entanglement comes from a CNOT ring rather than diagonal phases. Depth-two and depth-three IQP already fall outside the closed form—their $k = 1$ readout deviates from the depth-one cosine map (kernel alignment 0.77 at depth two)—so, like IQP+ R_Y , they are genuinely closed-form-free.

Figure 4 shows the best-over-locality advantage per dataset for each family. None opens a window: on the higher-signal datasets every family is ≤ 0 at its best locality, and across the 39 (dataset $\times$ locality) cells per family only 2–4 reach nominal significance—the false-positive rate expected at $\alpha = 0.05$ —with the largest advantage anywhere being +0.03. The no-go is therefore not an artifact of commutativity, of the diagonal entangling structure, or of the closed form: it holds empirically across families where dequantization cannot be proven. The scale result of Section 5 reproduces on the closed-form-free depth-two map as well—the planted weight-3 advantage decays from +0.19 ($q=4$, $p < 0.01$) to +0.03 ($q=8$, not significant)—so the obstruction also *strengthens*

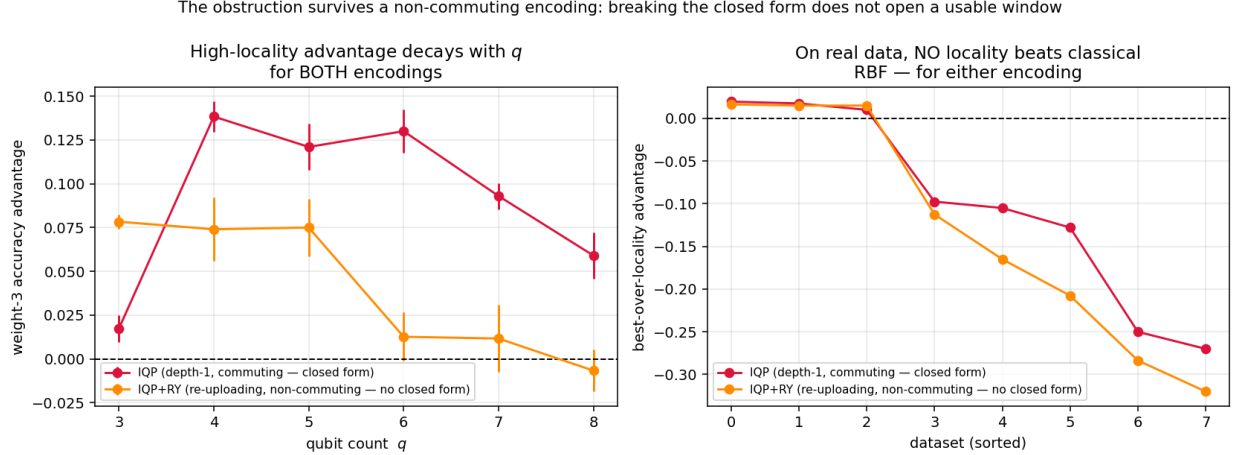


Figure 3: **The obstruction survives a non-commuting encoding.** Left: weight-3 advantage decays with q for both the commuting IQP map and its non-commuting re-uploading extension. Right: on real data, no locality beats the classical RBF for either encoding.

with system size beyond the case where dequantization can be proven.

8 A Pre-Training Screen for Any Encoding

That advantage appears only where the kernel is classically cheap (low k) or concentrated (high k) is, by itself, a statement about two encodings. To make it a protocol, we use the centered kernel-target alignment gap

$$\Delta(k) = \text{KTA}(K_k, y) - \text{KTA}(K_{\text{RBF}}, y), \quad (3)$$

which requires only the Gram matrices and training labels—no SVM fit. Across the study Δ tracks where on the locality axis measured advantage appears: it is at most marginally positive at the localities and scales where the planted signal is representable and unconcentrated, and ≤ 0 everywhere on real data, mirroring the accuracy curves. A practitioner with a candidate encoding can therefore sweep k , compute $\Delta(k)$ from labels alone, and read off whether *any* locality admits a window—before committing quantum resources. The obstruction is thus a measurable property of an encoding, not a property of one circuit we happened to study.

Δ predicts advantage without a closed form. The screen was motivated on the depth-one IQP kernel, which dequantizes. It is only useful if it also tracks advantage on kernels we *cannot* reduce to a known classical object. The non-commuting R_Y kernel (Section 6) has no closed form and is exactly that test. Pooling all measured conditions for each encoding—the per-dataset locality cells and the planted-signal points— Δ correlates with the measured accuracy advantage at Spearman $\rho = 0.81$ ($p = 6 \times 10^{-13}$, $n = 52$) for the non-commuting kernel, comparable to $\rho = 0.87$ for IQP (Figure 5). The screen is therefore not an artifact of dequantization. (At $q \leq 8$ every kernel here remains classically simulable; “no closed form” is the meaningful step we can take, not genuine intractability.)

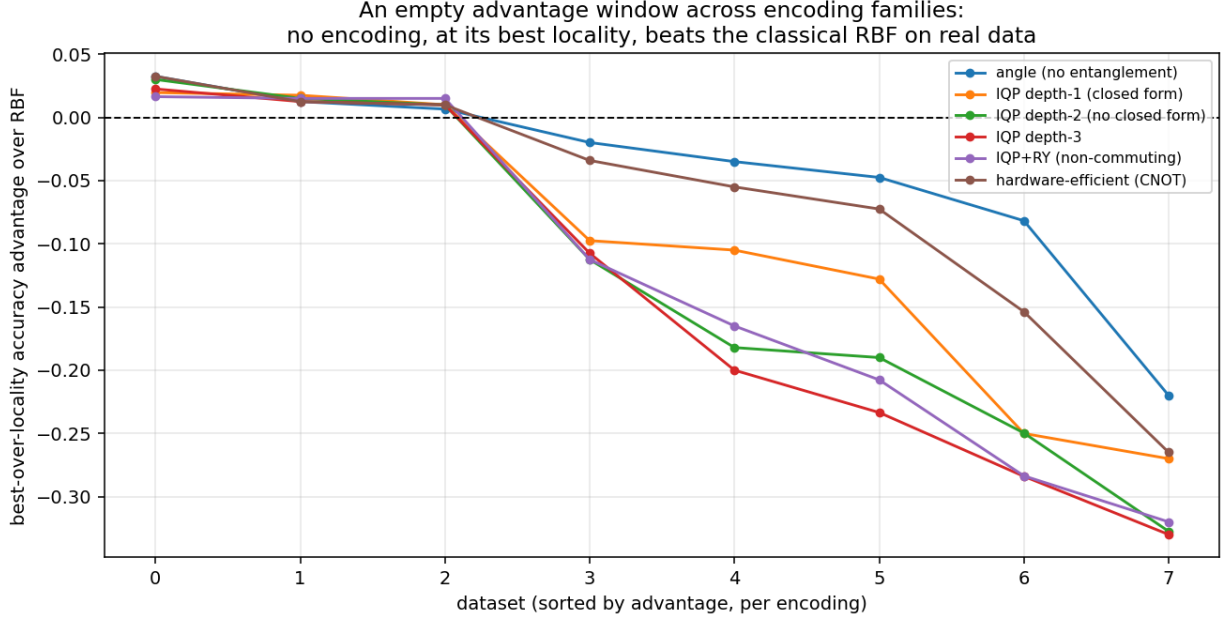


Figure 4: **An empty advantage window across encoding families.** Best-over-locality accuracy advantage over the classical RBF, per dataset, for six encoding families (including three with no closed form). No family beats the classical baseline at any locality on the higher-signal datasets; the few positive points are at noise level on the lowest-signal datasets.

9 The Obstruction Extends to Trained Models

A natural objection is that we study *kernels*, whereas a trained variational model might evade the concentration. It does not, on this encoding. Fixing the depth-one IQP map and appending a hardware-efficient variational ansatz $V(\theta)$, we measure the gradient variance $\text{Var}[\partial_\theta C]$ of the resulting cost as a function of q (Figure 6). For both a global cost ($O = Z_0 \cdots Z_{q-1}$) and a local cost ($O = Z_0$) the gradient variance decays geometrically in q , at a rate comparable to the global-kernel concentration measured in Section 5—a barren plateau. Notably the usual escape (adopt a local cost to avoid the plateau [6]) does *not* help here: because the IQP map is all-to-all entangling, even a 1-local observable is fed by a globally scrambled state, so the local cost concentrates too. The same concentration that empties the kernel’s advantage window also obstructs gradient-based training on the encoding; the result is not specific to the kernel method.

10 Scope and Limitations

Our no-go is demonstrated, not proven universal. (i) The dequantization anchor is exact only for the commuting depth-one IQP family; the broader encoding sweep (Section 7) is empirical, as it must be without a closed form. We claim the window is closed for the families tested and provide the test to check others, not that no encoding can open it. (ii) The “real data has no representable signal” half is specific to classical tabular data; quantum-native data, where provable advantage is known [1], is out of scope—the present claim concerns classical data fed to quantum kernels. (iii) We study quantum *kernels*; Section 9 shows the same concentration also obstructs gradient training of a variational model on this encoding, but a full optimization study (landscapes, plateau-

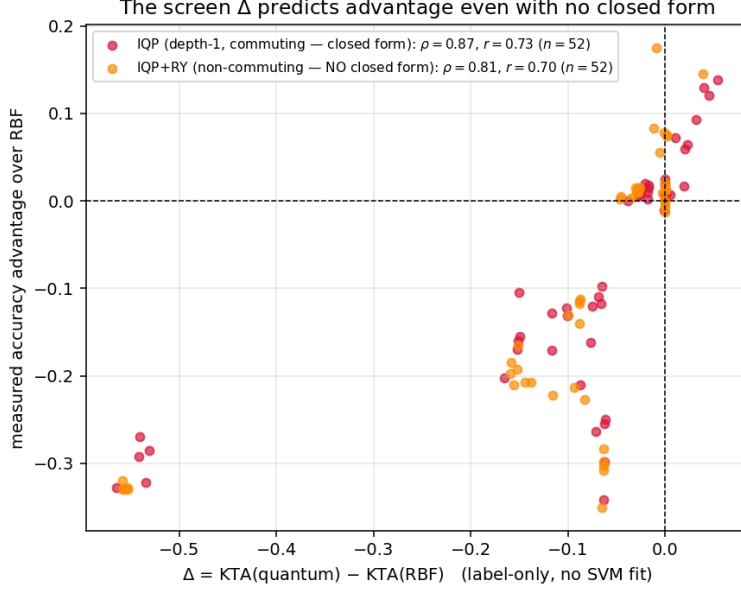


Figure 5: **The label-only screen Δ predicts measured advantage even with no closed form.** Each point is one measured condition; Δ uses only the Gram matrices and labels. The non-commuting R_Y kernel (orange, no closed form) is tracked as well as the closed-form IQP kernel (red).

mitigation strategies) is beyond our scope. (iv) The Δ screen is validated here including on a kernel with no closed form (Section 8); because every encoding at $q \leq 8$ remains classically simulable, its behavior on genuinely super-polynomial kernels is the natural next test.

11 Conclusion

The two established obstructions to quantum kernel advantage—dequantization of local readouts and concentration of global readouts—are the two ends of a single readout-locality axis, and for the encodings studied the interval between them holds no usable advantage on tabular data. The null is not fragile: it survives moving to a non-commuting, closed-form-free encoding, and it *strengthens* as system size grows, because the only readout that can represent genuinely high-weight signal concentrates it away geometrically in q . We reposition the question from “does this quantum kernel win?” to “where on the locality axis can advantage appear?”—and provide, in Δ , a cheap pre-training test for answering it on any future encoding. Where advantage might still appear—encodings whose dequantization limit extends to higher locality while their concentration limit recedes, or quantum-native data—is now a sharply localized target rather than an open benchmark.

References

- [1] H.-Y. Huang, M. Broughton, M. Mohseni, R. Babbush, S. Boixo, H. Neven, and J. R. McClean, “Power of data in quantum machine learning,” *Nature Communications* **12**, 2631 (2021).
- [2] M. Schuld, R. Sweke, and J. J. Meyer, “Effect of data encoding on the expressive power of variational quantum-machine-learning models,” *Physical Review A* **103**, 032430 (2021).

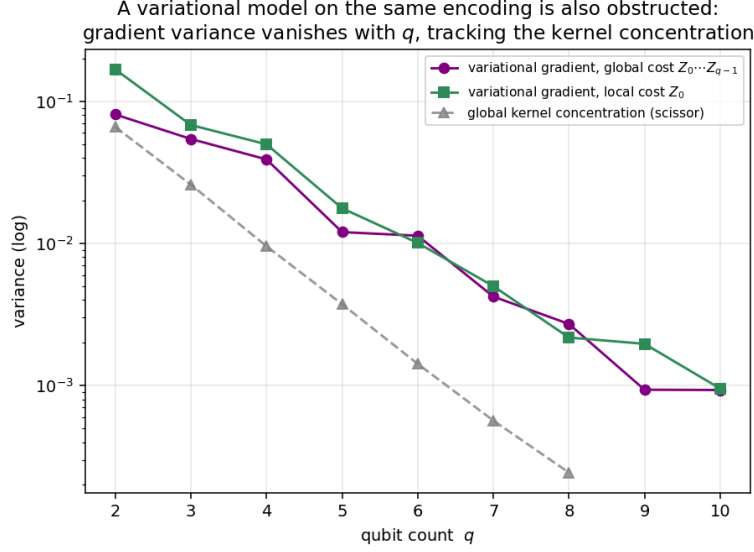


Figure 6: **A variational model on the same encoding is also obstructed.** Gradient variance of a hardware-efficient ansatz on the fixed IQP encoding vanishes geometrically in q for both global and local costs, tracking the global-kernel concentration (dashed). The globally entangling encoding leaves no trainable local escape.

- [3] J. M. Kübler, S. Buchholz, and B. Schölkopf, “The inductive bias of quantum kernels,” in *Advances in Neural Information Processing Systems (NeurIPS)* **34** (2021).
- [4] S. Thanasilp, S. Wang, M. Cerezo, and Z. Holmes, “Exponential concentration in quantum kernel methods,” *Nature Communications* **15**, 5200 (2024).
- [5] J. Bowles, S. Ahmed, and M. Schuld, “Better than classical? The subtle art of benchmarking quantum machine learning models,” arXiv:2403.07059 (2024).
- [6] M. Cerezo, A. Sone, T. Volkoff, L. Cincio, and P. J. Coles, “Cost function dependent barren plateaus in shallow parametrized quantum circuits,” *Nature Communications* **12**, 1791 (2021).