

When Will a Projected Quantum Kernel Help?

A Predictable, Mechanistic, and Structural Answer for Tabular Data

Author

June 7, 2026

Abstract

Projected quantum kernels (PQKs) avoid the exponential concentration of fidelity kernels by reading out only local observables, and are central to the hope that quantum kernels might one day help. We show that for the standard depth-one IQP feature map, whether a PQK can beat a classical kernel is (i) mechanistically fixed, (ii) cheaply predictable before any training, and (iii) structurally unattainable on tabular data. *Mechanistically*, the projected features admit an exact closed form— $\langle X_k \rangle = \cos(x_k) \prod_{j \neq k} \cos(x_k x_j)$ and likewise for $\langle Y_k \rangle$, for any number of qubits—so the PQK is identical, to machine precision, to a classical radial basis function (RBF) kernel over an explicit trigonometric feature map; entanglement contributes only fixed cross-feature cosines. *Predictively*, a single pre-training quantity, the centered kernel-target alignment gap $\Delta = \text{KTA}(\text{PQK}) - \text{KTA}(\text{RBF})$, predicts the measured accuracy advantage across 18 real and synthetic conditions (Spearman $\rho = 0.92$), whereas the geometric difference does not ($\rho = -0.01$); on eight tabular datasets, evaluated under one shared split and a single symmetric tuning protocol, $\Delta \leq 0$ and the PQK never wins. *Structurally*, optimizing feature selection to force $\Delta > 0$ succeeds but necessarily lowers accuracy, because the only PQK-favorable representations are those in which the classical kernel—and hence the available signal—is weak. The failure of the PQK on tabular data is therefore measurable in advance and intrinsic, not a tuning artifact.

1 Introduction

Quantum kernel methods embed inputs x into a state $|\psi(x)\rangle$ and classify with a support vector machine (SVM) on the resulting kernel. The *fidelity* kernel $|\langle \psi(x) | \psi(x') \rangle|^2$ concentrates exponentially with qubit count and becomes uninformative [4]; the *projected* quantum kernel (PQK) was introduced to mitigate this by reading out only local reduced observables and taking a classical RBF over them [1], and it is the construction most relevant to practical claims of advantage.

Rather than add another benchmark showing that quantum kernels seldom win on tabular data—already established [1, 5]—we ask a sharper question: *can we say in advance whether a PQK will help, explain why, and determine whether the answer can be changed?* For the standard depth-one IQP feature map we answer all three. Our contributions:

1. **Mechanism (Section 3).** An exact, arbitrary-qubit closed form for the depth-one IQP projected features. The PQK is therefore a classical RBF over an explicit trigonometric feature map (verified to $\sim 10^{-15}$); its inductive bias is a fixed, low-order trigonometric basis, and entanglement enters only as fixed cross-feature cosines.

2. **Prediction (Section 6).** A single pre-training quantity, the alignment gap $\Delta = \text{KTA}(\text{PQK}) - \text{KTA}(\text{RBF})$, predicts the measured advantage (Spearman $\rho = 0.92$ over 18 real and synthetic conditions), where the geometric difference of [1] does not ($\rho = -0.01$). It needs only the kernels and training labels—no SVM fit—and is estimable from a small subsample.
3. **Structure (Section 7).** Optimizing feature selection to force $\Delta > 0$ succeeds but necessarily reduces accuracy: the only PQK-favorable representations are signal-poor ones in which the classical kernel is weak. The mismatch is intrinsic, not a feature-engineering accident.

The mechanism (1) is specific to the commuting depth-one IQP map; the predictive (2) and structural (3) results concern any kernel pair and do not depend on it.

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 projected quantum kernel as a robust alternative to the fidelity kernel. They report that, on conventional data, engineered relabelings are needed to elicit any separation. We sharpen this: we observe large but inert g , and replace it with a quantity that actually predicts advantage.

Encoding and Fourier expressivity. Schuld, Sweke and Meyer [2] show that a data-encoding circuit realizes a truncated Fourier series whose frequencies are fixed by the encoding. Our closed form is a concrete instance: the depth-one IQP map produces exactly the frequencies $\{x_k\}$ and the pairwise products $\{x_k x_j\}$.

The inductive bias of quantum kernels. Kübler, Buchholz and Schölkopf [3] argue that quantum kernels generalize well only when the target aligns with the kernel’s spectrum, and otherwise require prohibitive data. We make this lens quantitative and predictive (a validated screening metric) and show its limit is structural via a feature-selection experiment.

Exponential concentration. Thanasilp et al. [4] establish that quantum kernels concentrate exponentially with system size. Our closed form exhibits the mechanism directly: a factor $\prod_{j \neq k} \cos(x_k x_j)$ of sub-unit terms drives the projected features toward zero as the qubit count grows.

Benchmarking. Bowles, Ahmed and Schuld [5] document across many models and datasets that quantum methods rarely match well-tuned classical baselines and that apparent advantages dissolve under matched evaluation. Our fairness protocol (shared splits, one symmetric tuning routine for every kernel) is in this spirit.

3 Mechanism: the IQP Projected Kernel Is a Classical Trigonometric Kernel

Given an m -dimensional input x , the depth-one IQP feature map prepares

$$|\psi(x)\rangle = U(x) H^{\otimes m} |0\rangle^{\otimes m}, \quad U(x) = \exp\left(-\frac{i}{2} \sum_k x_k Z_k - \frac{i}{2} \sum_{(a,b) \in E} x_a x_b Z_a Z_b\right),$$

with entangling graph E (the complete graph in the standard implementation). The PQK reads out the one-particle reduced Pauli expectations $\phi(x) = (\langle X_k \rangle, \langle Y_k \rangle, \langle Z_k \rangle)_k$ and sets $K_{\text{PQK}}(x, x') = \exp(-\gamma \|\phi(x) - \phi(x')\|^2)$.

Proposition 1 (Closed form of the depth-one IQP projected features). *For all m and all x ,*

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

Proof sketch. The circuit yields a uniform superposition with phases $\varphi(z) = \frac{1}{2}(\sum_k x_k z_k + \sum_{(a,b) \in E} x_a x_b z_a z_b)$, $z_k \in \{\pm 1\}$, so all amplitudes have equal modulus and $\langle Z_k \rangle = 0$. Flipping qubit k changes the phase by $\Delta_k(z) = x_k + \sum_{j:(k,j) \in E} x_k x_j z_j$, hence $\langle X_k \rangle = \mathbb{E}_z[\cos \Delta_k]$ over the neighbor bits. Applying $\mathbb{E}_z[\cos(a + \sum_j b_j z_j)] = \cos(a) \prod_j \cos(b_j)$ with $a = x_k$, $b_j = x_k x_j$ gives the result; the conjugate quadrature gives $\langle Y_k \rangle$. (The sine sign is a global reflection and does not affect the kernel.) $\square$

Corollary 1 (Dequantization). *The depth-one IQP PQK equals a classical RBF over the explicit feature map of Proposition 1. Entanglement enters only through the factors $\cos(x_a x_b)$, i.e. fixed pairwise cross-features. The product $\prod_{j \neq k} \cos(x_k x_j)$ of sub-unit terms drives ϕ toward 0 as m grows—the concentration mechanism of [4].*

We verified Proposition 1 against a state-vector simulator: across uniformly random inputs for $m = 2, \dots, 5$ the maximum discrepancy between the PQK Gram and the closed-form RBF Gram was $\leq 1.3 \times 10^{-15}$, and across the eight benchmark datasets below it was $\leq 3.2 \times 10^{-15}$. The quantum kernel’s inductive bias is thus an explicit, low-order trigonometric basis.

4 Experimental Setup

Pipeline and kernels. Each dataset is reduced per fold by standardization, mutual-information feature ranking, and PCA to q components (the qubits), with q chosen by a saturation criterion and floored at $q \geq 2$. Three kernels, all “RBF over a feature map” on the identical q -dimensional representation, are compared: (i) **classical RBF** over the PCA coordinates; (ii) **PQK**, the RBF over the depth-one projected IQP features; (iii) **Fourier twin**, the RBF over the closed-form features of Proposition 1.

Fairness invariants. A single stratified 10-fold split is built once per dataset and reused, with the identical preprocessed representation, by every kernel. All three are tuned by one nested cross-validation routine over the *same* bandwidth grid (12 points, log-spaced in $[10^{-3}, 10^2]$), the same $C = 1$, and the same inner folds; only the feature map differs. No kernel receives a tuning step or representation another lacks.

Alignment. The centered kernel-target alignment is $A(K, y) = \langle K_c, y_c y_c^\top \rangle_F / (\|K_c\|_F \|y_c y_c^\top\|_F)$, with K_c the doubly-centered Gram and y_c the centered ± 1 label vector; per kernel we take the best alignment over the shared bandwidth grid. It uses only the kernel and labels—no SVM.

Datasets and control. Eight pure-numeric binary datasets with genuine signal span 3–60 features and 208–1372 samples (Table 1). A synthetic control fixes the geometry and varies only the label structure: $y = \text{sign}(w^\top \phi(\omega R) - \text{median})$ with R a standardized q -dimensional representation and ω a frequency, so the target lies in the trigonometric (Fourier twin / PQK) RKHS by construction.

Dataset	$n \times d$	q	Acc RBF	Acc PQQ=twin	ΔAcc	KTA RBF	KTA PQQ
haberman	306×3	2	0.699	0.716	+0.017	0.059	0.037
blood	748×4	2	0.759	0.758	-0.001	0.090	0.062
diabetes	768×8	3	0.736	0.665	-0.070	0.153	0.038
ilpd	583×10	2	0.705	0.705	0.000	0.104	0.096
breast_cancer	569×30	2	0.939	0.859	-0.079	0.739	0.362
ionosphere	351×34	4	0.917	0.874	-0.043	0.319	0.315
qsar_biodeg	1055×41	2	0.766	0.672	-0.094	0.229	0.028
sonar	208×60	3	0.749	0.538	-0.210	0.176	0.072

Table 1: The PQQ never beats the classical RBF and equals its Fourier twin to machine precision (one column). $\Delta\text{Acc} = \text{Acc}_{\text{PQQ}} - \text{Acc}_{\text{RBF}}$; classical alignment $\geq$ quantum on all eight.

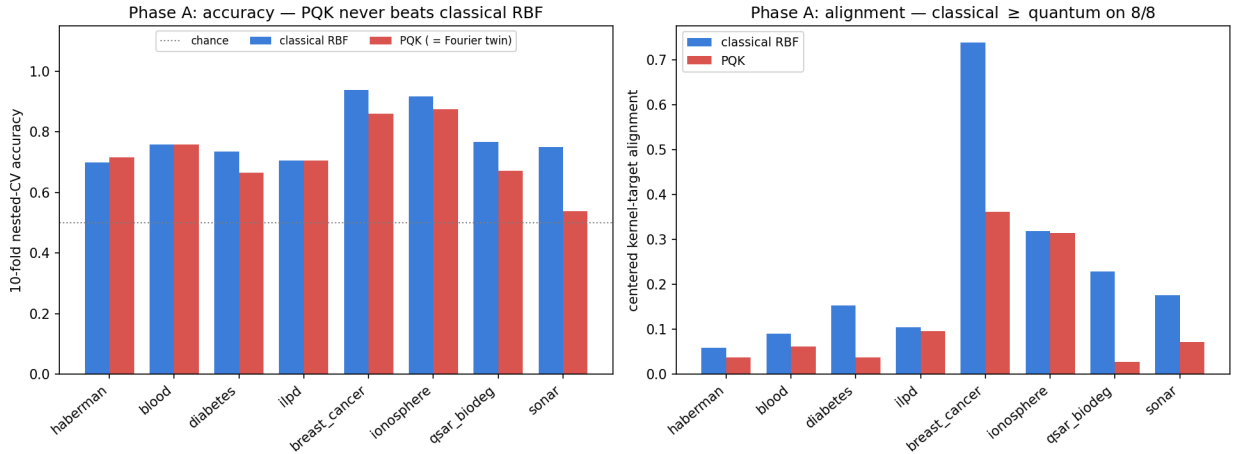


Figure 1: Left: nested-CV accuracy; the PQQ (= Fourier twin) never exceeds the classical RBF. Right: centered kernel-target alignment; classical $\geq$ quantum on all eight datasets.

5 No Advantage, and Alignment Explains It

On all eight datasets the PQQ and its closed-form twin produce identical Gram matrices (maximum difference 3.2×10^{-15} , mutual alignment 1.0000) and hence identical accuracies (Table 1, single PQQ/twin column). The PQQ never beats the classical RBF; deficits reach -0.21 (sonar). The classical kernel’s alignment meets or exceeds the quantum kernel’s on 8/8 datasets, and the accuracy deficit broadly tracks the alignment gap—large gap, large classical win (breast_cancer, qsar_biodeg); negligible gap, a tie (ilpd, blood) (Figure 1). The relationship is a trend with scatter (ionosphere loses despite a near-zero gap), which we report as observed.

6 A Predictive Suitability Metric

Calibration by a control. Whether the alignment gap can be made to matter is shown by the synthetic control. As the label frequency ω rises, a plain RBF over the coordinates degrades toward chance while the Fourier twin and the PQQ—which contain the relevant frequencies—retain accuracy (Figure 2, Table 2); the feature-map advantage grows from ≈ 0 to $+0.24$ (blood geometry) and $+0.37$ (Gaussian geometry), and the PQQ equals the Fourier twin to 0 throughout.

freq ω	blood geometry			Gaussian geometry		
	RBF	PQK=twin	Δ	RBF	PQK=twin	Δ
0.5	0.977	0.977	+0.000	0.975	0.972	-0.003
1.0	0.963	0.982	+0.020	0.915	0.930	+0.015
2.0	0.925	0.982	+0.057	0.810	0.965	+0.155
4.0	0.835	0.975	+0.140	0.698	0.927	+0.230
8.0	0.660	0.903	+0.243	0.565	0.935	+0.370

Table 2: Synthetic control: as ω rises the plain RBF collapses while the trigonometric map (Fourier twin, and the PQK identically) holds; the feature-map advantage Δ grows.

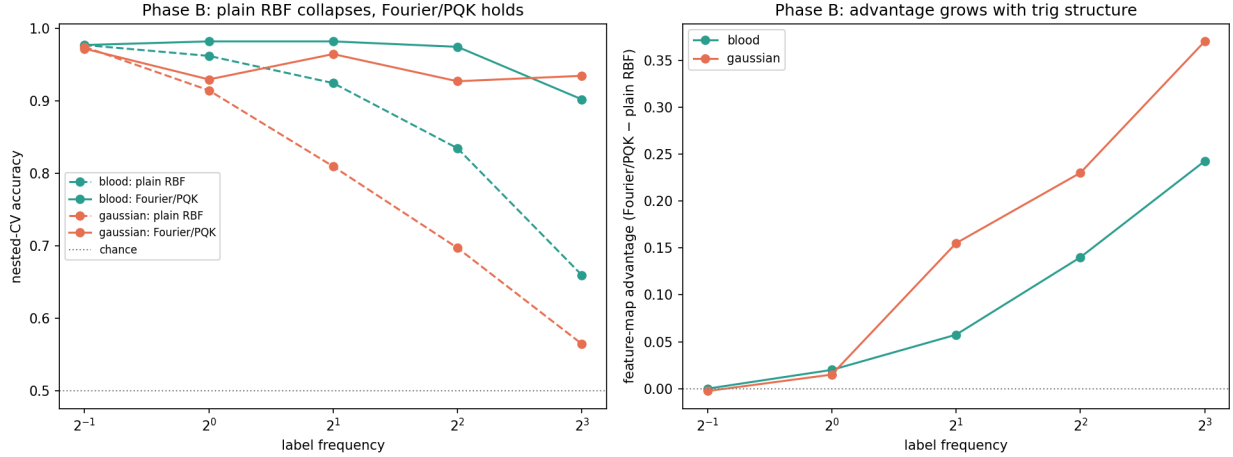


Figure 2: Control. Left: accuracy versus label frequency; the plain RBF (dashed) collapses while the Fourier/PQK kernel (solid) holds. Right: the feature-map advantage grows with trigonometric structure.

This populates the regime of *positive* quantum-minus-classical alignment that real data (Section 5) never reaches.

The metric. Define $\Delta = \text{KTA}(K_{\text{PQK}}, y) - \text{KTA}(K_{\text{RBF}}, y)$, computed from the kernels and *training* labels alone—a genuine pre-training screen. Pooling all 18 conditions (eight real datasets, ten control points), Δ predicts the measured nested-CV advantage with Spearman $\rho = 0.92$ ($p = 8 \times 10^{-8}$; Pearson $r = 0.82$, fit $R^2 = 0.66$), with an empirical zero-crossing near $\Delta = -0.04$ (Figure 3, left); a simple rule is to try the PQK only when $\Delta > 0$. By contrast the geometric difference g [1] is uncorrelated with advantage over the same conditions ($\rho = -0.01$, $p = 0.96$; Figure 3, right)—indeed within the control sweep g is non-monotone, near its minimum exactly where the PQK advantage is largest. The alignment gap is the sufficient-direction screen that the necessary-but-not-sufficient g is not. The Gram is the $O(n^2)$ cost; since Δ is only a screen it is accurately estimable from a random subsample of $m \ll n$ points (Appendix A).

7 Optimizing for Suitability: No Free Lunch

If Δ predicts advantage, can we engineer a representation that the PQK wins on? We test this leakage-free: per fold we select which q of the top- $K=6$ PCA components maximize the *train* gap Δ

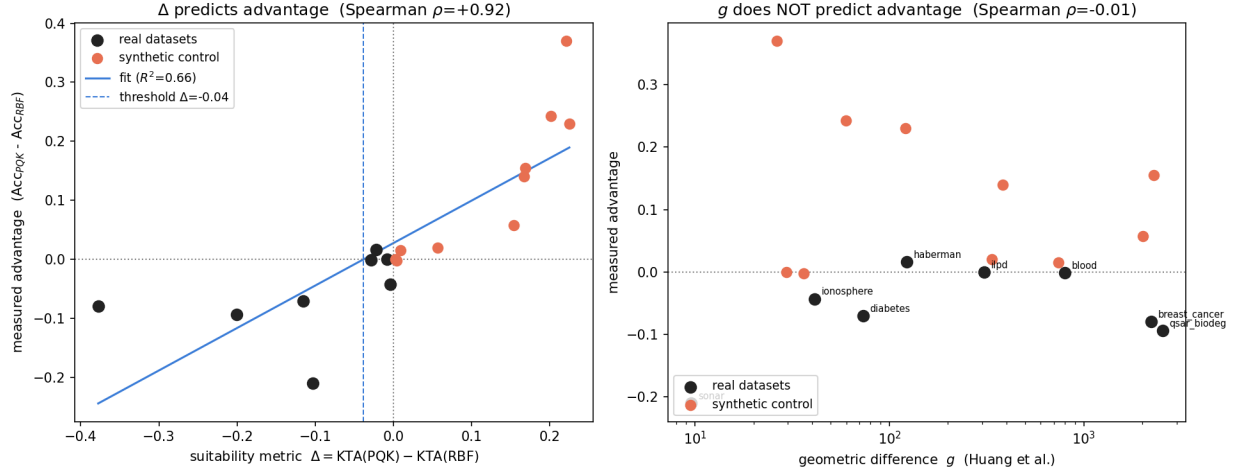


Figure 3: Left: the pre-training alignment gap Δ predicts the measured PQK advantage across real (black) and synthetic (orange) conditions (Spearman $\rho = 0.92$, fit $R^2 = 0.66$). Right: the geometric difference g [1] is uncorrelated with advantage ($\rho = -0.01$); the highest-advantage control point has near-minimal g .

Dataset	q	test Δ (def $\rightarrow$ opt)	PQK-RBF (def $\rightarrow$ opt)	best acc (def $\rightarrow$ opt)
blood	2	$-0.027 \rightarrow -0.020$	$+0.004 \rightarrow +0.007$	$0.774 \rightarrow 0.765$
diabetes	3	$-0.068 \rightarrow -0.003$	$-0.086 \rightarrow -0.046$	$0.737 \rightarrow 0.697$
ilpd	2	$-0.004 \rightarrow -0.003$	$+0.003 \rightarrow -0.005$	$0.715 \rightarrow 0.708$
breast_cancer	2	$-0.424 \rightarrow -0.001$	$-0.114 \rightarrow +0.000$	$0.942 \rightarrow 0.736$
ionosphere	4	$-0.049 \rightarrow +0.031$	$-0.077 \rightarrow -0.034$	$0.915 \rightarrow 0.906$
sonar	3	$-0.066 \rightarrow +0.003$	$-0.087 \rightarrow +0.005$	$0.712 \rightarrow 0.625$

Table 3: Δ -guided feature selection (leakage-free: select on train, score on test). The held-out gap rises and the PQK deficit closes, but best accuracy falls—the favorable regime is the low-signal one.

(exhaustive over $\binom{K}{q}$), then evaluate accuracy on the held-out fold, comparing against the default (top- q variance) selection.

The metric is a usable, generalizing target: train-selected subsets raise the *held-out* alignment gap on every dataset (breast_cancer $-0.42 \rightarrow 0$, ionosphere $-0.05 \rightarrow +0.03$, sonar $-0.07 \rightarrow +0.003$; Table 3, Figure 4 left), and the PQK deficit closes or slightly flips as Δ predicts (sonar test advantage $-0.087 \rightarrow +0.005$). But the gap turns positive only by selecting components on which the classical kernel is weak—that is, by discarding discriminative signal—so absolute accuracy falls, sharply where the default gap was most negative (breast_cancer $0.94 \rightarrow 0.74$, sonar $0.71 \rightarrow 0.63$; Figure 4 right). Optimized-PQK beats the *default classical* baseline on no dataset. PQK-favorable means signal-poor: the inductive bias mismatch is intrinsic, not a feature-engineering accident.

8 Conclusion

For the standard depth-one IQP projected quantum kernel, the question “will it help?” has a complete answer on tabular data. It is *mechanistic*: the kernel is exactly a classical RBF over an explicit trigonometric feature map (Proposition 1), with entanglement reduced to fixed cross-

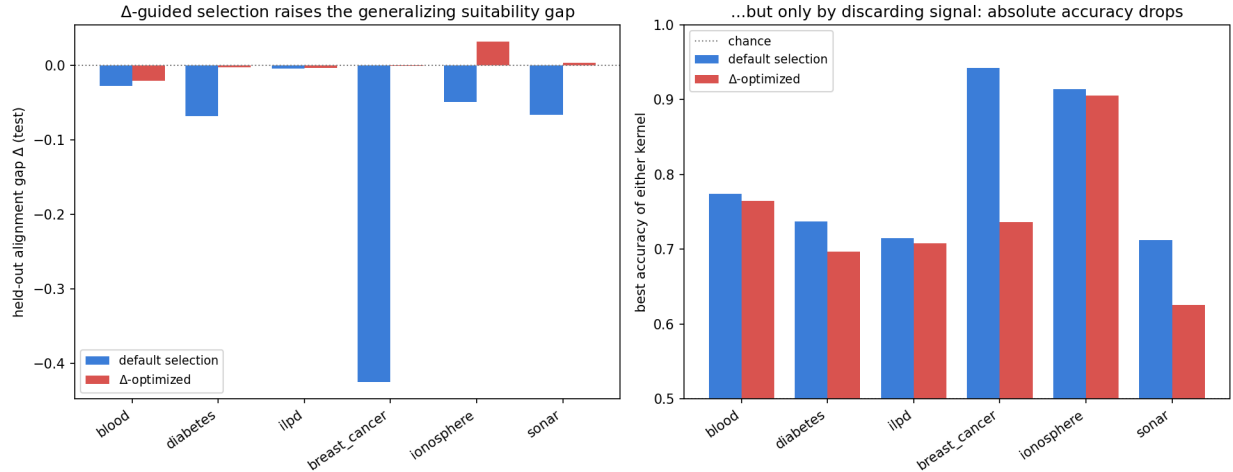


Figure 4: Left: Δ -guided selection raises the generalizing (held-out) suitability gap. Right: but only by discarding signal—the best accuracy of either kernel drops.

feature cosines. It is *predictable*: the pre-training alignment gap Δ forecasts the measured advantage (Spearman $\rho = 0.92$) where the geometric difference cannot, giving a cheap, label-only screen that flags every dataset in our study as unsuitable before any training. And it is *structural*: forcing $\Delta > 0$ by feature selection costs more accuracy than it buys, because PQK-favorable representations are signal-poor. These tie together prior accounts—the large-but-inert geometric difference [1], the Fourier-fixed expressivity [2], the alignment-limited generalization [3], the concentration visible in the closed form [4], and the empirical benchmarking pattern [5].

Practical implications. The results yield concrete guidance for anyone applying projected quantum kernels to tabular data.

1. **Screen with Δ before computing the kernel.** Estimate the alignment gap on a subsample (Appendix A: ~ 80 points, $\approx 1\%$ of the full Gram cost). A value $\Delta \leq 0$ is a validated, label-only “no” obtained before any training.
2. **Do not treat the geometric difference as evidence of advantage.** Large g is necessary but not sufficient and was uncorrelated with realized advantage here ($\rho = -0.01$); a result justified by large g alone is unsupported.
3. **Always benchmark against the dequantized twin under identical tuning.** Since the depth-one IQP PQK *is* a classical RBF over an explicit trigonometric map, any claimed quantum result should be compared to that matched classical kernel under one shared split and one tuning protocol; agreement (to machine precision, as here) indicates no quantum content.
4. **Do not feature-engineer toward suitability.** If $\Delta \leq 0$ on reasonable features it is intrinsic: forcing $\Delta > 0$ selects signal-poor representations and loses more accuracy than it gains (Section 7).

Limitations. The closed form relies on the commuting, diagonal structure of the depth-one IQP map; it does not cover non-commuting or data-reuploading encodings, where computing the pro-

jected features may itself be classically hard and the dequantization need not hold. The predictive metric and the no-free-lunch result do not depend on the closed form and apply to any kernel pair, but were evaluated on small tabular data and the projected (not fidelity) kernel. The natural next step is to test whether a non-commuting encoding either dequantizes via a higher-order Fourier map or resists dequantization while concentrating. This is the real frontier: a local-versus-global readout dilemma in which local readout dequantizes (as shown here) while global, fidelity-style readout concentrates [4]. Any genuine advantage must escape both.

A Cheap estimation of Δ by subsampling

The Gram matrix is the $O(n^2)$ bottleneck. Because Δ is only a screen, it can be estimated from a random subsample of $m \ll n$ points at cost $O(m^2)$. On the synthetic control the subsampled estimate $\hat{\Delta}$ converges to the full-data value with noise falling as $m^{-1/2}$, and the suitable cases ($\Delta \approx 0.2\text{--}0.3$) separate from the unsuitable case ($\Delta \approx 0$) by $m \approx 80$ —roughly 1% of the full kernel cost (Figure 5). At very small m the sample alignment underestimates its population value, so the screen errs toward declaring data unsuitable—the safe direction.

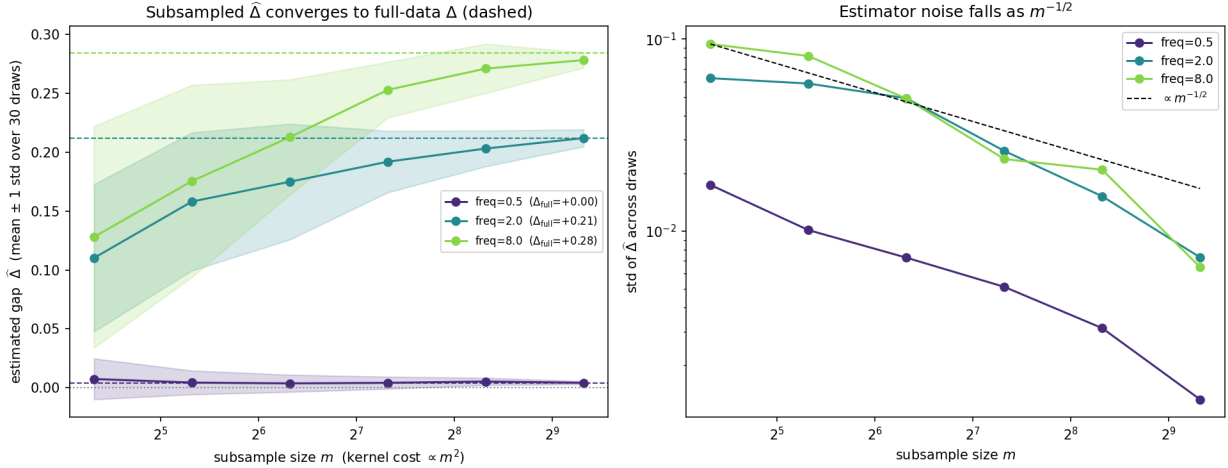


Figure 5: Subsampled estimation of Δ on the synthetic control. Left: $\hat{\Delta}$ from m points converges to the full-data Δ (dashed); suitable and unsuitable cases separate by $m \approx 80$. Right: the standard deviation of $\hat{\Delta}$ falls as $m^{-1/2}$.

B Advantage–alignment synthesis

Figure 6 places both phases on one axis (valid because the PQK equals its classical twin): real datasets sit at or left of zero alignment gap, and only manufactured trigonometric structure crosses into the positive, PQK-favorable quadrant.

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).

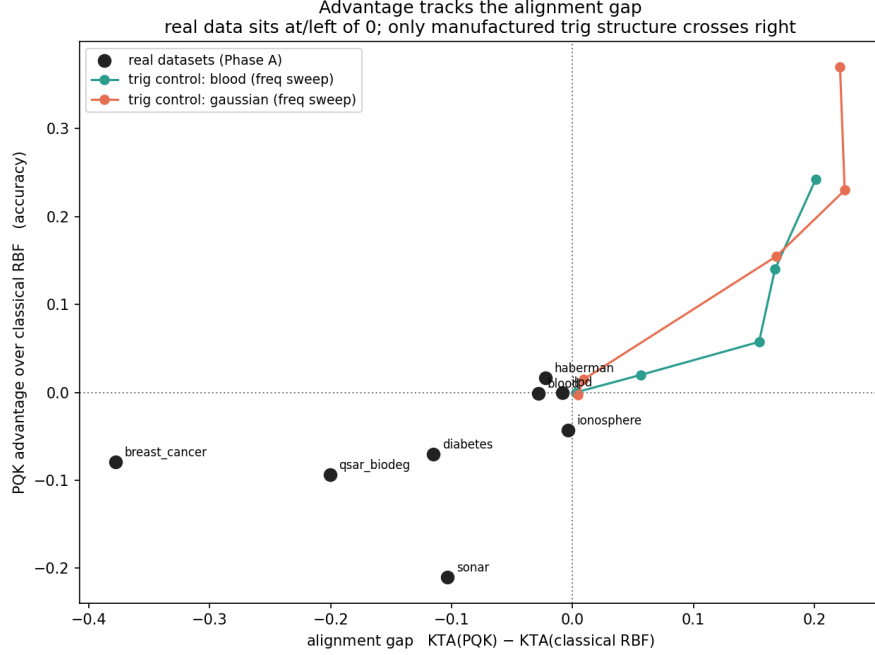


Figure 6: PQQ advantage versus the alignment gap $KTA(PQK) - KTA(RBF)$. Real datasets (black) lie at/left of zero; the trigonometric control sweeps (colored) cross into the positive quadrant.

- [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).
- [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).