

An Exact Classical Form for the Depth-One IQP Projected Quantum Kernel

Daniel Felps

June 10, 2026

Abstract

The projected quantum kernel is a widely cited near-term proposal for a quantum advantage in machine learning: each input is embedded in a quantum state, and the kernel is constructed from local measurements of that state. We show that, for the standard depth-one IQP embedding with single-qubit readout, this kernel admits an *exact closed form*: it is a classical kernel over the explicit trigonometric feature map $\langle X_k \rangle = \cos(x_k) \prod_{j \neq k} \cos(x_j x_k)$. The entanglement in the circuit contributes *only* the fixed product of cosines $\prod_j \cos(x_j x_k)$: a set of classical cross-feature interactions rather than a quantum resource. The derivation is elementary and we give it in full. The same expression has two further consequences. Because each added qubit multiplies in another sub-unit cosine, it explains why the kernel *concentrates*—its features shrink toward zero as the qubit count grows. Read as a Hilbert-space geometry, it shows that the projected features are *unnormalized* and that the embedding’s entangling frequencies are pairwise *products* of inputs, leaving the kernel sensitive to data scaling and non-uniformly fragile to perturbations in a way a classical RBF is not. We stress the scope of these claims: they concern this one instantiation—the depth-one IQP embedding with single-qubit readout—and not the projected-kernel framework in general, whose local-projection strategy remains a sound response to the concentration of fidelity kernels. Read constructively, the same closed form points to where a genuine advantage could still reside: non-commuting embeddings and higher-weight local projections.

1 Introduction

A common starting point in quantum machine learning is the following construction: map a classical input x into the exponentially large Hilbert space of q qubits via a state $|\psi(x)\rangle$, and let a standard kernel method operate in that space. Because the embedding can be hard to simulate classically, one expects the induced kernel to capture structure that no efficient classical kernel can. The *projected quantum kernel* (PQK) of Huang et al. [1] is the most prominent concrete realization of this idea designed to remain usable at scale, and it is widely used as a baseline “quantum kernel.”

This note addresses a narrow but sharp question about it: *for the standard depth-one IQP embedding with single-qubit readout, what does the quantum part actually contribute?* The answer is clean enough to write on one line. The projected features have an exact closed form,

$$\langle X_k \rangle = \cos(x_k) \prod_{j \neq k} \cos(x_j x_k), \quad \langle Y_k \rangle = \sin(x_k) \prod_{j \neq k} \cos(x_j x_k), \quad \langle Z_k \rangle = 0, \quad (1)$$

so the PQK *is* a classical kernel over the explicit feature map (1)—an identity, not an approximation; a state-vector simulation confirms it to machine precision (Section 4). The entanglement enters only through the factors $\prod_j \cos(x_j x_k)$ —fixed cross-feature cosines, i.e. ordinary classical feature

interactions. There is no quantum resource hidden in this kernel. A weaker, additive-error version of this classicality is essentially known: states prepared by Hadamards and diagonal phases are *computationally tractable* in the sense of Van den Nest [11], so their local expectations can be estimated classically by sampling (see the Remark in Section 4). The contribution of the closed form is what it adds beyond estimability: it is *exact*—the classical surrogate outperforms the quantum device rather than merely matching its shot noise—and it is *explanatory*, since the same expression that dequantizes the kernel also accounts for its concentration and its geometry. We emphasize at the outset that this is a statement about one specific instantiation—the depth-one IQP embedding read out through single-qubit observables—and not about projected quantum kernels in general: as Section 8 discusses, the local-projection idea behind the PQK remains well motivated, and our result narrows rather than dismisses it.

The paper is organized around that one equation. Sections 2–3 fix the objects the result needs—the fidelity kernel and its concentration problem, the projected kernel introduced to address it, and the IQP circuit—so that the derivation in Section 4 stands on its own. We then establish three distinct consequences of Equation (1). The closed form shows that the kernel is classical (Section 5); that, because each qubit introduces another sub-unit factor, it *concentrates* as the system grows (Section 6); and that, read as a Hilbert-space geometry, those same factors leave its features unnormalized and non-uniformly fragile, so that its inductive bias is sensitive to how the data are scaled (Section 7). A single expression thus accounts for the kernel’s classical simulability, its scaling limitation, and its sensitivity to data normalization. Section 8 turns these into implications and a way forward: it scopes the result to the depth-one IQP embedding with single-qubit readout, explains why the general projected-kernel framework is untouched, and identifies the embeddings and diagnostics under which projected kernels could still deliver a quantum advantage. A brief conclusion closes the paper.

2 Setup: the fidelity and projected kernels

2.1 Kernel Methods

A kernel method classifies by comparing inputs through a similarity $K(x, x') = \langle \phi(x), \phi(x') \rangle$, where the *feature map* ϕ sends each input into some (possibly very high-dimensional) vector space. The central trick is that one never needs ϕ explicitly: a support vector machine (SVM) finds a separating hyperplane using only the *Gram matrix* $K_{ij} = K(x_i, x_j)$ of pairwise similarities. Any symmetric, positive-semidefinite K is guaranteed to arise from *some* feature map (Mercer’s theorem), so one is free to design the similarity directly and let the feature space be implicit. The familiar radial basis function (RBF) kernel $K(x, x') = \exp(-\gamma \|x - x'\|^2)$, for instance, corresponds to a fixed, infinite-dimensional smooth feature map; the bandwidth γ sets how quickly similarity decays with distance. A kernel helps on a task precisely when its feature map places same-class points near one another—when its geometry aligns with the labels. For the kernel and SVM background assumed here, see Schölkopf and Smola [7] or Shawe-Taylor and Cristianini [8].

2.2 Quantum Feature Maps

A *quantum* feature map takes the feature space to be the state space of a quantum computer. It sends x to a state $|\psi(x)\rangle = U(x)|0\rangle^{\otimes q}$ prepared by a data-dependent circuit $U(x)$ on q qubits, starting from the all-zeros state $|0\rangle^{\otimes q}$. A q -qubit pure state is a unit vector in $\mathbb{C}^{2^q}$, so the feature space dimension grows *exponentially* in the number of qubits—the basis for the expectation that a quantum kernel might encode structure no efficient classical feature map can. (Here $\langle A \rangle \equiv \langle \psi | A | \psi \rangle$)

denotes the expectation value, or average measurement outcome, of an observable A .) The idea of computing kernels from quantum states is due to Havlíček et al. [5] and Schuld and Killoran [6]; for the broader area see the review/text [9]. Two ways to turn states into a kernel are standard.

2.3 The Fidelity Kernel and Its Concentration Problem

The most direct choice is the *fidelity kernel* $K_{\text{fid}}(x, x') = |\langle \psi(x) | \psi(x') \rangle|^2$, the squared overlap of the two states—a natural “how aligned are these two unit vectors” similarity, equal to 1 when the states coincide and 0 when they are orthogonal. It is a valid kernel and uses the full Hilbert space. Unfortunately it *concentrates*. The intuition is the curse of dimensionality: in an exponentially large space, two generic unit vectors are almost surely nearly orthogonal, so as q grows $\langle \psi(x) | \psi(x') \rangle$ becomes exponentially small for essentially all $x \neq x'$. The Gram matrix then approaches the identity—every distinct pair looks equally (and maximally) dissimilar—and the kernel carries no usable signal [4]. Training on such a kernel is infeasible, and distinguishing its near-equal entries requires exponentially many measurement shots on real hardware.

2.4 The Projected Quantum Kernel

Huang et al. [1] proposed to sidestep concentration by reading out only *local* information. For each qubit k one measures the three single-qubit Pauli observables $\langle X_k \rangle, \langle Y_k \rangle, \langle Z_k \rangle$. These are exactly the coordinates of the *Bloch vector* of qubit k ’s reduced state (the 2×2 density matrix obtained by averaging out all other qubits), so the triple $(\langle X_k \rangle, \langle Y_k \rangle, \langle Z_k \rangle)$ is a point in the unit ball $\mathbb{R}^3$ summarizing what qubit k looks like in isolation. Stacking these across qubits gives a classical vector

$$\phi_{\text{PQK}}(x) = (\langle X_k \rangle, \langle Y_k \rangle, \langle Z_k \rangle)_{k=1}^q \in \mathbb{R}^{3q}, \quad (2)$$

and defines the projected quantum kernel as an ordinary RBF over these projected features,¹

$$K_{\text{PQK}}(x, x') = \exp(-\gamma \|\phi_{\text{PQK}}(x) - \phi_{\text{PQK}}(x')\|^2). \quad (3)$$

Because ϕ_{PQK} has only $3q$ bounded coordinates (a number that grows *linearly*, not exponentially, in q), the kernel cannot collapse to the identity the way the fidelity kernel does—there is simply not enough room for all pairs to look mutually orthogonal. The PQK is therefore the framework’s natural response to concentration—more robust by construction, though Section 6 will qualify how far that goes—and the question addressed here is what its features compute.

3 The IQP embedding

To answer that we need to fix a circuit $U(x)$. We use the embedding most common in this setting, the depth-one Instantaneous Quantum Polynomial (IQP) map. It consists of a layer of Hadamards followed by data-dependent *diagonal* phase rotations:

$$|\psi(x)\rangle = \underbrace{\exp\left(-\frac{i}{2} \sum_{(j,k) \in E} x_j x_k Z_j Z_k\right)}_{\text{entangling phases}} \underbrace{\exp\left(-\frac{i}{2} \sum_j x_j Z_j\right)}_{\text{single-qubit phases}} H^{\otimes q} |0\rangle^{\otimes q}, \quad (4)$$

¹Huang et al. [1] write the kernel as $\exp(-\gamma \sum_k \|\rho_k(x) - \rho_k(x')\|_F^2)$ over the single-qubit reduced density matrices ρ_k . Since $\rho_k = \frac{1}{2}(I + \vec{r}_k \cdot \vec{\sigma})$ with Bloch vector $\vec{r}_k = (\langle X_k \rangle, \langle Y_k \rangle, \langle Z_k \rangle)$, one has $\|\rho_k - \rho'_k\|_F^2 = \frac{1}{2}\|\vec{r}_k - \vec{r}'_k\|^2$, so (3) is the same kernel with the bandwidth halved.

where Z_j and X_j, Y_j are the Pauli matrices acting on qubit j , and E is the set of qubit pairs that are entangled (we take the complete graph, i.e. all pairs, the common default). The factor $-\frac{1}{2}$ in the exponents is the rotation-gate convention of standard library implementations of this embedding; other conventions rescale or reflect the features but affect no kernel quantity in this paper (see the footnote in Section 4). Reading (4) right to left: the Hadamards put every qubit into an equal superposition; the single-qubit phases inject the features x_j ; and the two-qubit $Z_j Z_k$ phases inject the products $x_j x_k$ and are the *only* entangling part of the circuit. The name “instantaneous” refers to the fact that all the phase gates are diagonal in the computational basis and hence *commute*—a property we will use directly. Despite this simplicity, sampling from the output of IQP circuits is believed to be classically intractable [10]; as we show below, however, the projected (local) readout discards the global entanglement information responsible for that hardness, so the resulting feature map can be analyzed by entirely classical means.

To make the diagonal structure concrete, recall two facts. First, the Hadamard layer creates the uniform superposition $H^{\otimes q}|0\rangle = 2^{-q/2} \sum_{z \in \{0,1\}^q} |z\rangle$ over all 2^q computational basis states $|z\rangle$ (bit strings z). Second, a diagonal gate does not move amplitude between basis states; it only multiplies each $|z\rangle$ by a phase. To track those phases, write $s_j = (-1)^{z_j} \in \{+1, -1\}$ for the eigenvalue of Z_j on $|z\rangle$ ($s_j = +1$ if bit j is 0, and -1 if it is 1), since $Z_j|z\rangle = (-1)^{z_j}|z\rangle$. Substituting $Z_j \rightarrow s_j$ and $Z_j Z_k \rightarrow s_j s_k$ into the exponents of (4), the amplitude of $|z\rangle$ in $|\psi(x)\rangle$ is $2^{-q/2} e^{i\Phi(s)}$ with the real phase

$$\Phi(s) = -\frac{1}{2} \left(\sum_j x_j s_j + \sum_{(j,k) \in E} x_j x_k s_j s_k \right). \quad (5)$$

The key observation is that every basis state has the *same* amplitude magnitude $2^{-q/2}$; only the phase $\Phi(s)$ depends on the data. This equal-magnitude structure is what makes the closed form possible, and it is special to embeddings built entirely from Hadamards and diagonal gates.

4 The closed form

Proposition 1. *For the depth-one IQP embedding (4) with entangling set E , the projected features are exactly*

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

Derivation. Index basis states by their sign vector $s \in \{\pm 1\}^q$ as above, so $|\psi\rangle = 2^{-q/2} \sum_s e^{i\Phi(s)} |s\rangle$ with Φ from (5).

$\langle Z_k \rangle = 0$. Because Z_k is diagonal ($Z_k|s\rangle = s_k|s\rangle$), its expectation depends only on the amplitude *magnitudes*, and those are all equal:

$$\langle Z_k \rangle = \sum_s |2^{-q/2} e^{i\Phi(s)}|^2 s_k = 2^{-q} \sum_s s_k = 0,$$

because exactly half of the sign vectors have $s_k = +1$ and half -1 , so the $+1$ and -1 terms cancel. Geometrically, the equal superposition leaves each qubit’s Bloch vector in the XY -plane (the equator), with no Z -component—a fact we did not need to compute coordinate by coordinate.

$\langle X_k \rangle$. The bit-flip X_k sends $|s\rangle$ to $|s^{(k)}\rangle$, where $s^{(k)}$ is s with s_k negated. Hence

$$\langle X_k \rangle = \langle \psi | X_k | \psi \rangle = 2^{-q} \sum_s e^{-i\Phi(s)} e^{i\Phi(s^{(k)})} = 2^{-q} \sum_s e^{i[\Phi(s^{(k)}) - \Phi(s)]}.$$

Only the terms of Φ that contain s_k change under the flip: the single-qubit term $-\frac{1}{2}x_k s_k$ and the entangling terms $-\frac{1}{2}x_j x_k s_j s_k$ for $j \sim k$. Negating s_k negates each, so

$$\Phi(s^{(k)}) - \Phi(s) = s_k \left(x_k + \sum_{j:(j,k) \in E} x_j x_k s_j \right) \equiv s_k A(s).$$

Summing first over $s_k \in \{\pm 1\}$ (only the exponent's prefactor s_k varies, since A does not contain s_k) gives $\sum_{s_k} e^{is_k A} = 2 \cos(A)$; averaging over the remaining $q - 1$ signs—those with $j \not\sim k$ do not appear in A and average out trivially—leaves

$$\langle X_k \rangle = \mathbb{E}_s \left[\cos \left(x_k + \sum_{j \sim k} x_j x_k s_j \right) \right],$$

the expectation over i.i.d. uniform signs. Now use the elementary identity, valid for i.i.d. uniform signs $s_j \in \{\pm 1\}$,

$$\mathbb{E}_s \left[\cos \left(\alpha + \sum_j \beta_j s_j \right) \right] = \cos(\alpha) \prod_j \cos(\beta_j). \quad (6)$$

To see why, note that averaging a single sign gives $\mathbb{E}_{s_j} [e^{i\beta_j s_j}] = \frac{1}{2}(e^{i\beta_j} + e^{-i\beta_j}) = \cos \beta_j$; since the signs are independent the expectation of the product factorizes, $\mathbb{E}[e^{i(\alpha + \sum_j \beta_j s_j)}] = e^{i\alpha} \prod_j \cos \beta_j$, and (6) is the real part. With $\alpha = x_k$ and $\beta_j = x_j x_k$ this is exactly the stated form.² The computation for $\langle Y_k \rangle$ is identical except that Y_k flips $|s\rangle$ with an extra factor is_k , which turns the leading cosine into a sine. $\square$

Remark: what the collapse needs. The product form relies on the *pairwise* structure of the phases, not merely on their commutativity. Flipping s_k produced an argument $A(s)$ that is *linear* in independent signs, so the cosine average factorized through (6). With diagonal phases of degree three (terms $x_j x_k x_l Z_j Z_k Z_l$), $A(s)$ contains *products* of signs, and the same average becomes an Ising-type sum with complex couplings that no longer factorizes and whose exact evaluation is in general intractable. This does not, however, reopen the door to a quantum resource, because classical *estimability* survives at any degree. The derivation used only two facts—the amplitudes are flat, and a Pauli operator permutes basis states while multiplying by signs—so for any circuit of one Hadamard layer followed by polynomially many diagonal phase gates, every Pauli expectation is a uniform average of a bounded, efficiently computable function of the signs (for X_k , $\langle X_k \rangle = \mathbb{E}_s [\cos(\Phi(s^{(k)}) - \Phi(s))]$; the same bookkeeping covers any Pauli string), and sampling the signs estimates it to additive error ε from $O(1/\varepsilon^2)$ draws—the same precision a quantum device attains from $O(1/\varepsilon^2)$ measurement shots. This is not a new observation: such states are *computationally tractable* in the sense of Van den Nest [11], and the celebrated hardness of IQP circuits concerns *sampling* from their output distribution [10], not local expectation values. We record it only to scope the boundary: the exact closed form is special to pairwise phases, but classical reach extends, in the estimation sense, to the entire commuting family with readout of any fixed weight (Section 8). The two statements also calibrate one another: estimability dequantizes every diagonal-phase embedding weakly but explains nothing about the kernel it dequantizes; the closed form is exact rather than ε -accurate, free rather than $O(1/\varepsilon^2)$ samples per feature, and it—not estimability—yields the consequences of Sections 5–7.

²The closed form depends on the gate convention only through frequencies and reflections: full-angle gates $\exp(ix_j Z_j)$ and $\exp(ix_j x_k Z_j Z_k)$ give $\langle X_k \rangle = \cos(2x_k) \prod_{j \sim k} \cos(2x_j x_k)$ and $\langle Y_k \rangle = -\sin(2x_k) \prod_{j \sim k} \cos(2x_j x_k)$. A uniform frequency rescaling is the same map evaluated at rescaled inputs, and a sign flip of the whole Y block is an isometry of the feature space, so no kernel quantity in this paper depends on the choice.

Numerical check. Equation (1) can be verified directly: building the PQK (3) from a state-vector simulation and comparing it to an RBF over the classical features (1) gives identical feature matrices (entrywise agreement at 2×10^{-15} for $q = 2\text{--}5$) and identical Gram matrices. Figure 1(A) overlays every Gram entry for $q = 6$; the largest discrepancy is 3×10^{-15} (floating-point noise) and the kernel alignment is 1.0000. They are the same kernel.

5 Consequence 1: the kernel is classical

Read (1) as a feature map. Each qubit contributes the two coordinates

$$(\langle X_k \rangle, \langle Y_k \rangle) = m_k(x) (\cos x_k, \sin x_k), \quad m_k(x) = \prod_{j \neq k} \cos(x_j x_k),$$

that is, a one-dimensional *Fourier feature* $(\cos x_k, \sin x_k)$ scaled by the scalar $m_k(x)$. The Fourier part is exactly the truncated trigonometric expansion that any data-encoding circuit produces, with frequencies fixed by the encoding [2]. The modulation $m_k(x)$ —the *only* contribution of entanglement—is a product of cosines of the pairwise products $x_j x_k$: fixed, classical cross-feature interactions. Nothing in (1) requires a quantum computer to evaluate. The feature map consists entirely of elementary trigonometric functions—products of sines and cosines—so computing the kernel reduces to standard classical arithmetic.

Consequently the PQK (3) is a *classical RBF kernel over an explicit trigonometric feature map*. Its inductive bias is fixed in advance, with every frequency hard-coded by the encoding: it can help only when the labels happen to align with this particular basis of single-feature sinusoids and pairwise-product cosines, exactly as a classical kernel with that feature map would. This is the sense in which the kernel *dequantizes*. It also locates the role of entanglement precisely: not absent (the m_k are genuine data-dependent cross terms), but *classical*—a deterministic interaction map, not a resource a classical model cannot reproduce.

6 Consequence 2: it concentrates as it scales

The same product reveals a scaling pathology. The modulation $m_k(x) = \prod_{j \neq k} \cos(x_j x_k)$ is a product of $q - 1$ factors, each of magnitude at most one and typically strictly less. As the number of qubits grows, $m_k(x) \rightarrow 0$ for all but the most special inputs, so the projected features (1) shrink toward zero and the PQK Gram matrix drifts toward a constant—the same concentration that afflicts the fidelity kernel [4], now visible directly in a closed form rather than argued abstractly. The local readout mitigates concentration relative to the fidelity kernel but does not escape it: locality buys a milder rate, not immunity.

The rate is set by how spread the encoded angles are. Figure 1(B) plots the feature variance $\text{Var}_x[\langle X_0 \rangle]$ against q for two input scales. At unit scale the cosines stay near one and concentration is slow; at scale $\sigma = 2$ the features lose more than half their variance by $q = 6$ and keep falling. Thus the mechanism that the closed form makes explicit—multiplying in one sub-unit cosine per qubit—is also the mechanism that degrades the kernel as the device grows. Increasing the number of qubits leaves the construction exactly as classical while steadily reducing its ability to separate data points.

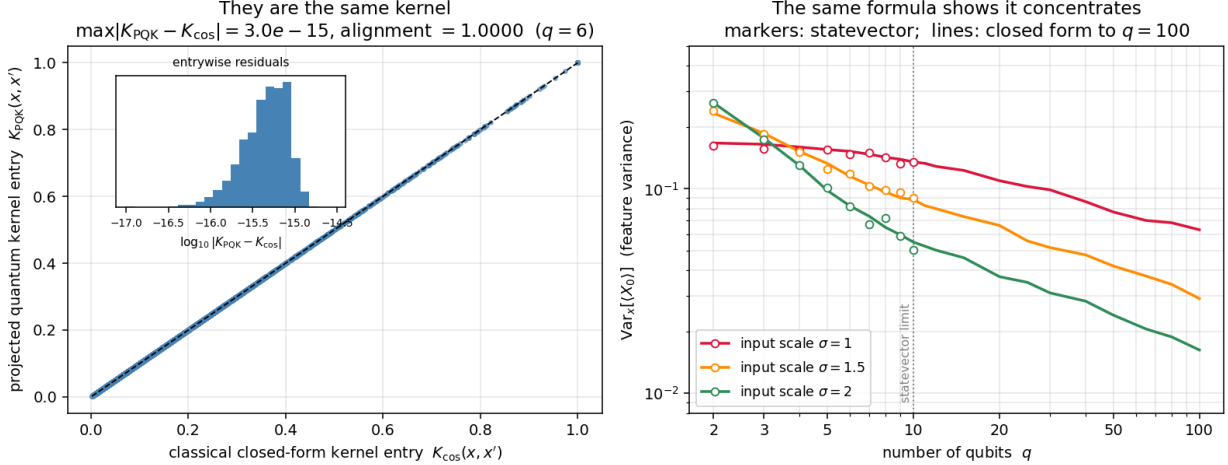


Figure 1: **The closed form: classical equivalence and concentration.** (A) The projected quantum kernel built from a simulated IQP circuit and the classical RBF over the closed-form features (1) produce identical Gram matrices: every entry lies on the diagonal, and the inset histogram of entrywise residuals shows pure floating-point noise ($\max|\Delta| = 3 \times 10^{-15}$, alignment 1.0000). The kernel is classical. (B) The same closed form predicts concentration: each qubit multiplies in another factor $\cos(x_j x_k) < 1$, so the feature variance $\text{Var}_x[\langle X_0 \rangle]$ falls as q grows, at a rate set by the input scale σ . Markers are statevector simulation ($q \leq 10$, the practical limit); lines are the closed form evaluated directly, at negligible cost, out to $q = 100$ —predicting the simulation and extending where simulation cannot go (slow decay at $\sigma = 1$, sharp at $\sigma = 2$).

7 Consequence 3: an unnormalized, scale-sensitive geometry

The first two consequences read (1) as a *feature map*. Reading it instead as a *Hilbert-space inner product* exposes a third lesson about how the kernel treats data—one that bears on whether a nearby test point is scored robustly. Because $\langle Z_k \rangle = 0$, the projected vector of (1) is, per qubit, $(\langle X_k \rangle, \langle Y_k \rangle) = m_k(x) (\cos x_k, \sin x_k)$ with the modulation $m_k(x) = \prod_{j \neq k} \cos(x_j x_k)$. The plain inner product of two such vectors collapses by the angle-difference identity to

$$\langle \phi(x), \phi(x') \rangle = \sum_k m_k(x) m_k(x') \cos(x_k - x'_k), \quad \|\phi(x)\|^2 = \sum_k m_k(x)^2, \quad (7)$$

the second formula being the $x' = x$ case. (Both match a state-vector simulation to $\sim 10^{-15}$.) Since the PQK (3) is a fixed decreasing function of $\|\phi(x) - \phi(x')\|^2 = \|\phi(x)\|^2 + \|\phi(x')\|^2 - 2\langle \phi(x), \phi(x') \rangle$, the pair (7) determines it entirely. Three properties of this geometry follow directly.

7.1 The Geometry Is Unnormalized

The fidelity kernel lives on unit states: $\|\psi(x)\| = 1$ for *every* x , so all inputs sit on one sphere. The projected vectors do not. By (7) the squared norm $\|\phi(x)\|^2 = \sum_k m_k(x)^2$ is data-dependent, sits well below its nominal maximum q (for standard-normal inputs at $q = 6$ it averages about 2), and—being a sum of the same shrinking products—drifts toward 0 as q grows, the concentration of Section 6 reappearing as a *vanishing norm*. Points therefore enter the kernel with unequal, data-dependent weight: an input with a single large product $x_j x_k \approx \pi/2$ has $m_k \approx 0$ and $m_j \approx 0$, so its j -th and k -th feature blocks all but vanish from the feature space. A classical kernel can be

normalized arbitrarily; here the magnitudes carry the entangling information, so normalizing them away is not without cost.

7.2 Normalization Selects the Operating Regime

The single-qubit phases encode frequencies x_k , but the entangling phases encode the *products* $x_j x_k$. Rescaling the inputs $x \mapsto cx$ —the most basic preprocessing decision—therefore shifts the single-feature frequencies by c but the cross-feature frequencies by c^2 . No global scale leaves the map neutral: for small c the cross cosines sit near 1, $m_k \approx 1$, and entanglement contributes essentially nothing (a pure Fourier map with no interactions); for c of order one the cosines oscillate and the interactions are strong but, as we show below, fragile; for large c the products $m_k \rightarrow 0$ and the kernel concentrates. For a classical RBF, rescaling the data and retuning the bandwidth are the *same* single knob— K depends only on $\gamma\|x - x'\|^2$ —so normalization is innocuous. The IQP encoding hard-codes its frequencies, so the data normalization scheme selects the operating regime, nonlinearly (c versus c^2), and a default scaler makes that choice implicitly. This is not a peripheral concern: for quantum kernels at large, the data scale (the kernel *bandwidth*) is known to decide between trivial and competitive generalization [14, 15]. The closed form shows *why* the IQP encoding is so exposed to that knob: the scale enters its interaction strength quadratically while entering its single-feature frequencies only linearly.

7.3 Spatial Variance in Robustness

Differentiating the modulation, $\partial \log m_k / \partial x_l = -x_k \tan(x_l x_k)$ for $l \neq k$: the sensitivity of feature k to a perturbation of feature l scales with the *partner* magnitude x_k and diverges as a product approaches $x_l x_k = \pi/2$, where its cosine crosses zero. Noise in one coordinate is thus amplified by the coordinates it multiplies, and points far from the origin sit on the steep flanks of the cosines and are fragile. The RBF again has no such position dependence: $\|\Phi_{\text{RBF}}(x) - \Phi_{\text{RBF}}(x + \eta)\|^2 = 2(1 - e^{-\gamma\|\eta\|^2})$ depends only on the perturbation η , never on where x is.

7.4 Numerical Experiment

We make this concrete with a noise-robustness test (Figure 2), comparing the projected geometry against a classical RBF *fairly*: we choose the RBF bandwidth γ so the two kernels have the *same* mean similarity between a clean point and a noisy copy at a reference noise level, so the comparison is about the *shape* of the degradation, not a tunable average rate. Panel (A) adds Gaussian noise of growing size to standard-normal inputs ($q = 6$) and plots the normalized Hilbert-space similarity $\hat{\kappa}(x, x + \eta) = \langle \phi(x), \phi(x + \eta) \rangle / (\|\phi(x)\| \|\phi(x + \eta)\|)$, the cosine of the angle between a point’s feature vector and its perturbation. (The PQK (3) is a fixed decreasing function of distances among these same feature vectors, so it inherits whatever inhomogeneity this geometry has; we measure the geometry directly.) The means are matched at the reference level and separate elsewhere—a single bandwidth can equate two differently shaped decay curves at only one point—with the PQK’s mean sitting higher at larger noise. That elevation is not resilience but a similarity *floor*: the feature distribution is anisotropic ($\|\mathbb{E} \phi\|^2 / \mathbb{E} \|\phi\|^2 \approx 0.42$), so even *unrelated* standard-normal inputs score $\hat{\kappa} \approx 0.40$ on average, where the matched RBF scores them ≈ 0.04 ; by $\sigma = 1$ a noisy copy (0.36) is barely more similar than a random stranger—the concentration of Section 6 reappearing as lost discrimination. The genuine deficit is the spread: the PQK’s 10–90 percentile band is at least $\sim 1.5\times$ wider than the RBF’s at every noise level, widening past $2\times$ at the largest, so a comparable *average* robustness hides a far less *uniform* one, and some test points are corrupted by noise that barely moves others. Panel (B) fixes the noise level and plots per-point similarity against the input

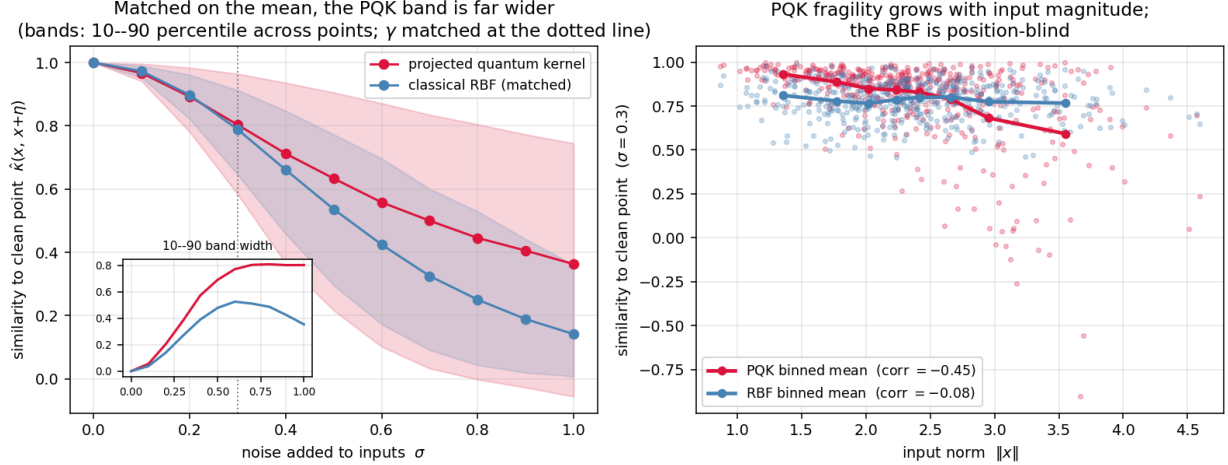


Figure 2: **Non-uniform robustness of the projected geometry.** The RBF bandwidth is matched so both kernels share the same *mean* robustness at the reference noise (dotted line). (A) Normalized Hilbert-space similarity between a clean point and a noise-perturbed copy versus noise size; lines are means, bands are the 10–90 percentile across points, and the inset plots the band width—the claimed quantity—directly. The means are matched at the reference noise and separate elsewhere (one bandwidth, one crossing); the PQK mean sits higher at large noise only because of its similarity floor—even unrelated inputs score ≈ 0.4 , versus ≈ 0.04 under the matched RBF—while its band is at least $\sim 1.5\times$ wider throughout: its robustness is inhomogeneous. (B) Per-point similarity at fixed noise versus the input norm $\|x\|$, with bold binned means over the scatter: PQK robustness decays with input magnitude (corr. -0.45), the RBF is position-blind (corr. -0.08). The entangling cross-frequencies $x_j x_k$ make fragility a function of where the data sit.

norm $\|x\|$: the PQK’s robustness falls markedly with $\|x\|$ (correlation -0.45) while the RBF’s is flat (-0.08). The fragility is set by the absolute input magnitude—which is precisely the normalization scheme—exactly as the sensitivity calculation predicts.

7.5 The Common Cause: Non-Stationarity

These symptoms have a single cause. A classical RBF is *stationary*: $K(x, x') = f(x - x')$ depends only on the displacement, so it treats every region of input space identically—which is exactly why its robustness is position-blind. The projected kernel is not. In (7) the factor $\cos(x_k - x'_k)$ is stationary, but the modulation $m_k(x) m_k(x')$ depends on x and x' separately and breaks translation invariance: a common shift $x \mapsto x + a$, $x' \mapsto x' + a$ moves the inner-product kernel (7) by an $O(1)$ amount—for standard-normal inputs, nearby pairs, and unit-normal common shifts at $q = 6$, by more than 5 against typical entry magnitudes of order one, where a stationary kernel would not move at all; Figure 3 shows the full distribution of this shift response. Non-stationarity is position-dependent treatment of inputs, so it is the very property that Figure 2(B) measures—the three symptoms above are one phenomenon viewed three ways; we formalize the shift measurement as a stationarity diagnostic in Section 8. Nor is non-stationarity generic to quantum kernels, or a consequence of their periodic features—stationary periodic kernels exist, and quantum encodings produce them. The entanglement-free angle encoding $\bigotimes_k R_Y(x_k)$ yields projected features $(\sin x_k, 0, \cos x_k)$ of unit norm per qubit, so its projected kernel depends only on $x - x'$ and satisfies $\Delta(a) = 0$ to machine precision. Two ingredients of the IQP map break this.

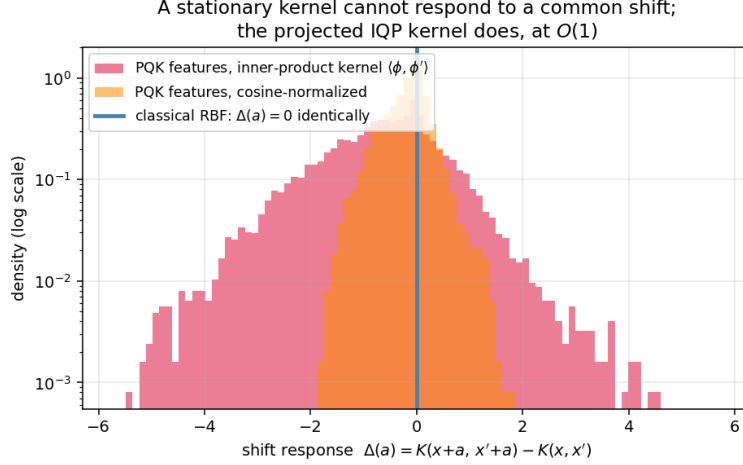


Figure 3: **Non-stationarity made visible.** Distribution of the shift response $\Delta(a) = K(x+a, x'+a) - K(x, x')$ over 500 standard-normal pairs ($x' = x + \eta/2$, $q = 6$) and 20 common unit-normal shifts a . A stationary kernel cannot respond to a common shift: the classical RBF sits in a spike at exactly zero (vertical line). The kernel of the projected features responds at $O(1)$ —reaching beyond ± 5 for the inner-product kernel (7) and ~ 1.8 after cosine normalization—which is the stationarity diagnostic of Section 8 applied to the depth-one IQP-PQK.

The entangling modulation $m_k(x)$ shrinks each Bloch vector by a data-dependent factor, which already makes the projected geometry position-dependent; and the entangling phases encode the data *nonlinearly*, through the monomials $x_j x_k$, so a translation of x does not even translate the encoded frequencies—the same root as the c -versus- c^2 coupling of Section 7.2. To be clear, non-stationarity is not a defect in itself: stationarity is an inductive bias—the right one when the task is translation-invariant, which is why the RBF is such a strong default on standardized data—and deliberately non-stationary classical kernels (polynomial, neural-tangent) earn their keep on tasks with position-dependent structure. The problem here is that the bias is *implicit*, selected by the data normalization rather than by design (Section 7.2), and arrives coupled to the fragility above. The same property also rules out the most obvious remedy. One might cosine-normalize the features to repair the vanishing norm; but the normalized kernel $\langle \phi(x), \phi(x') \rangle / (\|\phi(x)\| \|\phi(x')\|)$ remains non-stationary (it changes by up to ~ 1.8 under the same shifts; Figure 3), because the *direction* of $\phi(x)$ —the relative weights $m_k(x)/\|\phi(x)\|$ across qubits—is itself data-dependent. Normalization rescues the norm but not the scale-coupling or the fragility.

Read together with the first two consequences, the single product $m_k(x) = \prod_{j \neq k} \cos(x_j x_k)$ is responsible for all three properties: it makes the kernel classical (Section 5), causes it to concentrate (Section 6), and renders its geometry unnormalized, scale-coupled, and non-uniformly fragile (this section). The closed form exposes all three simultaneously.

8 Scope and outlook

The closed-form derivation provides a rigorous dequantization of the baseline projected quantum kernel under the depth-one IQP embedding. These findings should be scoped carefully, because the classical collapse demonstrated here is a property of that specific embedding rather than of the local-projection strategy itself. Separating the embedding from the framework reveals several

constructive implications for quantum machine learning. More broadly, the exercise illustrates a general program for assessing a proposed quantum kernel: derive its feature map exactly where the structure permits, and let the same closed form mark both the limits of the construction and the directions in which a genuine advantage could still lie.

8.1 The Enduring Necessity of Local Projections

The original motivation for the PQQ remains valid: the fidelity kernel concentrates exponentially (Section 2), and projecting the global state onto local observables—the single-qubit reduced density matrices (1-RDMs) read out by the PQQ—is a principled mechanism for preserving similarity structure at scale. That the depth-one IQP circuit fails to populate this projected space with hard-to-simulate correlations does not undermine the architectural role of local projections in scaling quantum kernels.

8.2 A Boundary for Genuine Quantum Advantage

Rather than invalidating the PQQ, the closed form acts as a filter on embedding design. Circuits built from a uniform superposition and *pairwise* diagonal phases—the depth-one IQP map—collapse to an exact classical product; and the sampling argument recorded in Section 4 extends the conclusion, in a weaker but broader form, to the whole commuting family: for diagonal phases of *any* degree, every local Pauli expectation remains classically estimable by sampling, to the same additive precision a quantum device achieves with finitely many shots (a special case of the simulability of computationally tractable states [11]). Within this family, single-qubit readout—and, since the sampling argument covers every Pauli string, any fixed-weight readout—therefore cannot yield a quantum prediction advantage. This is consistent with the finding of Huang et al. [1] that an advantage in their experiments required deeper, non-commuting embeddings (for example, Trotterized Hamiltonian evolution): a genuine geometric separation from classical models appears to require non-commuting, entangling circuits, whose amplitudes are not flat and which therefore resist both the algebraic reduction and the sampling argument. Expressivity is not the obstacle—in principle any kernel can be realized as an embedding quantum kernel [19]—so the binding question is which embeddings produce kernels that are simultaneously hard to simulate and matched to the data [13].

8.3 The Classical Surrogate Program

That a widely used “quantum” baseline is in fact a classical feature map is a useful outcome. As Huang et al. [1] note, assessing quantum advantage requires rigorous classical comparisons. Identifying the depth-one IQP-PQQ as a classical trigonometric kernel converts it from an expensive quantum simulation into an exact, efficient classical surrogate [3]: it can be evaluated classically at scale to benchmark data, reserving quantum hardware for the non-commuting embeddings that demonstrably require it. Related programs fit random-Fourier-feature surrogates to variational quantum models [17, 18]; here no fitting is needed—the surrogate is exact.

8.4 Pathways to Quantum Advantage in Projected Kernels

Two modifications to the construction are natural candidates for restoring a quantum contribution, and a third item gives a diagnostic for checking whether they succeed.

Non-commuting gates and circuit depth. Interleaving non-commuting rotations (for example R_X or R_Y gates) with the diagonal Z and ZZ phases breaks the commuting structure that drives the

cancellation in (6). Without that structure the equal-magnitude amplitude condition of Section 3 no longer holds, and the projected features need not collapse to a finite product of cosines—nor remain classically estimable by the flat-amplitude sampling argument of Section 4.

Multi-qubit projections. Reading out joint observables such as $\langle X_k X_j \rangle$, rather than only single-qubit expectations $\langle X_k \rangle$, prevents the partial trace from averaging away the entanglement generated by the circuit. Projecting onto higher-weight local observables retains more of the global correlation structure while still avoiding the exponential concentration of the full fidelity kernel. (For purely diagonal embeddings, however, even these higher-weight expectations remain flat-amplitude sign-averages and hence classically estimable—the Remark of Section 4—so this pathway is effective in combination with non-commuting structure, not as a substitute for it.)

A stationarity diagnostic. The shift response of Section 7 provides a cheap first screen. Sample input pairs (x, x') and common shifts a , and measure

$$\Delta(a) = K(x + a, x' + a) - K(x, x').$$

A stationary kernel satisfies $\Delta(a) = 0$ identically, so a candidate whose $\Delta(a)$ is concentrated near zero is, in practice, interchangeable with a stationary classical kernel and cannot offer the data anything an RBF could not. The converse direction must be read with care, because non-stationarity is not the boundary of classicality: the depth-one IQP-PQK itself has $\Delta(a)$ of order one (Figure 3) and is exactly classical, and stationarity is in any case an inductive bias rather than a virtue or defect—the right choice precisely when the task is translation-invariant (Section 7). A large $\Delta(a)$ therefore certifies only that the kernel treats input regions unequally, i.e. that it is not secretly an RBF; whether that non-stationary structure is classically reproducible is the question the second check answers. That check regresses the measured features onto the closed form (1) and its multi-qubit generalizations and reports the residual: a small residual indicates classical simulability, whereas a large, irreducible residual is a necessary condition for a genuine quantum contribution. The two tests thus order the diagnosis: $\Delta(a)$ asks whether the kernel differs from the stationary classical baseline at all; the regression asks whether that difference is anything a classical feature map cannot supply.

9 Conclusion

We asked what the quantum component of a projected quantum kernel computes, for the standard IQP embedding, and found an expression that can be written down explicitly: a classical trigonometric feature map whose only entangling content is a fixed product of cosines. The conclusion is not that quantum kernels are unhelpful in general, but that this particular and widely used construction is a classical kernel, and that its classical nature follows from an elementary algebraic reduction. The same expression also predicts its failure to scale and, read as a geometry, a third limitation: the projected features are unnormalized and their entangling frequencies are products of inputs, so the kernel’s behavior—and its robustness to noisy or shifted test points—depends strongly on the data normalization, in a way a classical RBF’s does not.

Two caveats delimit the scope. First, the exact closed form is special to the depth-one IQP map—one Hadamard layer, pairwise diagonal phases, single-qubit readout; deeper or non-commuting encodings need not admit one, though for purely diagonal encodings of any degree the features remain classically estimable by sampling (Section 4; cf. [11]). Second, the concentration shown here is for classical inputs of moderate scale; the precise rate depends on the data distribution and

the encoding bandwidth. The structural conclusion is robust: for this canonical embedding, the “quantum kernel” is a classical cosine kernel, and the cosines that make it classical are the cosines that make it concentrate. A practitioner applying a PQQ with this embedding to tabular data is, in effect, applying a fixed classical feature map, and is better served by selecting that feature map deliberately.

None of this implies that quantum kernels cannot help in principle: there exist constructed learning problems with a *provable* quantum kernel advantage [12], and a useful kernel must in any case have an inductive bias matched to the data [13]. Large-scale benchmarks likewise find that out-of-the-box classical models routinely match proposed quantum models on standard tasks [16], which makes exact statements about *why* a given construction is classical the more valuable. The contribution here is narrow and practical—a specific, popular construction is classical—and the method is the transferable part: when the status of an embedding is unclear, derive its feature map. The result sits within the broader *classical surrogate* program [3, 17, 18]; for the depth-one IQP kernel the surrogate is not fitted but exact and available in closed form.

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).
- [3] F. J. Schreiber, J. Eisert, and J. J. Meyer, “Classical surrogates for quantum learning models,” *Physical Review Letters* **131**, 100803 (2023).
- [4] S. Thanasilp, S. Wang, M. Cerezo, and Z. Holmes, “Exponential concentration in quantum kernel methods,” *Nature Communications* **15**, 5200 (2024).
- [5] V. Havlíček, A. D. Córcoles, K. Temme, A. W. Harrow, A. Kandala, J. M. Chow, and J. M. Gambetta, “Supervised learning with quantum-enhanced feature spaces,” *Nature* **567**, 209 (2019).
- [6] M. Schuld and N. Killoran, “Quantum machine learning in feature Hilbert spaces,” *Physical Review Letters* **122**, 040504 (2019).
- [7] B. Schölkopf and A. J. Smola, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond* (MIT Press, 2002).
- [8] J. Shawe-Taylor and N. Cristianini, *Kernel Methods for Pattern Analysis* (Cambridge University Press, 2004).
- [9] M. Schuld and F. Petruccione, *Machine Learning with Quantum Computers*, 2nd ed. (Springer, 2021).
- [10] M. J. Bremner, A. Montanaro, and D. J. Shepherd, “Average-case complexity versus approximate simulation of commuting quantum computations,” *Physical Review Letters* **117**, 080501 (2016).
- [11] M. Van den Nest, “Simulating quantum computers with probabilistic methods,” *Quantum Information and Computation* **11**, 784–812 (2011).

- [12] Y. Liu, S. Arunachalam, and K. Temme, “A rigorous and robust quantum speed-up in supervised machine learning,” *Nature Physics* **17**, 1013 (2021).
- [13] J. M. Kübler, S. Buchholz, and B. Schölkopf, “The inductive bias of quantum kernels,” in *Advances in Neural Information Processing Systems (NeurIPS)* **34** (2021).
- [14] R. Shaydulin and S. M. Wild, “Importance of kernel bandwidth in quantum machine learning,” *Physical Review A* **106**, 042407 (2022).
- [15] A. Canatar, E. Peters, C. Pehlevan, S. M. Wild, and R. Shaydulin, “Bandwidth enables generalization in quantum kernel models,” *Transactions on Machine Learning Research* (2023).
- [16] J. Bowles, S. Ahmed, and M. Schuld, “Better than classical? The subtle art of benchmarking quantum machine learning models,” arXiv:2403.07059 (2024).
- [17] J. Landman, S. Thabet, C. Dalyac, H. Mhiri, and E. Kashefi, “Classically approximating variational quantum machine learning with random Fourier features,” in *International Conference on Learning Representations (ICLR)* (2023).
- [18] R. Sweke, E. Recio-Armengol, S. Jerbi, E. Gil-Fuster, B. Fuller, J. Eisert, and J. J. Meyer, “Potential and limitations of random Fourier features for dequantizing quantum machine learning,” *Quantum* **9**, 1640 (2025).
- [19] E. Gil-Fuster, J. Eisert, and V. Dunjko, “On the expressivity of embedding quantum kernels,” *Machine Learning: Science and Technology* **5**, 025003 (2024).