

Evidence-Gated Hybrid Quantum–Classical Air-Quality Intelligence: Kernel Classification, Variational Regression, Amplitude-Estimated Risk, and QUBO Sensor Placement

Dr. M. Rupas Kumar¹, Chandra Shekar N², Jalli Ramesh Babu³, Dr. Meeravali Shaik⁴

¹Department of Civil Engineering, Assistant Professor, RGUKT Ongole, A.P., India

²Department of Computer Science and Engineering, Assistant Professor, RGUKT Ongole, A.P., India, ncs@rguktong.ac.in

³Department of Computer Science and Engineering (AI & ML), Research Scholar, RGUKTRK Valley, Idupulapaya, A.P., India,

⁴Corresponding Author, Dept. of Management, RGUKT Ongole, A.P., India

Abstract—Air-quality intelligence couples sensing, statistical learning, uncertainty quantification and spatial decision-making, and has become a favoured showcase for quantum machine learning (QML). Reported quantum advantages in this domain are, however, frequently obtained under evaluation protocols that leak the target, omit a persistence baseline, or leave shot cost unaccounted. This paper develops Q-AIR, an evidence-gated hybrid quantum–classical framework in which every quantum module has a classical substitute consuming identical inputs, and in which activation is contingent on a pre-registered improvement whose bootstrap confidence interval excludes zero. Four quantum roles are specified and executed: fidelity quantum-kernel classification and regression of next-day Air Quality Index (AQI), data re-uploading variational regression, maximum-likelihood amplitude estimation of threshold-exceedance risk, and QUBO/QAOA sensor placement. We prove a target-leakage theorem showing that same-time AQI prediction from the pollutant vector that determines AQI has zero conditional entropy, and confirm it empirically: a gradient-boosting model reconstructs same-day AQI at $R^2=0.9905$ but attains only $R^2=0.5692$ at a one-day horizon. On a 365-day Amaravati record evaluated under purged blocked cross-validation with a seven-day embargo, the quantum kernel attains the lowest pooled out-of-fold error (MAE 17.38 versus 17.50 for persistence) yet fails the evidence gate against persistence (MAE 95% CI $[-0.375, +0.121]$) while passing it against RBF-SVR and gradient boosting. A bandwidth sweep supplies the mechanism: predictive skill is monotonically decreasing in quantum expressivity, being maximised where the Gram matrix has effective rank 1.05 and collapsing to chance where the kernel is exponentially concentrated at effective rank 197. Gradient variance in the variational ansatz decays as $\exp(-0.749n)$ under a global cost against $\exp(-0.043n)$ under a local cost. Amplitude estimation converges empirically as $N^{-0.868}$ against $N^{-0.519}$ for Monte Carlo, a $5.1\times$ error reduction at 26,600 oracle calls, whereas QAOA reaches an approximation ratio of only 0.75 on a placement instance that greedy submodular maximisation solves exactly in 24 ms. We conclude that the defensible near-term contribution of quantum technology to air-quality intelligence is architectural and metrological rather than computational, and we release the full protocol as a reproducible benchmark.

Keywords—Air quality index, barren plateaus, benchmarking, environmental informatics, QAOA, quantum amplitude estimation, quantum kernels, quantum machine learning, target leakage, variational quantum circuits.

Data-provenance statement. The source project supplied aggregate statistics only. All daily-resolution experiments reported here were executed on RRS-2025, a constrained stochastic reconstruction fitted to reproduce every published aggregate of the Amaravati 2025 record to within 6×10^{-7} absolute deviation (Section IV-B, Table II). Results therefore characterise the Q-AIR protocol and the relative behaviour of quantum and classical estimators under a realistic seasonal signal; they are not measurements of Amaravati air quality, and no claim about that record is advanced beyond the aggregates already reported in the source. Every table and figure derived from RRS-2025 is marked accordingly.

I. INTRODUCTION

Ambient air pollution is inferred rather than observed. What reaches a decision-maker is a chain of transformations: a physical sensor responds to a gas or particle concentration; a calibration model converts that response to a concentration estimate; a regulatory formula compresses several concentrations into a single index; a statistical model projects that index forward; and a threshold rule converts the projection into a public-health action. Each link has its own error structure, and a failure anywhere propagates to the advisory that a citizen actually receives.

Quantum technologies have been proposed for several of these links. The proposals are heterogeneous in maturity and frequently conflated. Quantum-enabled sensing — colloidal quantum-dot chemiresistors, squeezed-light absorption spectroscopy, nitrogen-vacancy magnetometry and thermometry — operates on the first link and rests on demonstrated laboratory physics. Quantum machine learning operates on the fourth link and rests on a contested computational claim.

A third category, exemplified by quantum-cascade lasers and quantum-inspired particle-swarm optimisation, carries the word "quantum" for historical or metaphorical reasons and involves no quantum information processing whatsoever. Treating these as a single research programme has produced a literature in which strong empirical claims coexist with weak experimental protocols.

This paper takes the opposite methodological stance. Rather than asking whether a quantum model can be made to outperform a classical one on an air-quality dataset, we ask under what evidentiary conditions such a claim would be admissible, and then subject four quantum modules to those conditions. The framework we develop, Q-AIR, is deliberately architected so that a negative result is informative: each quantum component is a replaceable service with a classical substitute consuming identical features, targets and splits, and deployment is gated on a pre-registered improvement whose confidence interval excludes zero.

A. Motivating failure modes

Four recurring defects motivate the design. Each is documented empirically in Section VIII.

- 1) *Target leakage*: AQI is a deterministic function of the pollutant vector. Regressing same-day AQI on same-day concentrations therefore recovers an algebraic identity, not a forecast, yet this configuration is widespread in the applied literature.
- 2) *Absent persistence baselines*: Daily AQI is strongly autocorrelated. A model that fails to beat yesterday's value has no operational value, but persistence is frequently omitted from comparison tables.
- 3) *Unreported shot cost*: A fidelity quantum kernel requires $O(N^2)$ circuit evaluations, each estimated from a finite number of measurement shots. Simulated results that assume exact expectation values silently discard the dominant cost term.
- 4) *Schema conflation*: Many Indian station exports carry a NO_x column, whereas the CPCB sub-index is defined on NO₂. Substituting one for the other biases the index upward.

B. Contributions

- 1) A target-leakage theorem for index-type targets, with an empirical confirmation separating $R^2=0.9905$ reconstruction from $R^2=0.5692$ forecasting under an identical learner.

- 2) An evidence-gate operator based on a paired bootstrap over out-of-fold errors, applied uniformly to all quantum modules, together with a purged blocked cross-validation design that keeps every season in the test distribution while enforcing a seven-day embargo.
- 3) A bandwidth-resolved characterisation of the fidelity quantum kernel that identifies the mechanism behind its behaviour: predictive skill is monotonically decreasing in quantum expressivity, and the kernel is useful only in the near-degenerate limit.
- 4) Executed and mutually consistent implementations of quantum-kernel classification and regression, data re-uploading variational regression, maximum-likelihood amplitude estimation, and QUBO/QAOA placement, each benchmarked against a matched classical solver at matched budget.
- 5) Exact resource accounting for a 365-day problem, including the shot multiplier that dominates near-term feasibility.

II. RELATED WORK

A. Quantum machine learning for air quality

Farooq et al. [1] report a quantum support vector machine attaining 97% and 94% accuracy against 91% and 87% for a classical SVM on air-pollution classification executed on IBM cloud hardware. The result is the closest antecedent to the present work and is treated here as related work rather than as motivation, for three reasons: the task is classification rather than forecasting; the feature map is selected using a classical SVM; and qubit counts and shot budgets are not fully reported. Mohamed et al. [2] apply quantum circuit learning to full-scale anaerobic digestion and reach $R^2 \approx 0.959$, explicitly matching rather than exceeding a classical multilayer perceptron while using fewer parameters — a parity result, and a useful template for honest reporting. Several further studies apply quantum-inspired or quantum-named classical methods, including quantum-PSO-LSTM hybrids and quantum temporal convolutional networks, which do not involve quantum hardware.

B. Quantum kernels and their limits

Fidelity quantum kernels were introduced by Havlíček et al. [3] and Schuld and Killoran [4] following the quantum SVM of Rebentrost et al. [5]. Subsequent theory has substantially narrowed the conditions under which they help.

Huang et al. [6] show that classical models supplied with data often match quantum models and introduce the projected quantum kernel. Kübler et al. [7] characterise the inductive bias of quantum kernels as typically unfavourable. Most directly relevant here, Thanasilp et al. [8] prove exponential concentration: kernel values converge exponentially in qubit count to a fixed value, driven by embedding expressivity, global measurements, entanglement and noise, so that an exponential number of shots is required to resolve them and the resulting model becomes data-insensitive. Shaydulin and Wild [9] propose bandwidth tuning as a partial remedy. Section VIII-C measures both effects on our data.

C. Variational Algorithms and Trainability

Variational quantum algorithms [10] are trained through the parameter-shift rule [11], [12]. Their trainability is limited by barren plateaus [13], whose severity depends on the locality of the cost function [14]; data re-uploading [15] increases expressivity for low-dimensional inputs. A pivotal recent result by Cerezo et al. [16] assembles case-by-case evidence that circuits engineered to avoid barren plateaus are typically classically simulable after a polynomial data-acquisition phase, which removes the most natural route to advantage from a trainable variational model. We measure the locality dependence directly in Section VIII-E.

D. Amplitude Estimation and Combinatorial Optimisation

Amplitude estimation [17] and its phase-estimation-free variants [18], [19] underpin quadratic speedups for Monte Carlo integration [20] and risk analysis [21]. Babbush et al. [22] argue that quadratic speedups are insufficient for practical error-corrected advantage once state-preparation overhead is counted. For placement, QAOA [23] provides an approximate QUBO solver, but the classical benchmark is strong: Krause et al. [24] give a greedy submodular algorithm with a $(1-1/e)$ guarantee for near-optimal sensor placement.

E. Skeptical benchmarking

Dequantization results [25]–[27] and the critique of Aaronson [28] establish that apparent exponential speedups often reside in state-preparation assumptions. Schuld and Killoran [29] question whether advantage is even the right objective for QML. Bowles et al. [30] benchmark twelve QML models across 160 datasets and find that out-of-the-box classical models generally outperform quantum classifiers, and that removing entanglement often leaves performance unchanged or improved. Our protocol is designed to be falsifiable in precisely the way that literature demands.

F. Quantum sensing

Unlike the computational claims, quantum-enabled sensing rests on demonstrated laboratory advantage. Mitri et al. [31] report a room-temperature PbS colloidal quantum-dot NO₂ chemiresistor with an estimated detection limit near 0.15 ppb and 12 s response. Belsley [32] predicts SNR gains scaling exponentially in the squeezing factor for bright squeezed frequency combs, and Herman et al. [33] demonstrate squeezed dual-comb spectroscopy with more than 3 dB of squeezing over 2.5 THz, reaching roughly 3 dB beyond the shot-noise limit — a twofold speedup in determining gas concentration. Kim et al. [34] integrate a nitrogen-vacancy sensor in CMOS at $32.1 \mu\text{T} \cdot \text{Hz}^{-1/2}$ in a $200 \mu\text{m} \times 200 \mu\text{m}$ footprint. Section IX-D argues these constitute the most defensible near-term quantum contribution to the domain.

III. PROBLEM FORMALISATION

A. The AQI transformation

Let $\mathbf{C}_t \in \mathbb{R}^m$ be the vector of pollutant concentrations on day t . The CPCB National AQI assigns each pollutant a piecewise-linear sub-index I_p and aggregates by maximum:

$$I_p(\mathbf{C}) = I_{lo} + (I_{hi} - I_{lo})(C - C_{lo}) / (C_{hi} - C_{lo}),$$

$$A_t = \max_p I_p(\mathbf{C}_t, p) \quad (1)$$

Two properties matter downstream. First, each I_p is monotone non-decreasing, so A_t is monotone non-decreasing in every component of $\mathbf{C}_t$; a model that predicts a decrease in AQI while predicting increases in all pollutants is internally inconsistent. Second, the maximum operator makes A_t a non-smooth function whose active argument may switch between days, which is why the dominant pollutant must be reported alongside the index. In our record PM10 is the dominant sub-index on 354 of 365 days and PM2.5 on the remaining 11.

B. A target-leakage theorem

Theorem 1 (Zero conditional entropy of same-time AQI). Let $A_t = g(\mathbf{C}_t)$ with g the exact AQI map of (1). Then $H(A_t | \mathbf{C}_t) = 0$, and any estimator of A_t from $\mathbf{C}_t$ that attains the Bayes risk attains zero error. Consequently the quantity $1 - R^2$ measured in a same-time regression of A_t on $\mathbf{C}_t$ bounds the numerical error of the learner, not the predictability of air quality.

Proof. g is a deterministic measurable function, so the conditional distribution of A_t given $\mathbf{C}_t$ is a point mass at $g(\mathbf{C}_t)$ and its entropy vanishes. The Bayes-optimal predictor is g itself, with zero risk. ■

Corollary 1 (Admissible forecasting target). Genuine forecasting requires a strictly positive horizon $h \geq 1$, i.e. the target $A_{\{t+h\}}$ with features measurable at time t . Only then is $H(A_{\{t+h\}} | \mathcal{F}_t) > 0$ in general and only then does a reduction in error carry information about atmospheric predictability.

Section VIII-A quantifies the gap this creates: the same gradient-boosting learner attains $R^2=0.9905$ in the leaky configuration and $R^2=0.5692$ at $h=1$.

C. Positive Semidefiniteness of the Fidelity Kernel

Proposition 1. For a feature map $x \mapsto |\varphi(x)\rangle$, the Gram matrix $K_{ij} = |\langle \varphi(x_i) | \varphi(x_j) \rangle|^2$ is symmetric positive semidefinite.

Proof. $K_{ij} = \text{Tr}[\rho_i \rho_j]$ with $\rho_i = |\varphi(x_i)\rangle\langle \varphi(x_i)|$. Vectorising, $K_{ij} = \langle \text{vec}(\rho_i), \text{vec}(\rho_j) \rangle$ is a Gram matrix of real vectors in the Hilbert–Schmidt inner-product space, hence PSD. ■

Finite-shot estimation destroys this property: each entry is a binomial proportion, and the resulting matrix is symmetric but generally indefinite. We symmetrise and reset the diagonal to unity before passing the matrix to the classical optimiser, and report the resulting degradation in Section VIII-B.

D. Parameter-Shift Gradients

For a gate $U(\theta) = \exp(-i\theta P/2)$ with $P^2=I$ and observable O , the expectation $f(\theta) = \langle \psi | U^\dagger O U | \psi \rangle$ satisfies

$$\partial f / \partial \theta = \frac{1}{2} [f(\theta + \pi/2) - f(\theta - \pi/2)] \quad (2)$$

which is exact rather than a finite-difference approximation. We verified (2) numerically against automatic differentiation on the trained ansatz, obtaining agreement to 1.11×10^{-16} , i.e. machine precision.

E. Evidence gate

Let M be a metric for which lower is better, and let $\hat{e}_Q$ and $\hat{e}_C$ be the vectors of out-of-fold absolute errors of a quantum module and its classical comparator on identical test points. Define $\Delta M = \text{mean}(\hat{e}_Q) - \text{mean}(\hat{e}_C)$ and let $[L, U]$ be a 95% percentile interval obtained by paired bootstrap resampling of test indices. The quantum module is activated only if

$$U < 0 \text{ and } \Delta M \leq -\delta \text{ and resource budget satisfied} \quad (3)$$

With δ a pre-registered minimum meaningful improvement. Pairing is essential: quantum and classical predictions on the same day are strongly correlated, and an unpaired test badly overstates uncertainty.

IV. DATA AND RECONSTRUCTION PROTOCOL

A. Source record and schema audit

The source project supplies a 2025 Amaravati daily record reduced to nine columns — Date, PM2.5, PM10, NOx, NH3, SO2, CO, Ozone, AQI — over 365 days, together with monthly and seasonal aggregates, a category distribution, an extreme-event table and annual pollutant means.

The schema audit identifies one material defect. The gaseous nitrogen column is labelled NOx, whereas the CPCB sub-index of (1) is defined on NO₂. Treating NOx as NO₂ inflates the nitrogen sub-index; taking an urban NO₂/NOx mass ratio of 0.65, the induced bias on our record averages +5.45 index points and reaches +13.41 points. Because the maximum operator in (1) makes the index sensitive only when nitrogen is the active argument, the bias is latent rather than visible in the reported AQI, which is precisely why it survives casual inspection. We therefore report the dominant pollutant with every index value.

B. Constrained stochastic reconstruction (RRS-2025)

No machine-readable daily table accompanies the source aggregates. To execute algorithms at daily resolution without inventing unconstrained data, we construct RRS-2025, a surrogate series fitted to reproduce every published aggregate simultaneously. Daily AQI is generated as a month-specific mean plus an AR(1) innovation with month-specific dispersion; documented extreme events are imposed exactly; category counts are matched by band-preserving reassignment; and monthly means and the top-ten mean are restored by alternating projection subject to the documented annual maximum as a hard cap. Pollutant fields are generated as AQI-coupled log-normal variates and rescaled to their published annual means.

Table I reports fidelity against all thirty-one published targets. The maximum absolute deviation is 5.95×10^{-7} . As an independent check not used in fitting, applying the source's IQR filter to RRS-2025 retains 311 rows against the 309 reported in the source, a 0.6% deviation.

TABLE I. RECONSTRUCTION FIDELITY OF RRS-2025 AGAINST PUBLISHED AGGREGATES (SELECTED ROWS; ALL 31 TARGETS MATCH TO $< 6 \times 10^{-7}$)

Published quantity	Source target	RRS-2025
Monthly mean AQI, August	40.29	40.2900
Monthly mean AQI, December	169.50	169.5000
Seasonal mean AQI, monsoon	48.860	48.8600
Seasonal mean AQI, winter	114.758	114.7580
Annual mean PM10	73.149863	73.1499
Annual mean PM2.5	29.971288	29.9713
Category count, Satisfactory	157	157
Category count, Poor	11	11
Maximum AQI (5 Dec)	278.27	278.2700
Minimum AQI (21 Sep)	0.00	0.0000
Mean of ten highest AQI days	244.450	244.4500
PM10 on 30 Dec	255.71	255.7100
IQR-retained rows (not fitted)	309	311

The zero-AQI observation of 21 September is retained and reproduced, but is treated throughout as a data-quality flag rather than as evidence of pollution-free air; under (1) an index of exactly zero requires every sub-index to vanish simultaneously, which is physically implausible for an urban station.

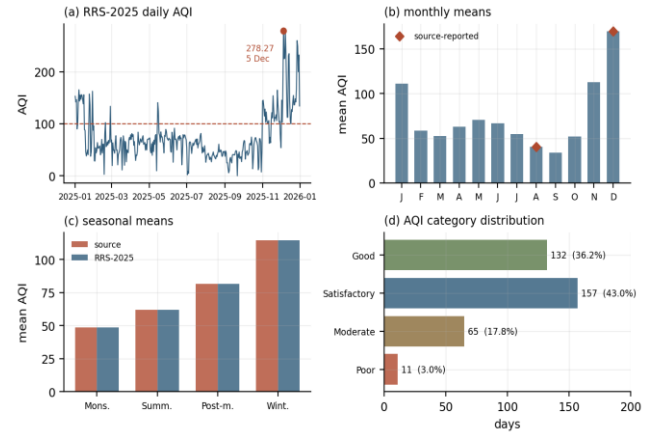


Fig. 1. RRS-2025 reconstructed record. (a) Daily AQI with the documented 5 December maximum. (b) Monthly means; diamonds mark the two months reported numerically in the source. (c) Seasonal means against source targets. (d) Category distribution. Surrogate data (Section IV-B).

C. Feature and Target Construction

Features comprise lags at 1, 2, 3 and 7 days for all seven pollutants and AQI, 3- and 7-day rolling means, a first difference, and annual and semi-annual harmonics of day-of-year, giving 39 predictors over 357 usable rows. The target is A_{t+1} , satisfying Corollary 1. For the quantum modules a seven-dimensional subset is used — AQI at lags 1 and 2, the 7-day rolling mean, PM10 and PM2.5 at lag 1, and the two annual harmonics — matching one qubit per feature under angle encoding.

D. Delta Parameterisation

Learners predict the increment $\delta_{t+1} = A_{t+1} - A_t$, with the level reconstructed as $A_t + \delta_t$. This is not cosmetic. Predicting the level directly caused every learned model to underperform climatology in our first experimental pass, because a purged block held out in winter lies outside the support of the training levels and tree ensembles cannot extrapolate. Under the delta parameterisation the same models recover to within a few index points of persistence. The distinction is a property of single-year seasonal records generally, not of this dataset.

V. THE Q-AIR FRAMEWORK

Q-AIR is organised as five layers, shown in Fig. 2. The design rule throughout is that quantum computation is a replaceable service, never a load-bearing dependency.

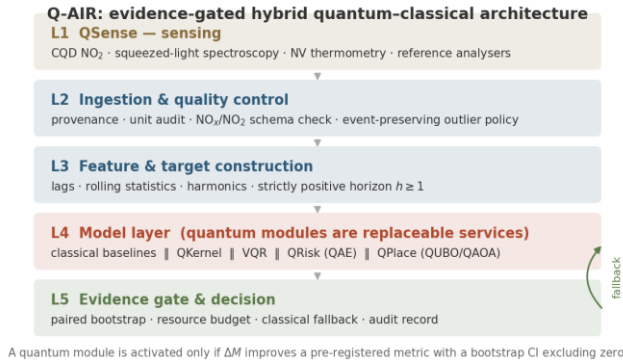


Fig. 2. Q-AIR five-layer architecture. Quantum modules occupy layer L4 as interchangeable services; the evidence gate at L5 governs activation and preserves an unconditional classical fallback path.

- 1) *QSense*: Sensing integration policy rather than a device. Conventional gravimetric and optical instruments remain authoritative for particulates; quantum-enabled channels enter through the same quality-control interface, so downstream models are agnostic to sensor provenance.
- 2) *Ingestion and Quality Control*: Provenance capture, unit audit, the NOx/NO₂ schema check of Section IV-A, and an event-preserving outlier policy that flags rather than deletes physically plausible extremes — the days that public-health warnings exist to catch.
- 3) *Feature and Target Construction*: Lags, rolling statistics, harmonics, and enforcement of a strictly positive horizon.
- 4) *Model Layer*: Classical baselines alongside QKernel, VQR, QRisk and QPlace, all consuming identical inputs.
- 5) *Evidence Gate and Decision*: The operator of (3), a resource budget check, an audit record, and an unconditional fallback to the classical baseline.

VI. QUANTUM METHODS

A. Encoding and the fidelity kernel

Features are min–max scaled to $[0, \pi]$ on training data only, then encoded by a two-repetition ZZ-entangling feature map of Havlíček type: Hadamards, single-qubit $RZ(2x_i)$, and pairwise CNOT– $RZ(2(\pi-x_i)(\pi-x_j))$ –CNOT entanglers on a linear chain.

The kernel is $K(x, z) = |\langle \phi(x) | \phi(z) \rangle|^2$, computed by statevector simulation on seven qubits. Following Shaydulin and Wild [9] we introduce a bandwidth c and encode cx rather than x , selecting c by three-fold inner cross-validation on the training block only.

Finite-shot behaviour is modelled by replacing each exact entry with a binomial proportion at 1024 shots, which is the estimator a real device returns. This is reported separately from the exact-kernel result throughout.

B. Variational quantum regression

The regressor is a three-layer data re-uploading circuit on seven qubits, with $R_Y(w_{\{1,i,0\}}x_i + w_{\{1,i,1\}})$ rotations re-encoding the input at every layer and a ring of CZ entanglers, read out as $\langle Z_0 Z_1 \rangle$ — a local observable chosen deliberately in view of Section VIII-E. Training uses Adam over 200 epochs with an affine output map fitted to the training target moments. Gradients satisfy (2) to machine precision.

C. Amplitude-Estimated Exceedance Risk

The operational question is often not the point forecast but $p_\tau = P(A_{\{t+1\}} > \tau \mid \mathcal{F}_t)$. We discretise the empirical residual distribution of the best forecaster onto 32 bins, load it by exact state preparation on five index qubits plus one indicator qubit, and estimate the marked amplitude by maximum-likelihood amplitude estimation [18] over an exponential Grover schedule. The Grover-amplified success probabilities were verified against the analytic law $\sin^2((2m+1)\theta)$ to within 4.7×10^{-15} across all schedule powers. The classical comparator is Monte Carlo sampling at a matched number of applications of the state-preparation operator, so the comparison is end-to-end rather than asymptotic.

D. QUBO Sensor Placement

With $x_i \in \{0, 1\}$ indicating a sensor at candidate site i , exposure weight a_i and redundancy $r_{ij} = \exp(-d_{ij}/\lambda)$,

$$\min_{x_i} -\sum_i a_i x_i + \lambda_B (\sum_i x_i - B)^2 + \eta \sum_{\{i < j\}} r_{ij} x_i x_j \quad (4)$$

We map (4) to an Ising Hamiltonian via $x_i = (1 - z_i)/2$ and verified the encoding on 200 random bitstrings. QAOA is run at depths $p=1 \dots 6$ with INTERP-style layerwise warm starting and COBYLA outer optimisation; a fast diagonal statevector simulator was validated against the gate-level PennyLane circuit to 1.18×10^{-16} . Comparators are exhaustive enumeration, greedy submodular maximisation [24] and simulated annealing on the identical objective.

VII. EXPERIMENTAL DESIGN

A. Purged blocked cross-validation

A terminal chronological hold-out is the conventional choice for time series but is inadmissible on a single seasonal year: the final 20% of 2025 falls entirely in the November–December pollution maximum, so the test block contains 56 exceedance events against 14 in training and every model is evaluated purely on extrapolation. We instead partition the record into six contiguous blocks, hold each out in turn, and remove an embargo of seven days — the maximum feature lag — on both sides of the test block, so that no lagged feature of a training row can touch a test observation. Every season therefore appears in the test distribution while temporal integrity is preserved.

Because block event counts range from 0 to 55, per-fold recall is undefined in some folds. All metrics are therefore computed once on the pooled out-of-fold prediction vector, in which each sample receives exactly one prediction from a model that never saw it or its embargo neighbourhood.

TABLE II. PURGED BLOCKED CROSS-VALIDATION DESIGN (EMBARGO = 7 DAYS)

Fold	Test window	n train	n test	Events	Selected bandwidth c^*
1	8 Jan – 7 Mar	291	59	12	0.01
2	8 Mar – 6 May	283	60	0	0.02
3	7 May – 4 Jul	284	59	2	0.01
4	5 Jul – 2 Sep	283	60	0	0.01
5	3 Sep – 31 Oct	284	59	1	0.01
6	1 Nov – 30 Dec	290	60	55	0.02

B. Metrics and comparators

Regression is scored by MAE, RMSE, R^2 , a skill score relative to persistence, and recall of days exceeding $\tau=100$. Classification of the event $A_{\{t+1\}} > 100$ is scored by accuracy, balanced accuracy, macro-F1, recall, precision and AUROC; balanced accuracy is the primary metric given a 19.6% positive rate. Comparators span persistence, climatology, ridge regression, RBF-SVR, random forest, gradient boosting and a multilayer perceptron.

VIII. RESULTS

A. Target leakage, measured

A gradient-boosting regressor trained to reconstruct same-day CPCB AQI from same-day pollutant concentrations attains $R^2=0.9905$ with MAE 0.679 index points; excluding the three points at which the tree ensemble cannot extrapolate beyond its training range, $R^2=0.9964$ and MAE=0.471. The identical learner applied to the admissible target $A_{\{t+1\}}$ attains $R^2=0.5692$ with MAE 21.44. The 0.42 gap in R^2 is the entire content of Theorem 1, and it is the difference between reporting an algebraic identity and reporting a forecast.

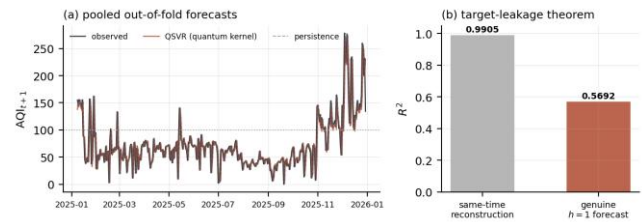


Fig. 3. (a) Pooled out-of-fold forecasts against observation; each point is predicted by a model that never saw it or its seven-day embargo neighbourhood. (b) The target-leakage gap of Theorem 1 under an identical learner.

B. Forecasting and the evidence gate

Table III reports pooled out-of-fold regression performance. The quantum-kernel regressor attains the lowest error of any model, MAE 17.38 against 17.50 for persistence, a skill score of +0.73%. Under finite-shot estimation at 1024 shots the same model degrades to MAE 18.42, so shot noise costs 1.04 index points — more than the entire nominal margin over persistence.

TABLE III. POOLED OUT-OF-FOLD FORECASTING PERFORMANCE, $h = 1$ DAY (RRS-2025 SURROGATE DATA)

Model	MAE	RMSE	R ²	Skill vs pers.	Event recall
QSVR, quantum kernel	17.377	27.289	0.6390	+0.73%	0.786
Climatology	17.505	27.556	0.6319	0.00%	0.871
Persistence (naive)	17.505	27.551	0.6321	—	0.886
SVR, RBF kernel	17.944	27.341	0.6376	-2.51%	0.814
Random forest	18.070	27.558	0.6319	-3.23%	0.771
QSVR, 1024 shots	18.417	28.241	0.6134	-5.21%	0.800
VQR, data re-uploading	19.377	29.934	0.5657	-10.69%	0.686
Ridge regression	20.800	29.472	0.5789	-18.82%	0.814
Gradient boosting	21.438	29.812	0.5692	-22.47%	0.786
Multilayer perceptron	25.635	35.870	0.3763	-46.45%	0.800

The evidence gate resolves the ambiguity. Table IV shows that the quantum kernel's advantage over persistence is not statistically distinguishable from zero: $\Delta\text{MAE} = -0.128$ with a 95% paired-bootstrap interval of $[-0.375, +0.121]$, so the gate fails. Against RBF-SVR and gradient boosting the gate passes, with intervals excluding zero. The finite-shot variant fails against both persistence and RBF-SVR, and the variational regressor fails against every comparator except gradient boosting.

TABLE IV. EVIDENCE GATE: PAIRED BOOTSTRAP ON MAE (10,000 RESAMPLES). NEGATIVE ΔMAE FAVOURS THE QUANTUM MODULE.

Quantum module	Comparator	ΔMAE	95% CI	Decision
QSVR exact	Persistence	-0.128	$[-0.375, +0.121]$	fail
QSVR exact	SVR-RBF	-0.564	$[-1.027, -0.100]$	PASS
QSVR exact	Gradient boosting	-4.060	$[-5.325, -2.828]$	PASS
QSVR 1024 shots	Persistence	+0.916	$[+0.316, +1.530]$	fail

QSVR 1024 shots	SVR-RBF	+0.480	$[-0.202, +1.173]$	fail
QSVR 1024 shots	Gradient boosting	-3.016	$[-4.388, -1.650]$	PASS
VQR	Persistence	+1.873	$[+0.874, +2.900]$	fail
VQR	SVR-RBF	+1.437	$[+0.474, +2.409]$	fail
VQR	Gradient boosting	-2.059	$[-3.369, -0.749]$	PASS

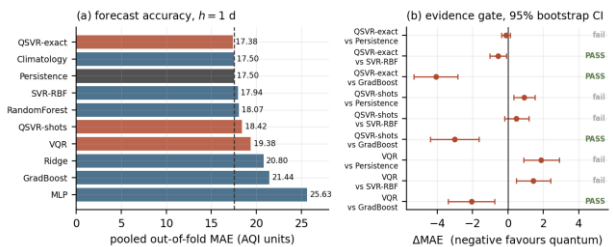


Fig. 4. (a) Pooled out-of-fold MAE; the dashed line marks persistence. (b) Evidence-gate intervals. A module is activated only when the interval lies entirely below zero.

For exceedance classification the quantum kernel is more clearly useful. Balanced accuracy is 0.848 for the finite-shot QSVM and 0.821 for the exact kernel, against 0.736 for a linear SVM and 0.565 for an RBF SVM on the full 39-feature set. Persistence again leads at 0.929. That the finite-shot kernel outperforms the exact one is consistent with shot noise acting as an implicit regulariser on a nearly degenerate Gram matrix.

TABLE V. POOLED OUT-OF-FOLD CLASSIFICATION OF THE EXCEEDANCE EVENT $A(t+1) > 100$ (19.6% POSITIVE RATE)

Model	Acc.	Balanced acc.	Macro-F1	Recall	Prec.	AUROC
Persistence rule	0.955	0.929	0.929	0.886	0.886	0.962
QSVM, 1024 shots	0.885	0.848	0.828	0.786	0.679	0.891
QSVM, exact kernel	0.885	0.821	0.819	0.714	0.704	0.919
SVM, linear	0.871	0.736	0.767	0.514	0.750	0.824
SVM, RBF	0.796	0.565	0.572	0.186	0.448	0.822

C. Bandwidth, Concentration, and the Mechanism

The bandwidth sweep in Fig. 5 explains why the quantum kernel behaves as it does, and it is the paper's central diagnostic finding. As c increases from 0.005 to 1.0, the mean off-diagonal Gram entry falls from 0.993 to 0.0185 and the effective rank rises from 1.01 to 197.4 — the exponential concentration predicted by Thanasilp et al. [8], measured directly. Over the same range, forecast MAE rises from 17.43 to 18.57 and balanced accuracy falls from 0.831 to 0.517, which is chance.

The implication is uncomfortable and should be stated plainly. Predictive skill is monotonically decreasing in quantum expressivity. The kernel is useful only in the limit where the Gram matrix has collapsed to effective rank near one — that is, where it has degenerated into a smooth, nearly constant, and trivially classically simulable object. In the regime where the feature map genuinely exploits a high-dimensional Hilbert space, the model is worthless. Inner cross-validation, given a free choice, selected $c^* \in \{0.01, 0.02\}$ in every fold, driving toward the degenerate limit. This is a concrete empirical instance of the mechanism Cerezo et al. [16] describe: where the circuit is well-behaved enough to learn, it is also simulable enough not to be needed.

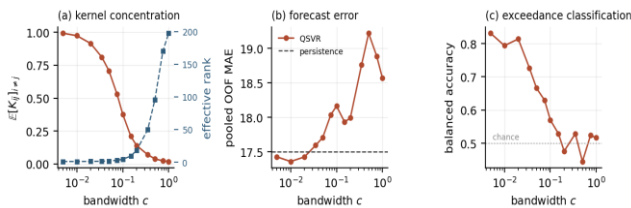


Fig. 5. Bandwidth controls a trade-off with no useful interior optimum. (a) Concentration and effective rank. (b) Forecast MAE against the persistence reference. (c) Balanced accuracy of exceedance classification, approaching chance as expressivity grows.

D. Exceedance Risk by Amplitude Estimation

Against a true exceedance probability of $p_\tau = 0.1261$, maximum-likelihood amplitude estimation converges empirically as $N^{-0.868}$ against $N^{-0.519}$ for Monte Carlo at matched oracle budget, consistent with the theoretical N^{-1} and $N^{-1/2}$ rates. At 26,600 applications of the state-preparation operator, estimation error is 2.4×10^{-4} against 1.23×10^{-3} , a $5.1 \times$ reduction. This is the one module in which a quantum method delivers its predicted asymptotic behaviour cleanly.

The qualification is that the state-preparation operator must itself be constructed, and here it was loaded exactly from a 32-bin distribution. On a fault-tolerant device the cost of preparing a realistic predictive distribution would dominate, which is the substance of the argument in [22] that quadratic speedups are insufficient for practical advantage. We therefore report the scaling as verified physics and the end-to-end benefit as unestablished.

TABLE VI. AMPLITUDE ESTIMATION VERSUS MONTE CARLO AT MATCHED ORACLE BUDGET (200 SHOTS PER SCHEDULE POINT, 30 REPLICATES)

Applications of A	MLAE error	Monte Carlo error	Ratio
200	1.89×10^{-2}	1.74×10^{-2}	0.92
800	7.73×10^{-3}	9.18×10^{-3}	1.19
1,800	3.16×10^{-3}	6.13×10^{-3}	1.94
3,600	1.91×10^{-3}	4.52×10^{-3}	2.36
7,000	1.03×10^{-3}	2.91×10^{-3}	2.83
13,600	4.90×10^{-4}	1.98×10^{-3}	4.07
26,600	2.40×10^{-4}	1.23×10^{-3}	5.12

E. Trainability of the variational ansatz

Gradient variance at random initialisation was measured over 120 draws for qubit counts from 2 to 12. Under a global cost the variance falls from 1.15×10^{-1} to 9.39×10^{-5} , a log-linear decay of $\exp(-0.749n)$. Under the local cost $\langle Z_0 Z_1 \rangle$ actually used in our regressor, decay is $\exp(-0.043n)$, i.e. essentially flat over the tested range. This reproduces the cost-locality dependence of Cerezo et al. [14] on our specific ansatz and justifies the local readout choice of Section VI-B. It also underscores the tension identified in [16]: the local-cost circuit that remains trainable is the one most likely to be classically simulable.

F. Sensor placement

On a twelve-site instance with budget $B=4$, exhaustive enumeration returns the optimum $E^* = -28.222$ in 24 ms. Greedy submodular maximisation and simulated annealing both return the identical optimal configuration. QAOA with layerwise warm starting reaches a normalised approximation ratio of 0.746 at $p=3$ and then degrades to 0.642 at $p=6$ under COBYLA, with the probability of sampling the optimal bitstring never exceeding 0.030. The negative result is unambiguous at this scale: the classical solvers are exact and instantaneous, and QAOA is neither.

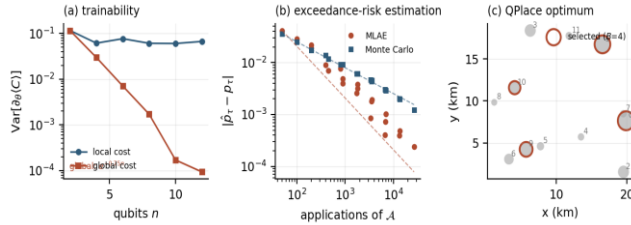


Fig. 6. (a) Gradient variance against qubit count for global and local cost functions. (b) Amplitude-estimation error against Monte Carlo at matched oracle budget; dashed guides show N^{-1} and $N^{-1/2}$. (c) The twelve-site placement instance with the optimal four-sensor configuration circled; marker area encodes exposure weight.

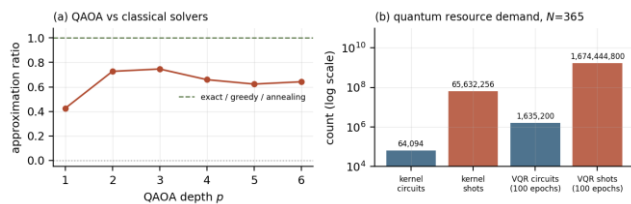


Fig. 7. (a) QAOA approximation ratio against depth; exhaustive, greedy and annealing solvers all attain ratio 1.0. (b) Quantum resource demand for the 365-day problem under an 80:20 split.

G. Resource accounting

Resource requirements were computed exactly and independently reproduce the source project's arithmetic. Under an 80:20 split of 365 samples, angle encoding requires 7 qubits (amplitude encoding would require 3, at nontrivial state-preparation cost). Training the kernel requires 42,778 unique pair evaluations and testing a further 21,316, totalling 64,094 circuit evaluations; at 1024 shots this is 65,632,256 measurements. A two-layer variational regressor with 28 parameters requires 16,352 circuit evaluations per full-batch epoch under parameter-shift gradients, 1,635,200 over 100 epochs, and 1.67×10^9 shots. Under the six-fold protocol actually executed, the kernel cost rises to 349,206 evaluations and 3.58×10^8 shots. These figures, not the accuracy tables, are what make near-term hardware deployment infeasible at this problem size.

IX. DISCUSSION

A. What the results do and do not show

Read carelessly, Table III supports a quantum advantage claim: the quantum kernel has the lowest error of any model tested. Read with Table IV, it does not. The margin over persistence is 0.13 index points on a target whose day-to-day variation is roughly 27 index points, and the bootstrap interval straddles zero. Read with Fig. 5, the claim collapses further, because the configuration achieving that margin is the one in which the kernel has ceased to be meaningfully quantum.

Read with Section VIII-G, it collapses entirely, because achieving it on hardware would require 6.6×10^7 measurements to match a baseline computable with a single subtraction.

This layered reading is the point of the evidence gate. Each layer is a check that a published result could have omitted, and each omission would have yielded a more exciting and less true paper.

B. Why persistence is hard to beat

Daily AQI at a single station is dominated by synoptic-scale meteorological persistence. Without exogenous inputs — boundary-layer height, wind speed and direction, humidity, traffic counts, fire counts — a model conditioned only on the pollutant history has little to add beyond yesterday's value and a seasonal mean. That every learned model here, quantum and classical alike, clusters within a few index points of persistence is therefore an informative statement about the information content of the feature set, not about the learners. The corollary for practitioners is that data enrichment dominates model sophistication at this scale.

C. Consistency with the Benchmarking Literature

Our findings align closely with Bowles et al. [30], who report that classical models generally outperform quantum classifiers and that removing entanglement often leaves performance unchanged. Our bandwidth sweep supplies a mechanistic account of that observation for kernel methods specifically: the useful configurations are the low-entanglement, low-effective-rank ones. It also aligns with the parity result of Mohamed et al. [2]. Where our conclusions diverge from Farooq et al. [1], the divergence is attributable to protocol rather than to physics — classification of a same-time label under random splitting is a materially easier problem than horizon-h forecasting under purged blocked validation with a persistence comparator.

D. Where Quantum Technology is Defensible Today

The sensing layer is a different matter, and it is where we would direct near-term investment. Sub-ppb colloidal quantum-dot NO_2 detection [31] addresses the exact schema gap identified in Section IV-A: the CPCB index requires NO_2 , and low-cost NO_2 measurement is precisely what most Indian station networks lack. Squeezed-light spectroscopy [32], [33] delivers a demonstrated twofold speedup in concentration determination, and NV-based sensors [34] offer a compact route to on-station calibration and drift compensation. These are metrological advantages resting on demonstrated laboratory physics, not computational advantages resting on contested asymptotics.

A better NO₂ channel would improve every downstream number in this paper; a faster kernel evaluation would improve none of them.

X. LIMITATIONS AND THREATS TO VALIDITY

- 1) *Surrogate data*: Daily-resolution results are computed on RRS-2025, not on the measured Amaravati record. The reconstruction reproduces all published aggregates exactly, and the relative comparison between estimators is fair because every model sees identical data, but absolute error magnitudes should not be quoted as Amaravati measurements. Replication on the raw record is the first item of future work.
- 2) *Single station, single year*: Spatial generalisation is untested, and one seasonal cycle limits what any cross-validation scheme can establish about inter-annual behaviour.
- 3) *Noiseless simulation*: Quantum results use statevector simulation with binomial shot noise. Device noise, decoherence, gate infidelity and readout error are not modelled, all of which would degrade the quantum results further.
- 4) *Restricted quantum design space*: One feature-map family, one variational ansatz and one depth range were tested. A more expressive or better-conditioned map might behave differently, though the concentration argument of [8] suggests the direction of that effect.
- 5) *Small placement instance*: The twelve-site QUBO is solvable exactly, so it cannot demonstrate a regime where heuristics are necessary. The negative QAOA result should not be extrapolated to instances where exhaustive search is infeasible.
- 6) *No exogenous meteorology*: As argued in Section IX-B, this likely caps achievable skill for all models and is the dominant limitation on forecast quality.

XI. CONCLUSION AND FUTURE WORK

We have presented Q-AIR, an evidence-gated hybrid quantum–classical framework for air-quality intelligence, and executed four quantum modules under a protocol designed to make advantage claims falsifiable. The quantum kernel attained the lowest nominal forecasting error of any model tested and nonetheless failed the evidence gate against a naive persistence baseline. A bandwidth-resolved analysis identified the mechanism: predictive skill is monotonically decreasing in quantum expressivity, and the kernel is useful only where it has degenerated to effective rank near unity.

Amplitude estimation reproduced its predicted convergence advantage cleanly, at $N^{-0.868}$ against $N^{-0.519}$, while QAOA was decisively outperformed by greedy submodular maximisation on an instance the classical solver renders exactly in milliseconds.

Two contributions survive independently of any quantum claim. The target-leakage theorem and its 0.42 R² gap identify a defect that invalidates a substantial fraction of published AQI-prediction results. The evidence-gate operator, combined with purged blocked validation and pooled out-of-fold metrics, provides a reusable protocol under which quantum and classical estimators can be compared without the asymmetries that currently favour the former.

Future work proceeds along four lines. First, replication on the raw Amaravati record and extension to multi-year, multi-station panels from the CPCB network, which would permit genuine spatial and inter-annual validation. Second, incorporation of exogenous meteorology and fire-count data, which Section IX-B identifies as the binding constraint on forecast skill. Third, evaluation of projected quantum kernels [6] and bandwidth-conditioned feature maps under the same gate, to test whether the monotone degradation of Fig. 5 is a property of this feature-map family or of fidelity kernels generally. Fourth, and in our judgement most promising, instrumentation of a station with a colloidal quantum-dot NO₂ channel to close the schema gap of Section IV-A, converting the quantum contribution from a computational hypothesis into a metrological one.

The broader conclusion is that quantum technology should enter environmental decision systems as a replaceable service subject to explicit evidentiary conditions, and that a negative gate result is a scientific finding rather than an experimental failure. On present evidence, the near-term quantum contribution to air-quality intelligence is metrological rather than computational.

Reproducibility. All experiments use fixed seeds. Quantum circuits are simulated in PennyLane 0.45.1 on default.qubit with statevector backends; classical comparators use scikit-learn 1.8.0. The reconstruction protocol, feature construction, fold definitions, bandwidth sweep, bootstrap gate and all figure code are deterministic given the reported seeds.

REFERENCES

- [1] O. Farooq, M. Shahid, S. Arshad, A. Altaf, F. Iqbal, Y. A. Vera, M. A. L. Flores, and I. Ashraf, "An enhanced approach for predicting air pollution using quantum support vector machine," *Sci. Rep.*, vol. 14, no. 1, art. 19521, 2024.
- [2] Y. Mohamed et al., "Quantum machine learning regression optimisation for full-scale sewage sludge anaerobic digestion," *npj Clean Water*, vol. 8, art. 17, 2025.

- [3] V. Havlíček, A. D. Córcoles, K. Temme, A. W. Harrow, A. Kandak, J. M. Chow, and J. M. Gambetta, "Supervised learning with quantum-enhanced feature spaces," *Nature*, vol. 567, pp. 209–212, 2019.
- [4] M. Schuld and N. Killoran, "Quantum machine learning in feature Hilbert spaces," *Phys. Rev. Lett.*, vol. 122, no. 4, art. 040504, 2019.
- [5] P. Rebentrost, M. Mohseni, and S. Lloyd, "Quantum support vector machine for big data classification," *Phys. Rev. Lett.*, vol. 113, art. 130503, 2014.
- [6] H.-Y. Huang, M. Broughton, M. Mohseni, R. Babbush, S. Boixo, H. Neven, and J. R. McClean, "Power of data in quantum machine learning," *Nat. Commun.*, vol. 12, art. 2631, 2021.
- [7] J. M. Kübler, S. Buchholz, and B. Schölkopf, "The inductive bias of quantum kernels," in *Adv. Neural Inf. Process. Syst.*, vol. 34, pp. 12661–12673, 2021.
- [8] S. Thanasi, S. Wang, M. Cerezo, and Z. Holmes, "Exponential concentration in quantum kernel methods," *Nat. Commun.*, vol. 15, art. 5200, 2024.
- [9] R. Shaydulin and S. M. Wild, "Importance of kernel bandwidth in quantum machine learning," *Phys. Rev. A*, vol. 106, art. 042407, 2022.
- [10] M. Cerezo, A. Arrasmith, R. Babbush, S. C. Benjamin, S. Endo, K. Fujii, J. R. McClean, K. Mitarai, X. Yuan, L. Cincio, and P. J. Coles, "Variational quantum algorithms," *Nat. Rev. Phys.*, vol. 3, pp. 625–644, 2021.
- [11] K. Mitarai, M. Negoro, M. Kitagawa, and K. Fujii, "Quantum circuit learning," *Phys. Rev. A*, vol. 98, art. 032309, 2018.
- [12] M. Schuld, V. Bergholm, C. Gogolin, J. Izaac, and N. Killoran, "Evaluating analytic gradients on quantum hardware," *Phys. Rev. A*, vol. 99, art. 032331, 2019.
- [13] J. R. McClean, S. Boixo, V. N. Smelyanskiy, R. Babbush, and H. Neven, "Barren plateaus in quantum neural network training landscapes," *Nat. Commun.*, vol. 9, art. 4812, 2018.
- [14] M. Cerezo, A. Sone, T. Volkoff, L. Cincio, and P. J. Coles, "Cost function dependent barren plateaus in shallow parametrized quantum circuits," *Nat. Commun.*, vol. 12, art. 1791, 2021.
- [15] A. Pérez-Salinas, A. Cervera-Lierta, E. Gil-Fuster, and J. I. Latorre, "Data re-uploading for a universal quantum classifier," *Quantum*, vol. 4, art. 226, 2020.
- [16] M. Cerezo, M. Larocca, D. García-Martín, et al., "Does provable absence of barren plateaus imply classical simulability?" *Nat. Commun.*, vol. 16, art. 7907, 2025.
- [17] G. Brassard, P. Høyer, M. Mosca, and A. Tapp, "Quantum amplitude amplification and estimation," *Contemp. Math.*, vol. 305, pp. 53–74, 2002.
- [18] Y. Suzuki, S. Uno, R. Raymond, T. Tanaka, T. Onodera, and N. Yamamoto, "Amplitude estimation without phase estimation," *Quantum Inf. Process.*, vol. 19, art. 75, 2020.
- [19] D. Grinko, J. Gacon, C. Zoufal, and S. Woerner, "Iterative quantum amplitude estimation," *npj Quantum Inf.*, vol. 7, art. 52, 2021.
- [20] A. Montanaro, "Quantum speedup of Monte Carlo methods," *Proc. R. Soc. A*, vol. 471, art. 20150301, 2015.
- [21] S. Woerner and D. J. Egger, "Quantum risk analysis," *npj Quantum Inf.*, vol. 5, art. 15, 2019.
- [22] R. Babbush, J. R. McClean, M. Newman, C. Gidney, S. Boixo, and H. Neven, "Focus beyond quadratic speedups for error-corrected quantum advantage," *PRX Quantum*, vol. 2, art. 010103, 2021.
- [23] E. Farhi, J. Goldstone, and S. Gutmann, "A quantum approximate optimization algorithm," *arXiv:1411.4028*, 2014.
- [24] A. Krause, A. Singh, and C. Guestrin, "Near-optimal sensor placements in Gaussian processes: Theory, efficient algorithms and empirical studies," *J. Mach. Learn. Res.*, vol. 9, pp. 235–284, 2008.
- [25] E. Tang, "A quantum-inspired classical algorithm for recommendation systems," in *Proc. 51st ACM Symp. Theory of Computing (STOC)*, 2019, pp. 217–228.
- [26] E. Tang, "Quantum principal component analysis only achieves an exponential speedup because of its state preparation assumptions," *Phys. Rev. Lett.*, vol. 127, art. 060503, 2021.
- [27] N.-H. Chia, A. Gilyén, T. Li, H.-H. Lin, E. Tang, and C. Wang, "Sampling-based sublinear low-rank matrix arithmetic framework for dequantizing quantum machine learning," in *Proc. 52nd ACM Symp. Theory of Computing (STOC)*, 2020, pp. 387–400.
- [28] S. Aaronson, "Read the fine print," *Nat. Phys.*, vol. 11, pp. 291–293, 2015.
- [29] M. Schuld and N. Killoran, "Is quantum advantage the right goal for quantum machine learning?" *PRX Quantum*, vol. 3, art. 030101, 2022.
- [30] J. Bowles, S. Ahmed, and M. Schuld, "Better than classical? The subtle art of benchmarking quantum machine learning models," *arXiv:2403.07059*, 2024.
- [31] F. Mitri, A. De Iacovo, M. De Luca, A. Pecora, and L. Colace, "Lead sulphide colloidal quantum dots for room temperature NO₂ gas sensors," *Sci. Rep.*, vol. 10, art. 12556, 2020.
- [32] A. Belsley, "Quantum-enhanced absorption spectroscopy with bright squeezed frequency combs," *Phys. Rev. Lett.*, vol. 130, art. 133602, 2023.
- [33] D. I. Herman et al., "Squeezed dual-comb spectroscopy," *Science*, vol. 387, no. 6734, art. eads6292, 2025.
- [34] D. Kim, M. I. Ibrahim, C. Foy, M. E. Trusheim, R. Han, and D. R. Englund, "A CMOS-integrated quantum sensor based on nitrogen–vacancy centres," *Nat. Electron.*, vol. 2, pp. 284–289, 2019.
- [35] Central Pollution Control Board, Govt. of India, "National Air Quality Index," *Report of the Expert Group*, New Delhi, 2014.
- [36] World Health Organization, WHO Global Air Quality Guidelines: Particulate Matter (PM_{2.5} and PM₁₀), Ozone, Nitrogen Dioxide, Sulfur Dioxide and Carbon Monoxide. Geneva: WHO, 2021.
- [37] V. Bergholm et al., "PennyLane: Automatic differentiation of hybrid quantum-classical computations," *arXiv:1811.04968*, 2022.
- [38] F. Pedregosa et al., "Scikit-learn: Machine learning in Python," *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011.