

Quantum Kernel Learning for Parkinsonian Data: A Five-Qubit Benchmark for Parkinson's Progression and Exploratory PSP Differentiation

Kumar Sanu¹, Madhuvan Dixit²

¹M Tech Scholar, ²Head and Professor,

^{1,2}Department of IT, NIIST, Bhopal, Madhya Pradesh, India

ABSTRACT

Early differentiation of Parkinson's disease (PD) from progressive supranuclear palsy (PSP), and reliable stratification of PD progression remain difficult because clinically heterogeneous measurements may carry limited class information. Quantum kernel methods are frequently proposed for biomedical classification, but their value depends on the data representation and must be established against tuned classical baselines. Methods: Two tabular analyses contained in the supplied research chapters were reconstructed as an integrity-preserving benchmark. The primary dataset comprised 500 records, 12 columns, no missing values, and three approximately balanced PD progression classes (154, 164, and 182 observations). Patient identifiers were excluded; categorical and numerical fields were encoded and standardized; an 80:20 stratified split and five-fold tuning were used. Principal component analysis reduced the predictors to five components (61.10% variance), enabling five-qubit fidelity-kernel classifiers using ZFeatureMap, ZZFeatureMap, PauliFeatureMap, and an Ry–Rz/CNOT map. Logistic regression, random forest, radial-basis-function SVM, and XGBoost served as classical comparators. A secondary feasibility dataset used three whole-blood probes from GSE6613 after metadata resolution, retaining 22 healthy controls, 50 PD cases, and six PSP cases for classical five-fold assessment. Results: On the progression task, the best classical accuracy was 35.00% (random forest), while the best quantum-kernel accuracy was 36.00% (ZZFeatureMap). ZFeatureMap achieved the highest macro-AUC (0.5488). All results were close to the 33.3% chance accuracy expected for three balanced classes; the proposed Ry–Rz/CNOT map achieved 33.00% accuracy and AUC 0.4868. In the secondary dataset, random forest achieved 64.10% overall accuracy but only 37.58% macro-sensitivity, reflecting severe PSP imbalance. Conclusion: Quantum feature-map complexity did not overcome weak endpoint signal, and the study provides no evidence of quantum advantage or clinical readiness. The principal contribution is a transparent negative benchmark and a validation roadmap for future multimodal, longitudinal PD–PSP studies.

How to cite this paper: Kumar Sanu | Madhuvan Dixit "Quantum Kernel Learning for Parkinsonian Data: A Five-Qubit Benchmark for Parkinson's Progression and Exploratory PSP Differentiation" Published in International

Journal of Trend in Scientific Research and Development (ijtsrd), ISSN: 2456-6470, Volume-10 | Issue-5, October 2026, pp.35-42,

www.ijtsrd.com/papers/ijtsrd142104.pdf



URL:

Copyright © 2026 by author (s) and International Journal of Trend in Scientific Research and Development Journal. This is an Open Access article distributed under the terms of the Creative Commons Attribution License (CC BY 4.0) (<http://creativecommons.org/licenses/by/4.0>)



KEYWORDS: Parkinson's disease; progressive supranuclear palsy; quantum kernel; quantum support vector classifier; disease progression; feature map; class imbalance; negative results.

I. INTRODUCTION

Parkinson's disease is clinically defined by parkinsonism-bradykinesia with rest tremor, rigidity, or both—followed by evaluation of supportive features, red flags, and exclusion criteria (Postuma et al. 2015). Progressive supranuclear palsy is a four-repeat tauopathy with multiple phenotypes, including PSP-

parkinsonism, that may resemble PD during early disease (Höglinger et al. 2017; Boxer et al. 2017). The overlap is clinically consequential: prognosis, levodopa responsiveness, fall and swallowing risk, counselling, and trial eligibility differ. Diagnostic uncertainty is greatest before vertical supranuclear

gaze abnormalities, early postural instability, axial-predominant rigidity, or characteristic imaging patterns become sufficiently evident.

The computational problem is related but not identical. A classifier may discriminate clean research cohorts while failing in the early, variant, or uncertain cases for which decision support is needed. Likewise, progression labels may be statistically balanced but clinically weak if they are not derived from repeated, validated outcomes. The Parkinson's Progression Markers Initiative was established as a longitudinal, multicentre resource integrating clinical, imaging, biological, genetic, and digital measures (Marek et al. 2011, 2018). Its design illustrates why progression modeling should ordinarily be anchored to within-person change rather than a poorly documented cross-sectional class field.

Support vector machines are attractive for modest biomedical datasets because margin regularization and kernels can represent nonlinear boundaries without training a large deep network. Quantum kernel learning extends this logic: a circuit maps classical features to quantum states, their overlaps define a kernel matrix, and a classical SVM learns the decision boundary (Havlíček et al. 2019; Schuld and Killoran 2019). Current implementations are hybrid rather than end-to-end quantum; Qiskit's QSVC, for example, extends the conventional support-vector classifier with a quantum-kernel argument. A large Hilbert space is mathematically interesting, but it does not guarantee better generalization on classical clinical data. Classical models can remain competitive even when quantum representations are expressive (Huang et al. 2021).

The supplied chapters originally proposed early HC–PD–PSP diagnosis and progression assessment. Their results section, however, contained a mixture of calculated experiments and explicitly hypothetical 82–82.5% diagnostic values. This article excludes all hypothetical confusion matrices, early-detection metrics, and superiority claims. It addresses the narrower empirical question: how do four five-qubit feature maps behave on the available three-class PD progression task, and what does a severely imbalanced HC–PD–PSP expression dataset reveal about the feasibility of PSP modeling? This reframing protects the contribution from overclaiming while retaining the methodologically informative result: quantum feature maps did not manufacture signal absent from the endpoint.

II. Related Work and Rationale

Clinical and data context

The 2015 Movement Disorder Society criteria systematized clinical PD diagnosis around a motor

core, supportive criteria, red flags, and exclusions (Postuma et al. 2015). PSP criteria later expanded diagnostic recognition beyond Richardson syndrome to variant phenotypes and earlier certainty levels (Höglinger et al. 2017). These advances improve clinical standardization but do not remove label uncertainty, especially in cross-sectional datasets and variant disease. For a machine-learning study, the label origin, adjudication date, disease duration, medication state, and follow-up horizon therefore matter as much as the algorithm.

PPMI provides a rich longitudinal research ecosystem, but the specific 500-record file summarized in the chapters cannot automatically be assumed to be an official PPMI curated extract. It contains a patient identifier, demographic and clinical variables, and a three-level Disease_Progression outcome, yet the chapter files do not preserve a data dictionary, acquisition date, endpoint derivation, or versioned extraction manifest. We therefore refer to it as the supplied PD progression dataset rather than asserting provenance that cannot be independently confirmed from the six chapter documents. This distinction is essential for reproducibility.

The secondary expression analysis uses GSE6613, an NIH Gene Expression Omnibus study of whole blood from PD, other neurodegenerative diseases, and healthy controls. The supplied table contained only three probes and lacked a diagnosis field. Sample identities were matched to GEO metadata, and six neurological-control samples were identified as PSP. This produces a technically analyzable but extremely fragile minority class; conclusions must therefore be restricted to feasibility and failure-mode analysis.

Quantum kernels

For a feature vector x , a circuit $U\phi(x)$ prepares $|\phi(x)\rangle = U\phi(x)|0\rangle^{\otimes n}$. A fidelity kernel can be written $K(x,z) = |\langle\phi(x)|\phi(z)\rangle|^2$. The kernel matrix is then supplied to a support-vector optimizer. ZFeatureMap encodes single-feature phase structure; ZZFeatureMap introduces pairwise interactions; PauliFeatureMap generalizes the chosen Pauli operators and entanglement pattern; and the proposed circuit uses Ry and Rz rotations with CNOT entanglement. Repetitions increase circuit depth and representation complexity but may also increase redundancy, kernel concentration, noise sensitivity, and computation.

Havlíček et al. (2019) established supervised learning in quantum-enhanced feature spaces, while subsequent work showed both the promise and caveats of quantum kernels. Huang et al. (2021) demonstrated that access to data can make classical learners competitive even when a target is generated

by a quantum process. Peters et al. (2021) implemented high-dimensional quantum-kernel learning on noisy hardware and highlighted scaling and error-mitigation challenges. Biomedical demonstrations have generally reported competitive

rather than consistently superior quantum performance (Vasques et al. 2023). A valid comparison must therefore tune classical baselines, lock test data, report uncertainty, and treat simulator performance separately from hardware utility.

III. Materials and Methods

Study design and reporting boundary

This work is a retrospective secondary analysis of two tabular datasets documented in the supplied chapters. No new participants were recruited and no clinical intervention was performed. The primary endpoint was the three-class Disease_Progression label in the 500-record PD dataset. The secondary endpoint was three-class HC–PD–PSP diagnosis derived by metadata matching for the limited GSE6613 probe table. Analyses labeled in the chapters as hypothetical, approximate, assumed, or expected were excluded a priori from this article.

Analysis	Records/classes	Predictors	Validation	Purpose
Primary PD progression	n=500; Class 1=154, Class 2=164, Class 3=182	Available demographic and clinical fields; Patient_ID excluded	80:20 stratified split; 5-fold tuning	Compare classical models and five-qubit feature maps
Secondary HC–PD–PSP	n=78; HC=22, PD=50, PSP=6	Three expression probes: 204252_at, 211803_at, 211804_s_at	Stratified 5-fold cross-validation	Assess imbalance and data sufficiency; classical models only

Table 1 Study datasets and empirically supported analysis boundaries.

Primary-data preprocessing

The supplied PD file contained 500 rows, 12 columns, and no missing values. Patient_ID was removed to prevent identity leakage. Numerical predictors were standardized; categorical predictors were one-hot encoded. A stratified 80:20 partition produced 400 training and 100 test observations while preserving the three-class proportions. Hyperparameters were selected within the training data using five-fold stratified cross-validation. The available chapter evidence does not state whether preprocessing was implemented inside a formal pipeline for every cross-validation fold; this is retained as a reproducibility limitation rather than silently assumed.

For quantum-kernel evaluation, the processed predictor matrix was reduced by principal component analysis to five components, matching a five-qubit circuit. The components retained approximately 61.10% of total variance and were scaled to $[0, \pi]$ for rotation encoding. Feature-map repetitions $r \in \{1, 2, 3\}$ and SVM regularization $C \in \{0.1, 1, 10\}$ were selected using training-set cross-validation with macro-F1 as the optimization target.

Models

Classical comparators were balanced logistic regression, random forest, radial-basis-function SVM, and XGBoost. The selected configurations reported in the chapters were: logistic regression $C=0.01$ with balanced weighting; random forest with 200 trees, minimum leaf size 5, and balanced class weighting; RBF-SVM with $C=1$ and $\gamma=0.1$; and XGBoost with 300 trees, maximum depth 6, and learning rate 0.03.

Four fidelity-kernel classifiers were examined: ZFeatureMap, ZZFeatureMap, PauliFeatureMap, and a custom Ry–Rz/CNOT map. The chosen repetitions and C values were determined by cross-validated macro-F1. The kernel matrix represents pairwise state fidelity and is passed to a classical support-vector classifier, consistent with the hybrid QSVC paradigm. The documents describe simulator-based evaluation and do not supply a reproducible hardware backend, noise model, shot count, software-version lockfile, or circuit transpilation record; no hardware claim is made.

Hybrid Quantum-Kernel Evaluation Pipeline



Fig. 1 Leakage-aware hybrid evaluation pipeline reconstructed from the supplied methodology. Preprocessing and model selection must remain training-fold operations.

Secondary GSE6613 feasibility analysis

The secondary file initially contained 105 samples and three expression probes but no class label. Sample identifiers were matched to official GSE6613 metadata. Twenty-two healthy controls, 50 PD cases, and six PSP

cases were retained; 27 samples representing other neurological diagnoses were excluded from the three-class analysis. Because an 80:20 split would place approximately one PSP case in the test set, stratified five-fold cross-validation was used. The same four classical model families were evaluated, and macro sensitivity, specificity, F1, and one-versus-rest AUC were used to reduce domination by the majority PD class.

Performance measures

Accuracy was the fraction of correct predictions. Multiclass sensitivity, specificity, F1, and AUC were macro-averaged so that each class received equal weight. Macro sensitivity is particularly important in the secondary dataset because overall accuracy can be high while the six-case PSP class is poorly recognized. For the primary task, the approximate chance accuracy is 33.3% because the classes are nearly balanced. Point estimates are reported exactly as preserved in Chapter 5. Raw fold predictions were not attached, so confidence intervals, paired significance tests, calibration measures, and confusion matrices could not be reconstructed.

IV. Results

Classical progression-classification baselines

Model	Accuracy (%)	Macro sensitivity (%)	Macro specificity (%)	Macro F1 (%)	Macro AUC
Logistic regression	30.00	29.73	65.02	29.81	0.4498
Random forest	35.00	34.68	67.54	34.74	0.4824
Classical RBF-SVM	26.00	24.94	62.64	23.02	0.5153
XG Boost	33.00	32.91	66.41	32.89	0.4694

Table 2 Classical performance on the held-out 100-record PD progression test set.

Random forest produced the highest classical accuracy (35.00%), macro sensitivity (34.68%), macro specificity (67.54%), and macro-F1 (34.74%). Classical SVM obtained the highest macro-AUC (0.5153) despite the lowest accuracy (26.00%). These values are close to chance and do not support a clinically meaningful progression classifier. The divergence between accuracy and AUC also cautions against selecting a model on one metric alone.

Quantum feature-map comparison

Feature map	Qubits	Repetitions	C	Accuracy (%)	Macro F1 (%)	Macro AUC
ZFeatureMap	5	3	1	35.00	34.36	0.5488
ZZFeatureMap	5	1	1	36.00	34.76	0.4401
PauliFeatureMap	5	3	1	34.00	32.82	0.5297
Ry-Rz + CNOT	5	3	10	33.00	32.48	0.4868

Table 3 Quantum feature-map results on the same progression task.

ZZFeatureMap achieved the highest accuracy (36.00%) and macro-F1 (34.76%), improving accuracy by one percentage point over the random-forest baseline. ZFeatureMap achieved the highest AUC (0.5488), but its accuracy was 35.00%. The custom Ry-Rz/CNOT map did not outperform standard maps: its 33.00% accuracy was essentially at chance and its AUC was below 0.50. Additional circuit complexity therefore provided no observable benefit under the evaluated configuration.

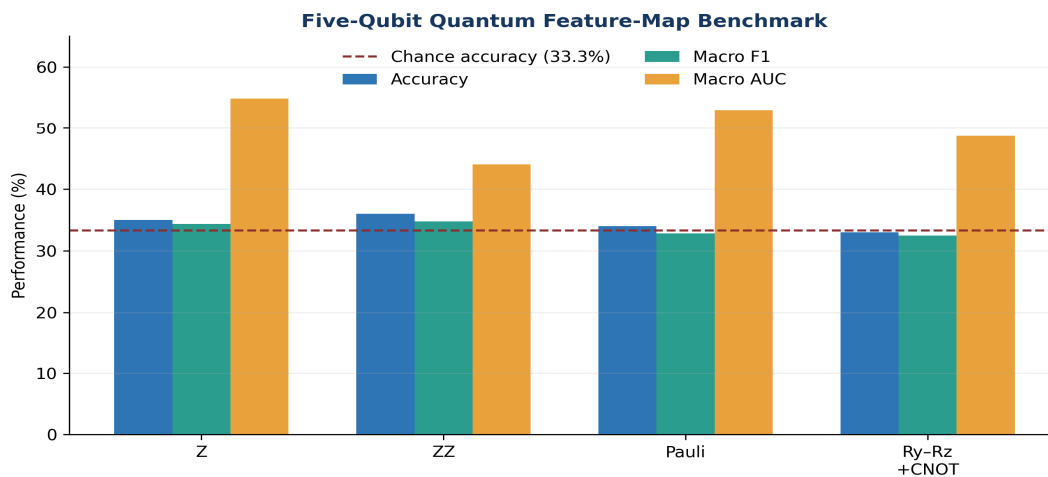


Fig. 2 Feature-map comparison. AUC is displayed as a percentage for visual comparability; the dashed line denotes chance-level accuracy, not chance AUC.

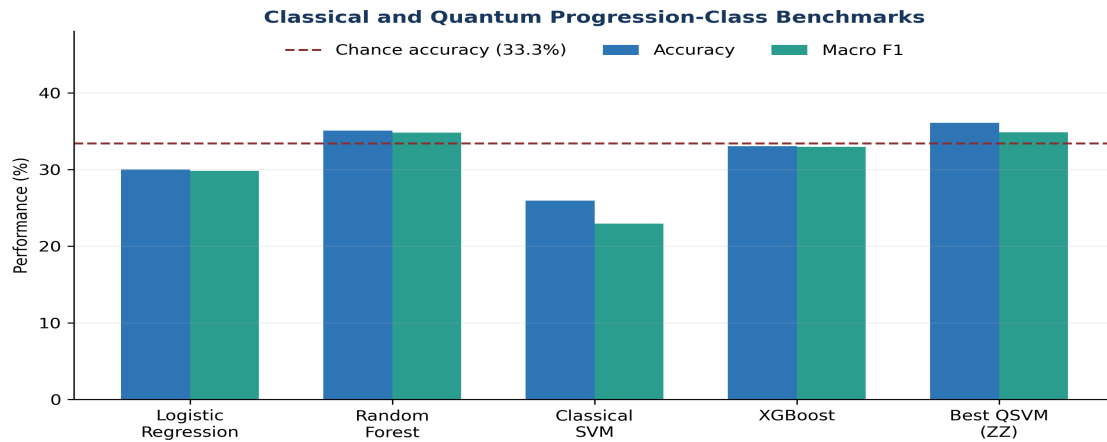
Classical-versus-quantum interpretation

Fig. 3 Best-performing quantum feature map compared with classical baselines. Differences are small relative to the low-signal operating range.

The best quantum accuracy was numerically higher than the best classical accuracy, but the chapter files do not provide repeated outer folds or test-level predictions from which uncertainty can be estimated. A one-point difference cannot therefore be interpreted as superiority. More importantly, both models remained near the 33.3% chance benchmark. This is a negative result: none of the evaluated learners extracted a reliable progression signal from the supplied predictors and endpoint.

Exploratory HC–PD–PSP expression baseline

Model	Accuracy (%)	Macro sensitivity (%)	Macro specificity (%)	Macro F1 (%)	Macro AUC
Logistic regression	29.49	17.88	63.23	21.19	0.3722
Random forest	64.10	37.58	70.97	35.96	0.4752
Classical SVM	33.33	25.82	67.72	27.75	0.5289
XGBoost	52.56	33.27	67.33	33.12	0.4290

Table 4 Exploratory five-fold classical baselines for the three-probe HC–PD–PSP dataset.

Random forest achieved 64.10% overall accuracy, but macro sensitivity was only 37.58% and macro-AUC was 0.4752. This pattern is consistent with majority-class success rather than robust PSP recognition. Classical SVM had the highest macro-AUC (0.5289), only slightly above chance. With six PSP cases and three probes, the data are insufficient for credible QSVM optimization, hold-out testing, or claims about early PSP diagnosis. The correct finding is a data-sufficiency warning, not a 64.1% PSP detection claim.

V. Discussion**Principal findings**

Three findings are central. First, all progression models operated close to chance, indicating weak predictive structure in the supplied outcome and feature set. Second, quantum feature maps changed the geometry enough to alter point estimates, but no map produced a clinically meaningful improvement; the custom entangling map performed worse than simpler alternatives. Third, the apparent 64.10% performance in the HC–PD–PSP expression dataset dissolved when class-balanced metrics were considered. Macro sensitivity revealed that the rare PSP class was not adequately recognized. These results directly contradict the hypothetical 82% diagnostic narrative present. Excluding those values is not a cosmetic downgrade; it is the transformation that makes the article scientifically coherent. The paper's defensible contribution is an integrity-preserving benchmark demonstrating that representation complexity cannot compensate for uncertain endpoints, narrow feature coverage, or extreme minority-class scarcity.

Why the quantum kernel did not create advantage

A quantum kernel can be expressive while remaining poorly aligned with the target labels. The five principal components retained only 61.10% of variance, so potentially relevant structure may have been discarded before quantum encoding. Conversely, retained variance is not necessarily outcome-relevant; PCA maximizes unsupervised variance, not class separation. Scaling every component to an angular interval also changes distance structure. Z, ZZ, Pauli, and Ry–Rz/CNOT circuits then impose different inductive biases, none of which is guaranteed to match the biological process represented by the progression label.

Kernel complexity may also reduce generalization. Deeper circuits can make examples appear uniformly dissimilar, producing concentrated or nearly identity-like kernels. In small classical datasets, a high-dimensional quantum space can therefore memorize idiosyncratic differences without learning a stable boundary. The best ZZ map used only one repetition, while the three-repetition custom map performed at chance; this pattern is consistent with the idea that more gates are not automatically more information.

Finally, the benchmark cannot establish computational advantage. The experiments were described as simulator-based, and kernel estimation scales with pairwise sample comparisons. No hardware runtime, shot budget, queue time, circuit depth after transpilation, error mitigation, or energy cost was available. A claim of quantum advantage would require not only predictive benefit against strong classical kernels but also a transparent computational comparison under realistic access assumptions.

Clinical and data implications

Progression is inherently longitudinal. A categorical field named Disease_Progression is not sufficient unless its clinical derivation is documented. Future PD outcomes should be defined using follow-up change in MDS-UPDRS domains, Hoehn and Yahr transition, functional loss, medication-state-aware motor scores, or time-to-event endpoints. PSP progression requires disease-appropriate measures such as PSPRS, falls, dysphagia, gaze metrics, and imaging change. Disease-specific scales should not be treated as interchangeable numerical targets.

For differential diagnosis, an adequate dataset must include the challenging spectrum: PSP-parkinsonism, early PSP, other atypical parkinsonian syndromes, clinically uncertain cases, and longitudinal adjudication. Training only on typical PD and classic PSP-Richardson syndrome would exaggerate deployable performance. Multicentre acquisition, scanner/site harmonization, medication-state documentation, and patient-level splitting are mandatory. Because PSP is rare, recruitment and evaluation should be planned around precision of minority-class sensitivity rather than overall sample size.

What would constitute a stronger test

Requirement	Minimum implementation	Why it matters
Endpoint integrity	Versioned data dictionary; longitudinal derivation; blinded clinical adjudication	Prevents a model from learning an arbitrary or circular label
Leakage control	Patient-grouped nested cross-validation; all preprocessing fitted within folds	Prevents optimistic estimates from shared patients or preprocessing
External validity	Hospital/site-held-out cohort with scanner and demographic shift	Tests transportability rather than random-split similarity
Fair quantum comparison	Equal tuning budgets; classical RBF/poly kernels; ideal and noisy quantum runs	Separates feature-map value from weak baselines or simulator effects
Clinical metrics	Minority sensitivity, calibration, confidence intervals, decision curves	Connects discrimination to patient-level decision utility
Reproducibility	Raw split indices, code, seeds, package versions, circuit/shot/backend records	Makes positive and negative findings auditable

Table 5 Minimum validation standard for a future PD–PSP quantum-kernel study.

VI. Strengths and Limitations

A strength of this article is the explicit separation of calculated from hypothetical evidence. The same performance metrics are reported across classical and quantum models, and both accuracy and class-balanced measures are interpreted relative to chance. The work also shows why negative findings are valuable: they identify endpoint and data limitations before expensive hardware experiments or clinical claims.

Several limitations remain. The raw datasets, full data dictionaries, code, fold assignments, and predictions were not included among the six attached chapters, so the reported experiments could not be independently rerun. The supplied 500-record file’s provenance as

an official PPMI extract cannot be confirmed from the documents. Only one train–test partition was reported for the primary experiment, preventing confidence intervals and significance testing. PCA retained 61.10% variance, leaving substantial information outside the five-qubit representation. The secondary analysis contained only six PSP cases and three probes. No calibrated probabilities, subgroup analyses, missingness stress tests, external cohort, quantum-hardware execution, or clinical decision analysis was available.

Accordingly, the results support only an exploratory computational benchmark. They do not demonstrate early detection, clinical progression prediction, PD–PSP diagnostic utility, quantum speedup, or quantum

advantage. Any submission should preserve this boundary in the abstract, title, highlights, cover letter, and response to reviewers.

VII. Conclusion

A five-qubit quantum-kernel benchmark was reconstructed from the supplied research chapters using four feature maps and four classical comparators. The best quantum model achieved 36.00% accuracy and 34.76% macro-F1 on three-class PD progression classification; the best classical accuracy was 35.00%. Both were close to chance. The proposed Ry-Rz/CNOT feature map did not outperform standard maps. In the small HC-PD-PSP expression dataset, overall accuracy was distorted by severe class imbalance and did not establish PSP recognition.

The appropriate scientific conclusion is that quantum feature mapping did not overcome weak endpoint signal or inadequate minority-class evidence. This negative result is more useful than an unsupported 82% claim: it directs future work toward longitudinal endpoint construction, multimodal biomarkers, multicentre PSP recruitment, leakage-safe validation, calibrated clinical utility, and hardware-aware reporting. Quantum kernels remain a testable representation strategy, but clinical relevance must be earned through stronger data and stricter validation.

References

- [1] Bartkiewicz K, Gneiting C, Černoch A, Jiráková K, Lemr K, Nori F (2020) Experimental kernel-based quantum machine learning in finite feature space. *Scientific Reports* 10:12356. <https://doi.org/10.1038/s41598-020-68911-5>
- [2] Berg D, Postuma RB, Adler CH et al (2015) MDS research criteria for prodromal Parkinson's disease. *Movement Disorders* 30:1600–1611. <https://doi.org/10.1002/mds.26431>
- [3] Boxer AL, Yu J-T, Golbe LI, Litvan I, Lang AE, Höglinger GU (2017) Advances in progressive supranuclear palsy: new diagnostic criteria, biomarkers, and therapeutic approaches. *Lancet Neurology* 16:552–563. [https://doi.org/10.1016/S1474-4422\(17\)30157-6](https://doi.org/10.1016/S1474-4422(17)30157-6)
- [4] Cortes C, Vapnik V (1995) Support-vector networks. *Machine Learning* 20:273–297. <https://doi.org/10.1007/BF00994018>
- [5] Havlíček V, Córcoles AD, Temme K et al (2019) Supervised learning with quantum-enhanced feature spaces. *Nature* 567:209–212. <https://doi.org/10.1038/s41586-019-0980-2>
- [6] Heinzel S, Berg D, Gasser T et al (2019) Update of the MDS research criteria for prodromal Parkinson's disease. *Movement Disorders* 34:1464–1470. <https://doi.org/10.1002/mds.27802>
- [7] Höglinger GU, Respondek G, Stamelou M et al (2017) Clinical diagnosis of progressive supranuclear palsy: the Movement Disorder Society criteria. *Movement Disorders* 32:853–864. <https://doi.org/10.1002/mds.26987>
- [8] Huang H-Y, Broughton M, Mohseni M et al (2021) Power of data in quantum machine learning. *Nature Communications* 12:2631. <https://doi.org/10.1038/s41467-021-22539-9>
- [9] Jabbari E, Holland N, Chelban V et al (2020) Diagnosis across the spectrum of progressive supranuclear palsy and corticobasal syndrome. *JAMA Neurology* 77:377–387. <https://doi.org/10.1001/jamaneurol.2019.4347>
- [10] Jäger J, Krems RV (2023) Universal expressiveness of variational quantum classifiers and quantum kernels for support vector machines. *Nature Communications* 14:576. <https://doi.org/10.1038/s41467-023-36144-5>
- [11] Marek K, Jennings D, Lasch S et al (2011) The Parkinson Progression Marker Initiative (PPMI). *Progress in Neurobiology* 95:629–635. <https://doi.org/10.1016/j.pneurobio.2011.09.005>
- [12] Marek K, Chowdhury S, Siderowf A et al (2018) The Parkinson's progression markers initiative (PPMI): establishing a PD biomarker cohort. *Annals of Clinical and Translational Neurology* 5:1460–1477. <https://doi.org/10.1002/acn3.644>
- [13] National Center for Biotechnology Information (2019) GSE6613: whole blood expression data from Parkinson's disease, other neurodegenerative diseases, and healthy controls. *Gene Expression Omnibus*. <https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE6613>
- [14] Parkinson's Progression Markers Initiative (2026) Study design and data access. <https://www.ppmi-info.org/study-design> and <https://www.ppmi-info.org/access-data-specimens/download-data>

- [15] Peters E, Caldeira J, Ho A et al (2021) Machine learning of high dimensional data on a noisy quantum processor. npj Quantum Information 7:161. <https://doi.org/10.1038/s41534-021-00498-9>
- [16] Postuma RB, Berg D, Stern M et al (2015) MDS clinical diagnostic criteria for Parkinson's disease. Movement Disorders 30:1591–1601. <https://doi.org/10.1002/mds.26424>
- [17] Qiskit Machine Learning (2026) Fidelity Quantum Kernel and Quantum Support Vector Classifier documentation, version 0.9.1. <https://qiskit-community.github.io/qiskit-machine-learning/>
- [18] Schuld M, Killoran N (2019) Quantum machine learning in feature Hilbert spaces. Physical Review Letters 122:040504. <https://doi.org/10.1103/PhysRevLett.122.040504>
- [19] Stamelou M, Respondek G, Giagkou N, Whitwell JL, Kovacs GG, Höglinger GU (2021) Evolving concepts in progressive supranuclear palsy and other 4-repeat tauopathies. Nature Reviews Neurology 17:601–620. <https://doi.org/10.1038/s41582-021-00541-5>
- [20] Vasques X, Paik H-Y, Cifci MA et al (2023) Application of quantum machine learning using quantum kernel algorithms on multiclass biomedical data. Scientific Reports 13:12347. <https://doi.org/10.1038/s41598-023-38558-z>
- [21] Whitwell JL, Höglinger GU, Antonini A et al (2017) Radiological biomarkers for diagnosis in PSP: where are we and where do we need to be? Movement Disorders 32:955–971. <https://doi.org/10.1002/mds.27038>

