

Technical White Paper · ethereal computing

Quantum-Enhanced Machine Learning for Fidelity Optimization in Non-Invasive Neurophysiological Biosignal Acquisition

A hybrid classical-quantum architecture and validation framework for brain-computer interfacing

AUTHOR

Eros Marcello Iuliano

Moffett Field, California, USA

eros@etherealcomputing.com

VERSION

Version 4.0 · June 2026 · Confidential — Research Work Product

Status of this document. This is a technical white paper. It presents a proposed hybrid classical-quantum architecture and a validation methodology for improving signal fidelity in non-invasive biosignal acquisition. It is a position-and-design document: the evaluation described in Section 8 specifies experiments to be performed and does not report completed empirical results from ethereal computing. Statements about a possible quantum advantage are framed as hypotheses to be tested, consistent with the present state of the peer-reviewed literature, which reports promising but inconsistent gains for hybrid quantum-classical models on EEG tasks [2,3,6]. Public release of this document constitutes a disclosure event; see the IP and Disclosure note at the end.

Contents

Abstract

Executive Summary

- 1 Introduction
- 2 The Fidelity Bottleneck in Non-Invasive Biosignal Acquisition
- 3 Primer: Variational Quantum Circuits and Hybrid Models
- 4 Related Work and Honest Positioning
- 5 Proposed Hybrid Architecture
- 6 Software and Model Stack
- 7 Deployment Model and the Distillation Question
- 8 Validation Methodology
- 9 Strategic and Mission Relevance
- 10 Competitive Landscape
- 11 Limitations and Open Questions
- 12 Conclusion

References

Appendix A — Glossary of Terms

Appendix B — Proposed Reproducibility Checklist

IP and Disclosure Note

Abstract

Non-invasive electroencephalography (EEG) is the most deployable sensing modality for brain–computer interfaces (BCIs), yet its utility is capped by a fidelity bottleneck: the neural signal of interest is microvolt-scale and is routinely buried under muscle (EMG) artifacts, mains interference, electrode-impedance drift, and motion, all of which worsen with the dry-electrode systems preferred for field use. Classical denoising and decoding pipelines—-independent component analysis, adaptive filtering, and compact convolutional networks such as EEGNet [1]—have improved matters but show diminishing returns as the noise floor, rather than model capacity, becomes the binding constraint. This white paper proposes a hybrid classical–quantum signal-processing layer that augments proven temporal neural architectures with trainable variational quantum circuits (VQCs) inserted at two specific points, and a deployment lifecycle that confines quantum computation to the training and design phases while exporting a classical model for low-latency edge inference.

We make three contributions. First, we specify a modular insertion scheme that places a VQC after the depthwise-separable convolutional feature stage (a quantum projection head) and, separately, as the value projection of a single attention block, motivated directly by QEEGNet [2,3] and the QASA transformer [6]. Second, we define a train-time-quantum / inference-time-classical lifecycle and give an explicit, honest analysis of when this preserves any quantum contribution and when it merely reproduces a classical operator. Third—and most importantly for a credible engineering claim—we convert the marketing assertion that “quantum helps” into a falsifiable validation protocol: matched-capacity classical baselines under equal compute and tuning, contribution-isolating ablations, multi-seed statistics with reported effect sizes, and pre-registered success criteria. We do not claim a demonstrated quantum advantage for EEG; the peer-reviewed evidence to date is promising but inconsistent [3,6], and we treat the magnitude of any gain as the central open question this program is designed to answer.

Executive Summary

ethereal computing is developing a software layer that improves the quality of brain and biopotential signals collected from wearable, electrode-based sensors. The problem it addresses is concrete and widely acknowledged: outside the laboratory, EEG is noisy, and the noise—not the algorithms—limits how reliably a system can read attention, fatigue, intent, or motor commands. This is the same limitation that defense and aerospace programs cite when they describe non-invasive neural interfaces as promising but not yet precise enough for operational use [7].

Our approach is deliberately pragmatic. Rather than promising an all-quantum system—which is not feasible on today’s hardware—we combine well-established classical neural networks with small quantum circuits, using the quantum component only where it is most plausibly useful: learning a richer mapping that separates signal from artifact. The quantum work happens during training and model design. The model that ships to a headset or field device is a conventional, optimized classical network exported through standard tooling (ONNX), so no quantum processor is required at the point of use.

This document is intentionally candid about evidence. The published research that inspires this work—most directly QEEGNet, a quantum-augmented version of the EEGNet architecture—reports gains that are real in places but inconsistent across datasets, and its own authors describe it as a forward-looking

approach [2,3]. We therefore frame our quantum claims as hypotheses and devote a full section to the experiments that would confirm or refute them, including the uncomfortable but necessary test of whether a purely classical network of the same size performs just as well. A white paper that skips that test is marketing; one that specifies it is engineering. We have chosen the latter.

The intended readers are technical evaluators, prospective research and development partners, and program managers in clinical neurotechnology, defense, and aerospace. For those audiences, the value here is a clear architecture, an honest map of what is known versus assumed, and a validation plan rigorous enough to support a credible go/no-go decision.

SECTION 1

Introduction

A brain–computer interface turns measured neural activity into a control or monitoring signal. The most accessible measurements come from electroencephalography: electrodes on the scalp record the summed electrical activity of cortical neuron populations. EEG is inexpensive, portable, and carries no surgical risk, which is precisely why it dominates real-world BCI and neuroadaptive applications. The cost of that accessibility is signal quality. By the time cortical activity reaches a scalp electrode it has been attenuated and smeared by intervening tissue, and it arrives mixed with interference that is frequently larger than the signal itself. The central engineering problem of practical BCI is therefore not primarily one of classification cleverness; it is one of fidelity—recovering a faithful estimate of the neural signal before any downstream decoder can act on it.

Two decades of signal processing and, more recently, deep learning have pushed EEG decoding forward [16,17]. Compact convolutional networks such as EEGNet [1] encapsulate domain knowledge—frequency-selective and spatial filtering—in a small, data-efficient model and have become a standard backbone. Yet the field increasingly encounters a ceiling: when the limiting factor is the noise floor rather than representational capacity, adding classical parameters yields little. This motivates interest in qualitatively different feature maps. Quantum machine learning offers one such avenue. A variational quantum circuit maps inputs into a high-dimensional Hilbert space whose structure is not efficiently reproducible by small classical layers [11,12], which in principle could separate signal and artifact components that a comparable classical map entangles.

This is a genuinely promising idea, and it is also one that the field has not yet settled. The most directly relevant result, QEEGNet [2], augments EEGNet with a quantum encoding layer and reports improved accuracy and noise robustness on a motor-imagery benchmark; a follow-up study across multiple datasets [3] found that the improvements “remain inconsistent,” and the authors are explicit that the model is forward-looking rather than a settled advance. An honest white paper must hold both facts at once: the direction is worth pursuing, and the magnitude of benefit is unproven.

Accordingly, this paper does three things. Section 5 specifies a concrete hybrid architecture—where the quantum circuits sit, how they are trained, and how the system is exported for deployment. Sections 6 and 7 describe the software stack and a deployment model suited to field and mission conditions, including a frank analysis of a design choice (freezing the trained quantum layer into a classical operator for inference) that is central to the pitch and that deserves scrutiny. Section 8 sets out the validation methodology that would establish whether the architecture delivers measurable gains over strong classical baselines. We state our contributions mechanistically. First, an insertion scheme that places trainable VQCs at the

convolutional projection and at a single attention value-projection to enrich the learned representation for signal/artifact separation. Second, a train-time-quantum / inference-time-classical lifecycle with an explicit account of when quantum structure survives export. Third, a falsifiable evaluation protocol that subjects every claimed gain to matched-capacity classical comparison, ablation, and multi-seed statistics.

SECTION 2

The Fidelity Bottleneck in Non-Invasive Biosignal Acquisition

Understanding why a quantum-augmented layer might help requires first being precise about what degrades EEG. The scalp EEG of interest is typically on the order of tens of microvolts and is organized into overlapping rhythmic bands—delta (roughly 1–4 Hz), theta (4–8 Hz), alpha (8–12 Hz), beta (13–30 Hz), and gamma (above 30 Hz)—that index different neural processes. Several categories of interference contaminate this signal, and they differ enough that no single filter removes them all.

Muscle activity (EMG) produces broadband, high-frequency energy that overlaps the beta and gamma bands and can dwarf the cortical signal during jaw, neck, or facial movement. Ocular activity—blinks and saccades—injects large, low-frequency transients concentrated at frontal sites. Mains interference adds a narrowband component at 50 or 60 Hz and its harmonics. Electrode–skin impedance introduces low-frequency drift and thermal noise, and crucially it is not stationary: as gel dries or contact shifts, the noise characteristics change over the course of a recording. Motion modulates all of the above by physically perturbing the electrode–skin interface.

Dry-electrode systems sharpen this problem. They are strongly preferred for wearable and field deployment because they need no conductive gel and can be donned quickly, but they exhibit substantially higher and more variable contact impedance than wet electrodes. The neural waveforms they capture are the same; the noise environment is worse. A method that is robust specifically to non-stationary, dry-electrode interference is therefore worth more in the field than a method tuned to clean laboratory recordings.

Classical practice addresses these sources with a layered pipeline: band-pass and notch filtering for mains and out-of-band energy; regression or blind-source-separation methods such as independent component analysis (ICA) to isolate and remove ocular and some muscle components [15]; and learned models—CNNs and recurrent or transformer networks—for the joint denoising-and-decoding stage [16,17]. Each step is effective within its assumptions and limited outside them. ICA assumes a sufficient number of channels and statistically independent sources; adaptive filters assume a usable reference; supervised networks assume that training noise resembles deployment noise. The residual that survives this stack—non-Gaussian, non-stationary interference for which no clean reference exists—is exactly the regime in which the unsupervised, distribution-learning behavior of quantum-inspired filters such as the recurrent quantum neural network (RQNN) [4,5] was originally motivated.

The practical consequence of residual noise is not abstract. Lower effective signal-to-noise ratio (SNR) translates into lower decoding accuracy, higher false-positive rates on event detection, longer calibration sessions, and faster degradation of performance as conditions drift. For a soldier-worn or astronaut-worn system, those costs are operational, not merely statistical. The fidelity bottleneck is, in this sense, the problem worth solving; decoding accuracy follows from it.

SECTION 3

Primer: Variational Quantum Circuits and Hybrid Models

Because the proposed architecture rests on variational quantum circuits, this section builds the concept up from first principles before the later sections rely on it. Readers fluent in quantum machine learning may proceed to Section 4.

A quantum computer manipulates qubits, whose joint state is a vector in a complex space of dimension 2^n for n qubits. Operations are unitary transformations of that state, and the only way to extract classical information is measurement, which yields a sample from a probability distribution determined by the state. A variational quantum circuit—also called a parameterized quantum circuit or a quantum neural network—is a sequence of such operations in which some gates carry tunable parameters, typically rotation angles. Concretely, a VQC has three stages. A data-encoding stage maps a classical input vector into the quantum state, for example by using each input feature as a rotation angle; this choice of feature map largely determines what functions the circuit can represent [11,12]. A parameterized stage applies trainable single-qubit rotations interleaved with entangling two-qubit gates (such as CNOTs), which couple the qubits so that the state can represent correlations no product of independent features could. A measurement stage reads out expectation values—commonly of Pauli-Z operators on each qubit—that become the circuit’s output vector.

Writing the trainable part as a unitary $U(\theta)$ built from layers, the circuit produces, for input x and parameters θ , a vector of measured expectations $f(x, \theta)$. The appeal is expressivity: because the encoding lifts data into a 2^n -dimensional space and entanglement correlates the coordinates, a VQC can represent feature maps whose classical simulation may be costly, which is the formal sense in which quantum models are sometimes argued to access richer hypothesis classes [11]. The caution is that expressivity alone is not advantage; a model class can be large and still be no more useful than a well-chosen classical map on a given task.

Crucially for engineering, VQCs train by ordinary gradient descent. The parameters θ are updated to minimize a loss, and the gradients of $f(x, \theta)$ with respect to θ can be computed on quantum hardware via the parameter-shift rule [10] or, in simulation, by automatic differentiation. This is what makes hybrid models practical: a classical network and a quantum circuit can be composed into one computational graph and trained end-to-end, with the classical optimizer adjusting both classical weights and quantum gate angles together [8,13]. The quantum layer behaves, from the training loop’s perspective, like any other differentiable module.

Three hard constraints bound what is achievable today, and the architecture is designed around them. We are in the noisy intermediate-scale quantum (NISQ) era [14]: available processors have tens to a few hundred noisy qubits without error correction, so circuits must be shallow and few-qubit to produce usable signal. Measurement is statistical, so each expectation value requires many repeated executions (shots), and shot noise adds variance to both the forward pass and the gradient. And training itself can fail through barren plateaus—a phenomenon in which gradients vanish exponentially as circuits grow in width or depth, leaving the optimizer with no usable signal [9]. A credible hybrid design keeps quantum circuits small and well-initialized precisely to stay clear of these failure modes, and treats them as the first risks to test, not footnotes.

SECTION 4

Related Work and Honest Positioning

Four lines of prior work bound this proposal, and the positioning below states for each what it established and what it did not, so that the present contribution is not overstated.

Compact classical EEG networks

EEGNet [1] introduced a compact convolutional architecture using depthwise and separable convolutions to encode spatial and frequency-selective filtering, generalizing across multiple BCI paradigms with limited training data. Related deep models for EEG decoding—DeepConvNet and ShallowConvNet [16] and the broader literature surveyed by Roy et al. [17]—define the strong classical baselines any new method must beat. These works do not address quantum feature maps; they are the comparison set, and notably any quantum-augmented model is typically built on EEGNet as its classical core [2].

Quantum-augmented EEG encoding (QEEGNet)

QEEGNet [2] is the most directly comparable work: it appends a variational quantum encoding layer to EEGNet and reports improved accuracy and noise robustness on the BCI Competition IV-2a motor-imagery benchmark. It is the closest precedent for the quantum projection head proposed in Section 5. Two qualifications are essential and are stated plainly. First, the cross-task and cross-dataset extension [3] found that the gains over classical EEGNet “remain inconsistent,” concluding that hybrid architectures require further optimization to realize a quantum advantage. Second, the original work characterizes itself as forward-looking and does not claim that quantum layers reliably surpass classical methods. Our proposal differs from QEEGNet in placement (we additionally evaluate a VQC as an attention value-projection, not only a terminal encoding layer) and, more substantively, in evaluation discipline: we make matched-capacity classical comparison and ablation a precondition for any reported gain.

Quantum-inspired EEG denoising (RQNN)

The recurrent quantum neural network of Gandhi and colleagues [4,5] filters raw EEG by modeling the signal as a time-varying wave packet evolving under a Schrödinger-equation-inspired update, learning the signal’s statistics in an unsupervised manner and thereby removing non-Gaussian noise without an explicit noise model. This is the precedent for treating denoising as distribution learning rather than supervised regression, and it targets exactly the non-stationary residual that defeats classical filters. It is, however, a single-purpose filter rather than an end-to-end decoding architecture, and its quantum formalism is a mathematical analogy implemented classically rather than a circuit executed on quantum hardware. We draw the denoising objective from it, not the implementation.

Quantum attention (QASA)

The QASA transformer [6] replaces the value projection of a single encoder layer with a parameterized quantum circuit—using on the order of tens of trainable quantum parameters—while keeping all other layers classical, and reports competitive results on time-series forecasting with a deliberately minimal quantum footprint. This minimal-integration philosophy is the precedent for our attention-block insertion and for our broader stance that the right amount of quantum is small. As with the other works, QASA’s reported advantages are task-dependent and, by its authors’ framing, contingent on the maturation of low-noise hardware; we inherit that caution.

Taken together, the literature supports a measured claim: hybrid quantum–classical models for biosignals are a legitimate and active research direction with isolated positive results, and no work has yet

demonstrated a robust, reproducible quantum advantage over strong classical baselines on EEG. Our contribution is positioned squarely inside that gap—as a concrete architecture plus the evaluation needed to determine, rather than assert, whether the direction pays off.

SECTION 5

Proposed Hybrid Architecture

The proposed system is a layered, modular pipeline that ingests raw multi-channel biosignal windows and outputs either cleaned feature representations or task predictions (for example, motor-imagery class or a continuous cognitive-state estimate). The classical components carry the bulk of the computation; the quantum components are small, optional modules inserted at two points where a richer feature map is most plausibly valuable. The design goal throughout is that every quantum element be ablatable—removable in a controlled experiment—so its contribution can be measured rather than assumed.

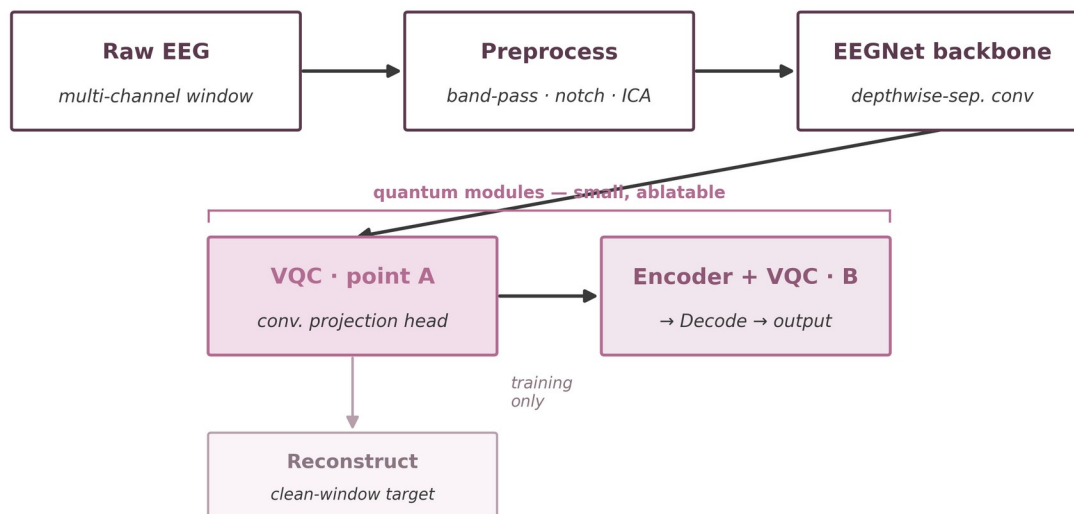


Figure 1. Hybrid pipeline. Classical stages carry the computation; two small variational quantum circuits are inserted where a richer feature map is most plausibly useful. Both are ablatable, and the auxiliary denoising branch is active during training only.

5.1 Classical backbone

Raw input first passes a standard preprocessing front end: band-pass filtering, notch filtering for mains interference, and, where channel count permits, an ICA or regression stage for ocular artifacts [15]. The cleaned window then enters a temporal feature extractor built on the EEGNet pattern [1]—depthwise-separable convolutions that perform learnable spatial and frequency-selective filtering at low parameter cost—optionally followed by a recurrent (LSTM) or transformer encoder when longer temporal dependencies matter. This backbone is deliberately a known-good classical model so that it can also serve, unmodified, as the control baseline in evaluation.

5.2 Quantum insertion point A — the convolutional projection head

After the convolutional feature stage, the network produces a compact feature vector. Insertion point A replaces (or augments, in an ablation) the subsequent dense projection with a variational quantum circuit, following the QEEGNet precedent [2]. The feature vector is dimensionally reduced to match the qubit count, angle-encoded into the circuit, transformed by a shallow parameterized-and-entangling block, and read out as Pauli-Z expectations that form the projected representation passed to the classifier. The hypothesis is that this projection learns a separation of signal and artifact subspaces that a same-sized classical dense layer does not; the validation protocol tests that hypothesis directly against a matched classical projection.

5.3 Quantum insertion point B — a single attention value-projection

When a transformer encoder is used for temporal context, insertion point B replaces the value projection of one attention block with a parameterized quantum circuit, following QASA [6]. Queries and keys remain classical; only the value pathway of a single layer is quantum, and a residual connection preserves the classical signal path. This keeps the quantum parameter count small—on the order of tens of angles—which is the regime least exposed to barren plateaus and shot noise [9,14]. The hypothesis is that a quantum-valued attention head captures non-classical inter-token correlations useful for tracking non-stationary signal structure; again, the ablation that removes it is part of the design.

5.4 Training and the denoising objective

The full graph—classical and quantum modules together—is trained end-to-end by gradient descent, with quantum-gate gradients obtained by the parameter-shift rule on hardware or by automatic differentiation in simulation [10,13]. Two objectives are combined. A supervised task loss (cross-entropy for classification, or a regression loss for state estimation) drives decoding performance. An auxiliary reconstruction or denoising loss—reconstructing a clean target window from a corrupted input, in the spirit of the RQNN distribution-learning objective [4,5]—drives the representation toward noise-robustness, and is most relevant to the dry-electrode and motion regimes identified in Section 2. The relative weighting of the two losses is a hyperparameter selected on validation data, never on the test set.

Algorithm: forward pass (one window)

- 1) Preprocess the raw window (band-pass, notch, optional ICA).
- 2) Apply the EEGNet convolutional stage to obtain spatial-temporal features.
- 3) At insertion point A, angle-encode the reduced feature vector, evaluate the VQC, and read out expectations as the projected features.
- 4) If a transformer is used, run the encoder with insertion point B supplying the quantum value-projection of one attention head.
- 5) Decode to a task output, and (during training only) also decode to a reconstructed window for the auxiliary loss.

Each quantum evaluation in steps 3 and 4 is the average over a fixed shot budget, and that budget is reported as a first-class experimental parameter because it controls the variance of both inference and gradients.

SECTION 6

Software and Model Stack

The stack is built entirely on community-validated, open frameworks so that every component can be independently audited—an explicit countermeasure to the “black-box quantum” concern that surrounds

some proprietary claims in this space. The classical side uses PyTorch or TensorFlow/Keras for model definition and GPU-accelerated training. The quantum side uses, depending on the experiment, PennyLane [13] for hybrid differentiable circuits and its hardware backends, TorchQuantum [20] for circuits embedded directly in the PyTorch graph with automatic differentiation and GPU-accelerated simulation, and IBM Qiskit [18] or Google Cirq [19] for circuit construction and execution on simulators and NISQ hardware. PennyLane and TorchQuantum are the primary development tools because they make the hybrid graph differentiable end-to-end; Qiskit and Cirq provide hardware-aware execution and validation.

Interoperability is handled through ONNX [21], the open neural-network exchange format, which provides a common operator set and portable model representation. After training, the deployable network is exported to ONNX (or TensorFlow Lite) so that it runs identically across the inference engines used at the edge—ARM CPUs, mobile SoCs, FPGAs, and small accelerators. The decision to keep the entire toolchain open is deliberate: it preserves the ability to verify behavior, avoids vendor lock-in, and aligns with procurement standards in regulated and defense settings, where provenance and auditability are requirements rather than preferences.

It is worth being precise about what “export” means for the quantum component, because this is where the deployment model in Section 7 turns. A trained VQC computes a fixed function of its input once its parameters are frozen. In simulation, that function can be evaluated by the same classical runtime as the rest of the network; on hardware, it requires a quantum processor and a shot budget. Which of these the deployed system uses is a design choice with real consequences, examined next.

SECTION 7

Deployment Model and the Distillation Question

The deployment model separates where quantum computation occurs from where the model is used. Two pathways operate over the same model lineage, linked by ONNX so that a model refined centrally translates identically to a field device.

7.1 On-device inference

Wearable and embedded targets—headsets, soldier-worn sensors, habitat or suit-mounted units—have tight power and latency budgets. On these, the system runs a classical network derived from the trained hybrid model, quantized or pruned as needed and executed via ONNX Runtime or TensorFlow Lite on ARM, FPGA, or small-accelerator hardware (for example Jetson-class or microcontroller-plus-accelerator devices). Preprocessing and artifact filtering run first; the network then produces cleaned features or class scores. The design target is end-to-end latency in the tens of milliseconds with continuous online operation; that figure is a requirement to be measured under load on representative hardware, not a measured result, and is reported as such.

7.2 Post-mission and cloud processing

Offline, the same data can be reprocessed on high-performance compute or quantum cloud backends: retraining on larger pooled datasets, fine-tuning circuit parameters, and running quantum simulators or NISQ hardware for feature extraction that is too expensive for the edge. This is also where calibration is refreshed—an alert threshold or a per-user model can be re-estimated from accumulated data and

pushed back to devices. For distributed operations, encrypted telemetry can carry summarized signals to a command center while devices run inference locally, fitting federated or edge-cloud architectures.

7.3 The distillation question, stated honestly

The pitch underlying this work asserts that a quantum advantage can be “captured during training” and then “frozen into classical operators” so that inference needs no quantum processor. This deserves careful scrutiny, because as commonly stated it can be misleading. If a trained VQC is replaced at inference by a fixed classical operator that reproduces its input–output map exactly, then by construction a classical layer is computing that map—and the question becomes why a classical network of equal capacity could not have learned the same map directly during training. There is no inference-time quantum effect in a frozen classical operator; whatever the quantum component contributed must have entered through the training dynamics—the optimization landscape, the inductive bias, or the representation the quantum module encouraged—not through anything quantum that survives in the shipped model.

This is not a fatal objection, but it sharply constrains the honest claim. Two cases must be distinguished. In the genuinely classical-at-inference case, any benefit is a training-time effect and must be demonstrated as such: the hybrid-trained-then-frozen model must outperform a matched classical model trained conventionally, or there is no benefit to speak of. In the quantum-at-inference case, the VQC is actually executed on hardware at run time, the benefit can in principle be a quantum one, but the NISQ constraints of Section 3—shot noise, latency, qubit count—apply in full and are incompatible with the tens-of-milliseconds edge budget today. The architecture supports both pathways, but this white paper claims neither as established. The validation protocol in Section 8 is built precisely to determine which case, if either, yields a real and reproducible gain.

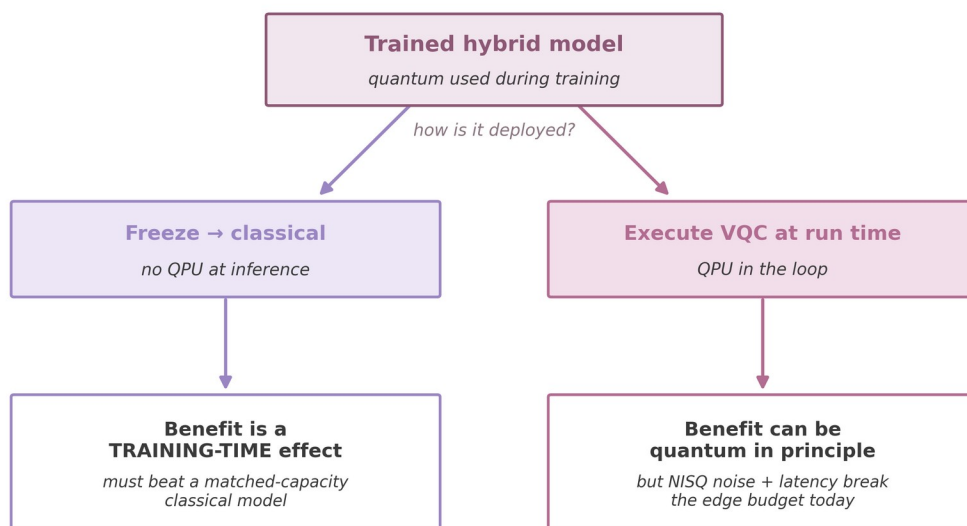


Figure 2. The distillation question (§7.3). A frozen classical operator carries no inference-time quantum effect, so any benefit must enter through training; the validation protocol is built to test which branch, if either, yields a real and reproducible gain.

7.4 Defense and mission deployment considerations

For defense and aerospace use, the deployment model assumes adversarial and degraded conditions by default: intermittent or denied connectivity, the need for fully offline operation, strict audit logging and data provenance, and well-defined security boundaries between edge and command-center tiers. The open, exportable toolchain supports these requirements—models and their training lineage can be inspected and version-controlled—but a complete threat model, including model-integrity protection, telemetry confidentiality, and safe-default behavior under sensor failure, is a deliverable of system engineering beyond this architectural white paper and is flagged here as required future work rather than a solved problem.

SECTION 8

Validation Methodology

This section is the core of the white paper's claim to be an engineering document rather than a brochure. It specifies how the hypotheses of Section 5 would be tested. The governing principle is that an apparent improvement is not an improvement until baseline fairness, contribution isolation, and statistical discipline have been established. No result described below has been performed; this is a protocol, and its value is that it is falsifiable.

8.1 Datasets

Evaluation begins on public motor-imagery benchmarks, principally BCI Competition IV-2a (the dataset on which QEEGNet reported its results [2], enabling direct comparison), supplemented by additional public EEG datasets spanning P300 and error-related-negativity paradigms to test cross-paradigm generalization in the spirit of EEGNet's original evaluation [1]. Because the operational target is dry-electrode field use, a dedicated noise-robustness track augments clean recordings with controlled synthetic and recorded artifacts—EMG bursts, simulated impedance drift, and motion—so that robustness can be measured as a function of degradation rather than only on pristine data. Any internally collected dataset would be added only with documented provenance and ethics review.

8.2 Baselines and fairness

The decisive comparison is against a classical model of matched capacity. For every hybrid configuration, a purely classical counterpart with a comparable parameter count and an equivalent dense or attention projection is trained under equal compute and equal hyperparameter-tuning effort on the same data and the same evaluation protocol. This is the comparison that distinguishes a real effect from the appearance of one, and it is non-negotiable. Additional baselines include the unmodified EEGNet [1], a deep convolutional model [16], and the classical denoising pipeline alone (ICA plus filtering) [15].

Comparison	What it isolates	Why it matters
Hybrid vs. matched-capacity classical	Whether the quantum module adds anything beyond equal parameters	The central test; guards against capacity confounds
Hybrid-frozen vs. classical (conventional)	Whether a training-time quantum effect survives export	Directly tests the distillation claim (§7.3)
Quantum-at-inference vs. frozen-	Whether run-time quantum	Separates the two deployment cases

Comparison	What it isolates	Why it matters
classical	execution helps	
With vs. without auxiliary denoising loss	Contribution of the denoising objective	Attributes gains to the right mechanism

8.3 Ablations

Each architectural element is removed or varied in isolation: insertion point A on or off; insertion point B on or off; qubit count and circuit depth swept across a small range; data-encoding scheme varied; and shot budget varied to expose the variance that finite sampling injects into both predictions and gradients. The purpose of the sweep over qubit count and depth is not only tuning but diagnosis—barren-plateau behavior [9] should appear as deteriorating trainability as the circuit grows, and observing it (or its absence) is itself a reported finding.

8.4 Metrics and statistics

Decoding performance is reported as classification accuracy and macro-F1 (with F1 emphasized under class imbalance), and where the application warrants it, as calibration of predicted confidence. Fidelity is reported through an SNR proxy and reconstruction error on the denoising task. Every headline number is computed over at least five random seeds and reported with a measure of uncertainty—a confidence interval, standard deviation, or interquartile range—never as a single run. Comparisons that assert superiority are accompanied by a named significance test (for example a paired test across seeds, with correction for multiple comparisons) and by an effect size, because a statistically significant but negligibly small gain does not justify the added complexity and the deployment cost of a quantum training pipeline. Selection-bias hazards—best-of-N reporting, seed cherry-picking, and any tuning on the test set—are excluded by protocol.

8.5 Success criteria and negative results

Success is defined before experiments are run: the hybrid model must show a statistically significant improvement over the matched-capacity classical baseline, of a pre-specified minimum effect size, that is consistent across more than one dataset and that persists under the dry-electrode noise track. A gain that appears on a single dataset, vanishes under matched capacity, or falls below the minimum effect size is recorded as a negative or null result and reported as such. Publishing the conditions under which the quantum component does not help is part of the methodology, both because it is honest and because it is the information a partner or program manager actually needs for a go/no-go decision.

SECTION 9

Strategic and Mission Relevance

The fidelity problem this work targets is the same one named by several public defense and aerospace programs. The alignment described below is with the stated goals of those programs; it is not a claim of participation in, funding from, or endorsement by any agency.

The clearest articulation is DARPA’s Next-Generation Nonsurgical Neurotechnology (N3) program, which seeks high-performance, portable, non-surgical neural interfaces for able-bodied service members and states explicitly that existing non-invasive modalities such as EEG “do not offer the precision, signal resolution, and portability required for advanced applications” in real-world settings [7]. A signal-processing layer that improves effective SNR and robustness addresses precisely the limitation N3 identifies, and does so in the non-invasive regime the program prioritizes. Related efforts toward reliable, long-duration BCI systems share the same dependence on stable interpretation of noisy EEG and EMG as electrodes shift and conditions change—the dry-electrode robustness emphasized throughout this paper.

In aerospace, EEG is used to monitor cognition, fatigue, and neurological status, and these uses are especially exposed to the electromagnetic and motion artifacts of operational environments; more robust denoising directly serves reliability there. Neuroadaptive training programs that infer attention or fatigue from EEG benefit in the same way: better fidelity yields more accurate feedback. Across these domains, the core technology—robust real-time interpretation of non-invasive biosignals—is dual-use, with civilian clinical-monitoring and consumer-neurotechnology applications that share the same technical requirements. The deliberately open stack (PyTorch, ONNX, and audited quantum frameworks) is what makes transition across these settings plausible without lock-in.

Two cautions keep this framing honest. First, programmatic alignment is not technical validation; the experiments of Section 8 remain the gate. Second, the underlying technology is signal processing for biosignal quality, and the appropriate ethical, legal, and regulatory considerations that accompany any neurotechnology—several of which the relevant programs themselves foreground—apply to its deployment and are outside the scope of this architectural document.

SECTION 10

Competitive Landscape

The landscape can be read along two axes that rarely intersect. On one axis sit established EEG hardware and software companies—for example g.tec, Emotiv, and OpenBCI—whose strength is sensing hardware and classical digital signal processing and machine learning, and who do not, to our knowledge, employ quantum methods. On the other sit quantum-computing firms and toolmakers—IBM, Google, Rigetti, and others—whose strength is hardware and general-purpose quantum software (Qiskit, Cirq, PennyLane), targeted at domains such as chemistry and optimization rather than neurotechnology. Academic work that bridges the two—QEEGNet [2,3] and the RQNN denoising line [4,5]—remains at the research stage.

The honest reading of this gap is twofold. The integration itself—quantum machine learning applied specifically to neurophysiological signal fidelity, in a deployable hybrid stack—is genuinely uncrowded, and we are not aware of a commercial product that offers it. That is the opportunity. But the same gap reflects that the value of the approach is not yet established even in the literature, which is the risk. The differentiator this program can credibly claim today is therefore not a proven performance edge; it is a specialist focus, an auditable open-stack implementation, and—the substance of Section 8—a validation discipline designed to determine whether the edge exists. Claiming more than that would repeat the overclaiming the field is already prone to.

Segment	Representative players	Strength	Gap relative to this work
Classical EEG	g.tec, Emotiv, OpenBCI	Sensors, classical DSP/ML,	No quantum feature

Segment	Representative players	Strength	Gap relative to this work
hardware/software		mature deployment	maps; noise-floor ceiling
Quantum platforms/tooling	IBM, Google, Rigetti	Hardware, general QML libraries	No neurotech specialization
Academic quantum-EEG	QEEGNet, RQNN groups	First demonstrations, peer-reviewed	Research-stage; gains inconsistent; not productized
This proposal	ethereal computing	Specialist hybrid stack; validation-first	Advantage not yet empirically demonstrated

SECTION 11

Limitations and Open Questions

Stating limitations plainly is part of the engineering claim, and the most important ones are not peripheral—they are central to whether the approach works at all.

No demonstrated quantum advantage. The foundational caveat is that no result, here or in the cited literature, establishes a robust, reproducible quantum advantage over strong classical baselines for EEG. QEEGNet’s cross-dataset gains are explicitly inconsistent [3]. Everything in this document is conditional on the Section 8 experiments returning a positive, matched-capacity result, and they may not.

The distillation tension. As argued in Section 7.3, a quantum benefit that is frozen into a classical operator at inference must be a training-time effect, and a classical model of equal capacity may achieve the same map directly. If the matched-capacity baseline matches the hybrid, the quantum pipeline’s added complexity is not justified. This is the single most likely way the program fails its own test.

Trainability and NISQ constraints. Barren plateaus can render quantum layers untrainable as they scale [9]; shot noise injects variance into inference and gradients; and current NISQ hardware [14] offers limited qubit counts and coherence. The architecture mitigates these by keeping circuits small and entry points few, but mitigation is not elimination, and trainability is an explicit experimental risk.

Latency on real hardware. The tens-of-milliseconds edge target is achievable with a classical exported model but is presently incompatible with executing VQCs on quantum hardware at inference. The quantum-at-inference pathway is therefore a research configuration, not a near-term field configuration, and should be described that way.

Data scale and distribution shift. EEG datasets are small relative to the data appetites of expressive models, and dry-electrode field noise differs from the laboratory recordings on which public benchmarks are built. Strong benchmark results need not transfer to deployment, which is why the noise-robustness track and documented field data are part of the plan rather than afterthoughts.

Scope of this document. This is an architecture and methodology white paper. It does not include a completed security/threat model, a regulatory pathway, or empirical results, each of which is required before operational claims can be made and each of which is identified here as future work.

SECTION 12

Conclusion

Non-invasive biosignal acquisition is limited less by decoding algorithms than by signal fidelity, and that bottleneck is what stands between today's EEG-based interfaces and the precision that clinical, defense, and aerospace applications require. This white paper has proposed a concrete way to attack the bottleneck: a hybrid classical-quantum architecture that keeps a proven classical backbone and inserts small, ablatable variational quantum circuits at a convolutional projection and at a single attention value-projection, trained end-to-end with a denoising objective and deployed through an open, exportable toolchain.

What distinguishes this document from the showcase material it is built on is its refusal to assert what it has not shown. The contributions are an insertion-point architecture grounded in QEEGNet and QASA, a deployment lifecycle whose central assumption—that quantum benefit can be captured at training and frozen for classical inference—is examined rather than asserted, and a falsifiable validation protocol whose decisive test is whether the hybrid beats a classical model of equal capacity under equal compute. The peer-reviewed evidence today is promising but inconsistent, and the honest position is that the magnitude of any quantum advantage for EEG remains open. The purpose of the architecture and the methodology presented here is to close that question with evidence—and to report the answer whichever way it falls.

References

- [1] V. J. Lawhern, A. J. Solon, N. R. Waytowich, S. M. Gordon, C. P. Hung, and B. J. Lance, "EEGNet: a compact convolutional neural network for EEG-based brain-computer interfaces," *Journal of Neural Engineering*, vol. 15, no. 5, art. 056013, 2018. DOI: 10.1088/1741-2552/aace8c.
- [2] C.-S. Chen et al., "QEEGNet: Quantum Machine Learning for Enhanced Electroencephalography Encoding," in *Proc. 2024 IEEE Workshop on Signal Processing Systems (SiPS)*, pp. 153–158, 2024. arXiv:2407.19214.
- [3] C.-S. Chen et al., "Exploring the Potential of QEEGNet for Cross-Task and Cross-Dataset Electroencephalography Encoding with Quantum Machine Learning," *Journal of Signal Processing Systems*, 2025. arXiv:2503.00080. DOI: 10.1007/s11265-025-01979-2.
- [4] V. Gandhi, V. Arora, L. Behera, G. Prasad, D. Coyle, and T. M. McGinnity, "EEG denoising with a recurrent quantum neural network for a brain-computer interface," in *Proc. International Joint Conference on Neural Networks (IJCNN)*, pp. 1583–1590, 2011. DOI: 10.1109/IJCNN.2011.6033413.
- [5] V. Gandhi, G. Prasad, D. Coyle, L. Behera, and T. M. McGinnity, "Quantum Neural Network-Based EEG Filtering for a Brain-Computer Interface," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 25, no. 2, 2014. PMID: 24807028.
- [6] C.-S. Chen and E.-J. Kuo, "Quantum Adaptive Self-Attention for Quantum Transformer Models," arXiv:2504.05336, 2025.
- [7] Defense Advanced Research Projects Agency, "Next-Generation Nonsurgical Neurotechnology (N3)," DARPA program page. Available: <https://www.darpa.mil/research/programs/next-generation-nonsurgical-neurotechnology>

- [8] 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,” *Nature Reviews Physics*, vol. 3, pp. 625–644, 2021.
- [9] J. R. McClean, S. Boixo, V. N. Smelyanskiy, R. Babbush, and H. Neven, “Barren plateaus in quantum neural network training landscapes,” *Nature Communications*, vol. 9, art. 4812, 2018.
- [10] K. Mitarai, M. Negoro, M. Kitagawa, and K. Fujii, “Quantum circuit learning,” *Physical Review A*, vol. 98, art. 032309, 2018.
- [11] 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*, vol. 567, pp. 209–212, 2019.
- [12] 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*, vol. 103, art. 032430, 2021.
- [13] V. Bergholm et al., “PennyLane: Automatic differentiation of hybrid quantum-classical computations,” *arXiv:1811.04968*, 2018.
- [14] J. Preskill, “Quantum Computing in the NISQ era and beyond,” *Quantum*, vol. 2, art. 79, 2018.
- [15] T.-P. Jung, S. Makeig, C. Humphries, T.-W. Lee, M. J. McKeown, V. Iragui, and T. J. Sejnowski, “Removing electroencephalographic artifacts by blind source separation,” *Psychophysiology*, vol. 37, no. 2, pp. 163–178, 2000.
- [16] R. T. Schirrmester et al., “Deep learning with convolutional neural networks for EEG decoding and visualization,” *Human Brain Mapping*, vol. 38, no. 11, pp. 5391–5420, 2017.
- [17] Y. Roy, H. Banville, I. Albuquerque, A. Gramfort, T. H. Falk, and J. Faubert, “Deep learning-based electroencephalography analysis: a systematic review,” *Journal of Neural Engineering*, vol. 16, no. 5, art. 051001, 2019.
- [18] Qiskit contributors, “Qiskit: An Open-Source Framework for Quantum Computing,” IBM Quantum. Available: <https://www.ibm.com/quantum/qiskit>
- [19] Cirq Developers, “Cirq: A Python framework for creating, editing, and invoking NISQ circuits,” Google Quantum AI. Available: <https://quantumai.google/cirq>
- [20] H. Wang, Y. Ding, J. Gu, Z. Li, Y. Lin, D. Z. Pan, F. T. Chong, and S. Han, “QuantumNAS: Noise-Adaptive Search for Robust Quantum Circuits” (TorchQuantum), in *Proc. IEEE International Symposium on High-Performance Computer Architecture (HPCA)*, 2022.
- [21] Open Neural Network Exchange (ONNX), Linux Foundation. Available: <https://onnx.ai>

Appendix A — Glossary of Terms

BCI. Brain-computer interface; a system that converts measured neural activity into a control or monitoring signal.

EEG / EMG. Electroencephalography / electromyography; scalp-recorded electrical activity of the brain / of muscle.

SNR. Signal-to-noise ratio; the ratio of signal power to noise power, the quantity this work aims to improve.

VQC / PQC. Variational (parameterized) quantum circuit; a quantum circuit with trainable gate parameters, used as a learnable module.

Feature map / encoding. The mapping that loads classical data into a quantum state; it largely determines a VQC's expressivity.

Parameter-shift rule. A method for computing exact gradients of quantum-circuit outputs with respect to gate parameters, enabling gradient-descent training.

Shots. The number of repeated circuit executions averaged to estimate a measurement expectation; finite shots cause sampling variance.

Barren plateau. A training pathology in which gradients vanish exponentially as a circuit grows, preventing optimization.

NISQ. Noisy intermediate-scale quantum; the current era of quantum hardware with limited, error-prone qubits.

ONNX. Open Neural Network Exchange; a portable model format enabling identical execution across inference engines.

Ablation. An experiment that removes or varies one component in isolation to measure its specific contribution.

Matched-capacity baseline. A classical model with comparable parameter count and training budget, used to test whether a quantum module adds value beyond capacity.

Appendix B — Proposed Reproducibility Checklist

The following items would be reported in full for any empirical result produced under the Section 8 protocol. The list adapts standard machine-learning reproducibility checklists to the hybrid quantum-classical setting. “Available on request” is not treated as reported.

Item	Reporting requirement
Code release	Public repository with license and, for blind review, anonymized mirror
Data	Dataset identifiers, splits, availability statement, and ethics review status for any collected data
Preprocessing	Exact filtering, notch, and ICA settings shipped with code

Item	Reporting requirement
Hyperparameters	Full table of values, ranges searched, and the selection method (on validation, never test)
Quantum configuration	Qubit count, circuit depth, encoding scheme, entangling pattern, and shot budget
Seeds and uncertainty	At least five seeds; confidence interval, standard deviation, or IQR for every headline number
Statistics	Named significance test with multiple-comparison correction, and reported effect size
Hardware / simulator	CPU/GPU and any QPU or simulator used, including versions
Compute budget	Training and inference wall-clock and, where relevant, shot counts per experiment
Negative results	Configurations in which the quantum component did not help, reported explicitly

IP and Disclosure Note

This white paper describes subject matter that may overlap with patentable inventions. Public release of a technical disclosure can constitute a printed-publication disclosure event that affects patent rights, and in many jurisdictions disclosure prior to filing forecloses or weakens later applications. Consistent with standard practice and with ethereal computing's own portfolio workflow, any provisional or non-provisional patent application covering this subject matter should be on file before this document is published or otherwise distributed externally. This note is procedural guidance, not legal advice; counsel should confirm the disclosure posture before release.