

Which settings of the trainable quantum model can be trained

Most of them — if you measure locally. The product offers five ansätze, four entanglement patterns, three measurements, four depths, warm starts and readout heads. We screened **every combination** before training, trained 80 of them against a full classical referee, and turned the answer into a **one-minute pre-training screen** and guardrails. Design tagged before the first run.

7,488

pre-training screens

gradient and output, three instruments, 4–16 qubits

80

refereed trainings

ten classical variants, whole-CI win rule

86%

trainable at 16 qubits

with a local measurement and the product's default start

Swipe →

A circuit with adjustable angles, a gradient, and a referee.

A variational quantum classifier encodes the data, turns it through trainable layers, measures a few numbers and lets a small classical head decide. Training moves the angles by **gradient descent** — and the theory warns that some settings make the gradient vanish before training starts, so that hours of training move nothing: the **barren plateau**.

The study measured the gradient and the output spread **before training** on every setting the product offers (seconds each), then trained 80 pre-registered settings for real, each scored against **ten tuned classical models** with the whole 95% confidence interval above parity.

Screen in a minute — train for hours only where it can pay.

The product's default region is trainable — and wins where it should.

Trainable at 16 qubits: **86%** of local-measurement settings

Under the product's default initialization the gradient stays flat with depth. The default configuration beat the classical referee on its native labels at 4, 8 and **16 qubits** (0.93 vs 0.86) — the largest register this series has ever won — and outscored every other trainable setting there in 14 of 14 comparisons.

The protection is the initialization — our warm-start bet **lost**

In the theory's regime (random angles over the full circle) local gradients fall about 100-fold at 16 qubits; the product's small-angle start keeps them where they are. We bet an explicit warm start would rescue 90% of the dead settings: it rescued **29%**, because the product had already built the protection in.

One measurement is dead beyond a few qubits, and nothing rescues it.

With the **global parity** observable the initial gradient halves with every qubit (about 4,000-fold from 4 to 16) under every initialization. The warm start lifts 20 of 336 such settings, 2 of 8 rescue trainings improve, and **0 of 22** refereed runs with it beat the referee.

Local measurements do not have this problem. The product now limits the global measurement to 11 qubits, with the reason printed — the same discipline as the dead encoding the kernel atlas found.

The textbook barren plateau is real — and, in this product, a measurement choice.

Structure matters less than the advice says. The screen works.

Entanglement, depth, re-uploading, readout head: **second order**

The entanglement pattern does not order the gradients (17 of 72 cells monotone); depth does not hurt local measurements; data re-uploading **lowers** the output spread; the quadratic head shifts the gradient by 22%. Under the product's start these are design choices, not risks.

9 wins of 80 — none from a concentrated screen

0 of 18 settings that screened as concentrated won or reached parity with the kernel method; 11 of 42 trainable ones did. Our registered yes/no cutoff predicts poorly (AUC 0.68); the **graded gradient reading** (0.88) and a **25-iteration probe** (0.82) predict well. Both now print in every run.

Measure locally, trust the default start, screen before training.

Measure locally

the global observable is dead at scale and nothing rescues it — the product now rejects it from 12 qubits.

Trust the default start

the product's initialization is the real protection; the warm start is a second attempt, not a guarantee — the advice now says so.

Screen before training

58 seconds per 16-qubit screen against about two hours of training — a concentrated screen is a stop sign, a trainable one a permit to try.

Eight of eleven registered hypotheses fell — mostly community advice, not the product.

The full study, unedited.

Ten pages: the gradient maps, eleven pre-registered hypotheses with their outcomes, the validation layer, the product changes, where the potential lies, and the registration — design tagged **before the first run**.

The study — 10 pages

qubit-lab.ch/evidence

Study PDF, published benchmark records and the QML Rapid Quantum Prototype are on qubit-lab.ch/evidence. Evidence, not promises.