

Trabajo fin de máster

Caracterización en hardware de la Estimación Iterativa de Amplitud Cuántica en procesadores IQM, con aplicación a la optimización de carteras bajo riesgo CVaR



Ignacio López Leis

Escuela Politécnica Superior
Universidad Autónoma de Madrid
C/ Francisco Tomás y Valiente, 11

UNIVERSIDAD AUTÓNOMA DE MADRID
ESCUELA POLITÉCNICA SUPERIOR



Máster en Máster en Computación Cuántica

TRABAJO FIN DE MÁSTER

**Caracterización en hardware de la Estimación
Iterativa de Amplitud Cuántica en procesadores
IQM, con aplicación a la optimización de carteras
bajo riesgo CVaR**

**Un pipeline híbrido cuántico-clásico: de la estimación de
amplitud en hardware NISQ a la validación financiera bajo
Basilea IV**

Autor: Ignacio López Leis
Tutor: Luis de Pedro Sánchez

agosto 2026

Todos los derechos reservados.

Queda prohibida, salvo excepción prevista en la Ley, cualquier forma de reproducción, distribución comunicación pública y transformación de esta obra sin contar con la autorización de los titulares de la propiedad intelectual.

La infracción de los derechos mencionados puede ser constitutiva de delito contra la propiedad intelectual (*arts. 270 y sgts. del Código Penal*).

DERECHOS RESERVADOS

© TODO: fecha de defensa sin confirmar por UNIVERSIDAD AUTÓNOMA DE MADRID
Francisco Tomás y Valiente, nº 1
Madrid, 28049
Spain

Ignacio López Leis

Caracterización en hardware de la Estimación Iterativa de Amplitud Cuántica en procesadores IQM, con aplicación a la optimización de carteras bajo riesgo CVaR

Ignacio López Leis

RESUMEN

Un banco que gestiona una cartera de inversión necesita responder, de forma rutinaria y bajo obligación regulatoria (Basilea IV), a una pregunta concreta: en los peores escenarios de mercado de una franja de probabilidad fijada, ¿cuánto se pierde en promedio? Esa magnitud es el valor en riesgo condicional (CVaR), y calcularla con precisión sobre una cartera de cientos de activos mediante el método clásico estándar –simulación de Monte Carlo– es caro: el coste crece como $O(\varepsilon^{-2})$ con la precisión ε exigida, así que duplicar la precisión cuadruplica el trabajo. La estimación cuántica de amplitud, y en particular su variante iterativa (IQAE), promete una escala $O(\varepsilon^{-1})$ para el mismo problema. Este TFM no se limita a asumir esa ganancia asintótica: construye un pipeline híbrido cuántico-clásico completo que la instancia sobre optimización de cartera bajo CVaR, y lo usa para *caracterizar experimentalmente*, en hardware cuántico ruidoso real (procesadores superconductores IQM Emerald y Garnet), en qué condiciones IQAE puede ejecutarse de forma fiable y hasta qué precisión –sin reclamar que la ventaja computacional cuántica ya esté disponible a las escalas probadas ($N \leq 100$ activos)–.

El pipeline consta de tres etapas. La Etapa 1 selecciona el universo de activos mediante un problema QUBO formulado para hardware de *annealing* D-Wave –resuelto, en las selecciones reportadas, con *annealing* simulado clásico–, favoreciendo activos estructuralmente periféricos en una red de correlaciones filtrada (TMFG). La Etapa 2 estima el subgradiente de CVaR respecto a los pesos mediante un circuito IQAE de *anchura fija* –cuatro qubits de distribución más un ancilla–, independiente de N , apoyándose en la formulación de Rockafellar–Uryasev, que reduce ese gradiente a una esperanza condicionada en la cola de pérdidas. La Etapa 3 actualiza los pesos con Adam sobre un vector de parámetros libres del que los pesos se recuperan por softmax, de modo que las restricciones se cumplen por construcción, sin proyección a posteriori.

El pipeline completo se valida de extremo a extremo **en simulación** sobre 374 constituyentes del S&P 500 y carteras de hasta 100 activos, a lo largo de 46 experimentos sistemáticos. Ese conjunto permite además derivar y verificar una desigualdad de *brecha de ranking* –consecuencia del error *multiplicativo* (no aditivo) de IQAE– bajo la cual el orden relativo de riesgo entre activos, que determina la dirección de actualización de Adam, se preserva pese al ruido: la similitud coseno entre subgradientes ruidosos e ideales supera 0,9998 en el punto de partida equiponderado de los 46 experimentos. Es una propiedad comprobable de cada universo y no un producto de la selección de la Etapa 1 –se cumple con el mismo margen ($\approx 1,83 \times \alpha \varepsilon = 0,005$) para la cartera periférica y para la central de comparación–, y en los pesos convergidos se viola necesariamente, porque los CVaR marginales se igualan en el óptimo: es una condición suficiente de inicio de trayectoria, no una garantía de diseño.

El hardware IQM real, por separado, caracteriza la dinámica de rondas de Grover de forma aislada, no el bucle de optimización ni el circuito de cinco qubits del pipeline: las ejecuciones usan una *sonda*

mínima de dos qubits (un giro R_y más una CNOT) que reproduce la misma amplitud de cola codificada. Sobre ella se establece la fórmula de profundidad $\text{depth}(m) = 4 + 6m$ y se identifica un umbral de *aliasing* –predicho analíticamente desde la geometría de la cola $((2m+1)\theta \approx \pi/2$ en $m \approx 3$ para $\theta \approx 0,2265$) y confirmado en ambos procesadores– que, por depender solo de la amplitud codificada, se transfiere sin cambios al circuito real, nunca ejecutado en hardware y que transpila *offline* a 1,476 de profundidad ya en $m = 0$. La precisión fiable de la sonda queda acotada a $\varepsilon_{\text{hw}} \approx 0,20$: dos órdenes de magnitud por encima del objetivo $\varepsilon = 0,005$ de simulación, y una cota inferior –no una medida– del suelo del circuito real. Se compara también, en simulador y en Emerald, la extrapolación a ruido cero (ZNE): el *folding* global introduce un sobrecoste de profundidad de $4,5\times$, y por debajo del umbral de *aliasing* ninguna variante mejora sobre la estimación cruda.

Aplicado a la optimización financiera, el pipeline logra, en la parte *in-sample*, una reducción de CVaR del 25,6 % frente a pesos iguales para la cartera PA, comparable a Markowitz (25,1 %) y a 1,4 puntos porcentuales de SLSQP (27,0 %). En la ventana *fuera de muestra* (501 días, 2024–2025), el CVaR obtenido es el más bajo de los métodos comparados, pero su respaldo estadístico es débil: el $p \approx 0,12$ de *bootstrap* a una cola (cartera PA frente a SLSQP) no es significativo ni siquiera a una cola, y la tasa de victorias del 54,3 % (25/46) a través de todos los experimentos tampoco lo es (binomial $p = 0,329$ a una cola). Frente a la batería de Basilea IV, el test unilateral de especificación de Acerbi–Székely (estadístico Z_2 , convención de signo propia de sus autores) no se rechaza en los cuatro niveles de confianza probados –el modelo es conservador, sobreestimando las pérdidas de cola en vez de subestimarlas–; los componentes de Kupiec y Christoffersen fallan de forma conservadora en la mayoría de los niveles, aunque ambos pasan en $\alpha = 10\%$ (Christoffersen también en $\alpha = 1\%$), y no se reclama superioridad regulatoria sobre esta base.

En conjunto, este TFM entrega una caracterización experimental de *dónde* funciona IQAE en hardware NISQ real, *dónde* deja de funcionar y por qué, usando un problema financiero regulado como instrumento de validación externa, no como demostración de ventaja computacional cuántica.

PALABRAS CLAVE

Estimación iterativa de amplitud cuántica, procesadores superconductores IQM, caracterización de hardware NISQ, análisis de profundidad de circuito, extrapolación a ruido cero, valor en riesgo condicional, algoritmos híbridos cuántico-clásicos, QUBO

ABSTRACT

A bank managing an investment portfolio must routinely answer, under a real regulatory obligation (Basel IV), a concrete question: within the worst market scenarios inside a fixed probability band, how much is lost on average? That quantity is the conditional value-at-risk (CVaR), and computing it accurately over a portfolio of hundreds of assets with the classical standard method –Monte Carlo simulation– is expensive: the estimation error shrinks as $O(\varepsilon^{-2})$ in the target precision ε , so doubling the required precision quadruples the work. Quantum amplitude estimation, and in particular its iterative variant (IQAE), promises an $O(\varepsilon^{-1})$ scaling for the same problem. This thesis does not merely assume that asymptotic gain: it builds a complete hybrid quantum-classical pipeline that instantiates it on CVaR portfolio optimisation, and uses the result to *experimentally characterise*, on real noisy quantum hardware (IQM Emerald and Garnet superconducting processors), under which concrete conditions IQAE can run reliably and up to what precision –without ever claiming that a quantum computational advantage is already available at the scales tested ($N \leq 100$ assets)–.

The pipeline has three stages. Stage 1 selects the asset universe via a QUBO problem formulated for D-Wave quantum-annealing hardware –the reported selection runs used classical simulated annealing–, favouring structurally peripheral assets on a correlation network filtered down to a Triangulated Maximally Filtered Graph (TMFG). Stage 2 estimates the subgradient of CVaR with respect to the portfolio weights with an IQAE circuit of *fixed width* –four distribution qubits plus one ancilla–, independent of N , relying on the Rockafellar–Uryasev formulation, which reduces that gradient to a conditional expectation in the loss tail. Stage 3 updates the weights via Adam acting not on the weights directly but on a vector of unconstrained parameters from which they are recovered through a softmax reparametrisation, so that the portfolio constraints hold by construction with no a-posteriori projection.

The full pipeline is validated end to end **in simulation** on 374 S&P 500 constituents and portfolios of up to 100 assets, across 46 systematic experiments. That same experiment set also lets us derive and verify a *rank-gap* inequality: a consequence of the *multiplicative* (not additive) error structure of IQAE under which the relative risk ordering between assets –the quantity that determines Adam’s update direction– is preserved despite noise, with a cosine similarity between noisy and ideal subgradients above 0,9998 at the equal-weight starting point of each of the 46 experiments. That condition is a checkable property of each concrete universe, not a product of the Stage-1 selection: computed on the real marginal-CVaR vectors, it holds at that starting point by the same margin ($\approx 1,83\times$ relative to the $\varepsilon = 0,005$ target) for the QUBO-selected peripheral portfolio and for the central comparison portfolio alike, so the verified benefit of the QUBO selection is the structural diversification of the universe and of its loss tail, not the safety margin against noise; and it is necessarily violated at the converged weights, where the marginal CVaRs of the held assets equalise, so it is a start-of-trajectory sufficient condition

rather than a design guarantee. Real IQM hardware, separately, characterises IQAE’s Grover-round dynamics in isolation, not the full optimisation loop nor the pipeline’s own five-qubit circuit: the hardware runs instead execute a minimal two-qubit *probe* (R_y rotation plus one CNOT) that carries the same encoded tail amplitude. On that probe, this thesis establishes the circuit-depth formula $\text{depth}(m) = 4 + 6m$ and identifies, predicting it analytically from the geometry of the distribution tail ($(2m+1)\theta \approx \pi/2$ at $m \approx 3$ for $\theta \approx 0,2265$) and confirming it empirically on both processors, an *aliasing* threshold that—depending only on the encoded amplitude—transfers unchanged to the pipeline’s own circuit, never executed on hardware and which transpiles offline to a native depth of 1,476 already at $m = 0$ (against 4 for the probe). The probe’s reliable precision on current hardware is bounded to $\varepsilon_{\text{hw}} \approx 0,20$ —two orders of magnitude above the $\varepsilon = 0,005$ target used in simulation, and a lower bound, not a measurement, of the pipeline circuit’s own precision floor—. Zero-noise extrapolation (ZNE) is also compared, in simulator and on Emerald hardware: global circuit folding introduces a $4,5\times$ depth overhead, and below the aliasing threshold no ZNE variant improves on raw IQAE estimation.

Applied to portfolio optimisation, the pipeline achieves, in-sample, a 25,6 % CVaR reduction relative to equal weighting for portfolio PA, comparable to Markowitz (25,1 %) and within 1,4 percentage points of SLSQP (27,0 %). Out of sample (501 days, 2024–2025), the resulting CVaR is the lowest among all methods compared, but the statistical support for that improvement is weak, and is reported with the same honesty as the rest of this work: the one-sided bootstrap $p \approx 0,12$ (portfolio PA versus SLSQP) is not statistically significant even one-sided, and the 54,3 % (25/46) win rate across all experiments is not statistically significant either (one-sided binomial $p = 0,329$). Against the Basel IV test battery, the one-sided Acerbi–Székely specification test (statistic Z_2 , in Acerbi–Székely’s own sign convention) is not rejected at any of the four confidence levels tested—the model is conservative, over-predicting tail losses rather than under-predicting them—; the Kupiec and Christoffersen components fail conservatively at most levels, though both pass at $\alpha = 10\%$ (Christoffersen also at $\alpha = 1\%$), and no regulatory superiority is claimed on this basis.

Taken together, this thesis delivers an experimental characterisation of *where* IQAE works on real NISQ hardware, where it stops working, and why, using a regulated financial problem as an external validation instrument, not as a demonstration of quantum computational advantage.

KEYWORDS

Iterative quantum amplitude estimation, IQM superconducting processors, NISQ hardware characterisation, circuit depth analysis, zero-noise extrapolation, conditional value-at-risk, hybrid quantum-classical algorithms, QUBO

ÍNDICE

1	Introducción	1
1.1	Motivación	1
1.2	Glosario de siglas	2
1.3	Tres preguntas abiertas en hardware IQM	2
1.4	Vehículo de validación: optimización de cartera bajo riesgo CVaR	3
1.5	Contribuciones	4
1.6	Relación con el trabajo previo: cómo situar esta tesis	5
1.7	Alcance y advertencias	7
2	Estado del arte	9
2.1	De la estimación de amplitud clásica a la cuántica	9
2.2	Estimación Iterativa de Amplitud	10
2.3	Valor en riesgo condicional y la formulación de Rockafellar–Uryasev	11
2.4	Basilea IV, VaR y CVaR regulatorio desde cero	14
2.5	QUBO y métodos de penalización	15
2.6	Selección de universo: el árbol de expansión mínima de Mantegna y la perifericidad de Pozzi et al.	16
3	Metodología	19
3.1	El pipeline híbrido de tres etapas	19
3.2	Etapas 1: selección de universo vía QUBO	20
3.3	Etapas 2: estimación del subgradiente de CVaR vía IQAE	23
3.4	Etapas 3: optimización con Adam	26
3.5	Diseño experimental: qué varía entre las 46 configuraciones	29
3.6	Justificación de hiperparámetros	29
4	Resultados experimentales	31
4.1	Diseño experimental: la batería de 46 experimentos	31
4.2	Geometría del QUBO y calidad del universo seleccionado	33
4.3	Robustez al ruido y estabilidad del gradiente	35
4.4	Caracterización en hardware IQM: Emerald y Garnet	41
4.5	Síntesis del capítulo	50
5	Aplicación financiera	51
5.1	Backtesting regulatorio desde cero	51

5.2 Validación in-sample: ¿optimiza bien el pipeline cuántico?	53
5.3 Escalabilidad a anchura de circuito fija	55
5.4 Validación fuera de muestra	56
5.5 Resultado honesto: la batería de Basilea IV falla parcialmente	57
6 Discusión y limitaciones	61
6.1 La selección QUBO como motor dominante del rendimiento	61
6.2 Similitud coseno: alcance del resultado	62
6.3 Siete amenazas a la validez	62
6.4 Implicaciones y camino hacia una ventaja práctica	65
7 Conclusiones y trabajo futuro	67
7.1 Síntesis de resultados	67
7.2 Trabajo futuro: cuatro condiciones concretas para la ventaja práctica	69
7.3 Las tres preguntas del Capítulo 1, respondidas	70
Bibliografía	74
Acrónimos	75
Apéndices	77
A Apéndices	79
A.1 Resultados completos por experimento	79
A.2 Descomposición de Euler del riesgo de cartera: cartera PA	80
A.3 Configuración completa de parámetros	81
A.4 Identificadores de trabajos en hardware	82

LISTAS

Lista de figuras

3.1	El pipeline híbrido de tres etapas	20
4.1	Dirección del gradiente y concentración de riesgo bajo ruido	36
4.2	Curvas de convergencia de CVaR bajo ruido	40
4.3	Umbral de aliasing de IQAE en IQM Emerald y Garnet	44
4.4	Comparativa de ZNE en simulador	48
4.5	Comparativa de ZNE en hardware IQM Emerald	49
5.1	Escalabilidad de la cartera PA con circuito fijo	55
5.2	Riqueza acumulada fuera de muestra	56
5.3	Backtesting regulatorio de Basilea IV	58

Lista de tablas

1.1	Comparación con el trabajo previo más cercano	6
3.1	Términos del QUBO de selección de universo	22
3.2	Justificación de los hiperparámetros principales	30
4.1	Taxonomía de experimentos	33
4.2	Métricas de calidad de universo	34
4.3	Métricas de concentración de riesgo de Euler (cartera PA)	34
4.4	Similitud coseno, error de CVaR y HHI bajo ruido (cartera PA)	36
4.5	Profundidad post-transpilación en IQM Emerald	44
4.6	Resultados de IQAE en hardware IQM	45
4.7	Métodos de ZNE sobre IQAE en simulador	48
4.8	Comparación de ZNE en hardware IQM Emerald	49
5.1	CVaR in-sample para la cartera PA	54
5.2	Experimento de escalabilidad (familia S1)	55
5.3	Rendimiento OOS de la cartera PA	56
5.4	Tests regulatorios de Basilea IV (cartera PA)	57
5.5	Detalle completo de Acerbi–Székely por nivel de confianza	58

A.1	Métricas clave de las ejecuciones de optimización	80
A.2	Descomposición de Euler del CVaR (cartera PA)	81
A.3	Configuración completa de parámetros	82

INTRODUCCIÓN

1.1. Motivación

El *valor en riesgo condicional* –CVaR, definido con precisión en la Sección 1.4– es la pérdida media esperada de una cartera en sus peores escenarios de mercado dentro de una franja de probabilidad acordada; la regulación bancaria internacional exige calcularlo, porque condiciona el capital que un banco debe reservar si ese escenario ocurre.

Calcularlo con precisión sobre cientos de activos es caro. El método clásico estándar –Monte Carlo– genera muchos escenarios aleatorios, se queda con los peores y promedia; converge al valor real, pero *despacio*: por el teorema central del límite el error decae como $1/\sqrt{N}$, es decir, coste $O(\varepsilon^{-2})$, así que duplicar la precisión cuadruplica el trabajo. La *estimación cuántica de amplitud* (QAE) resuelve un problema estructuralmente equivalente –el valor esperado de una pérdida condicionada a un evento de cola– explotando interferencia en vez de muestreo repetido, con escala $O(\varepsilon^{-1})$: duplicar la precisión solo duplica el trabajo. Esa ganancia –consecuencia del mismo mecanismo de amplificación que hace célebre al algoritmo de Grover, derivado en el Capítulo 2– motiva todo el trabajo. El cociente $\varepsilon^{-2}/\varepsilon^{-1} = \varepsilon^{-1}$ da su magnitud: para $\varepsilon = 0,005$, la precisión objetivo usada en simulación, unas $200\times$ consultas menos por llamada.

¿Por qué combinar riesgo de cartera, estimación cuántica de amplitud y hardware cuántico real –no solo simulado– en un mismo trabajo? La combinación no es obvia:

- **El riesgo de cartera no es el problema “de cartel” de la computación cuántica.** La literatura cuántico-financiera se concentra mayoritariamente en valoración de derivados (*option pricing*), el caso de uso con el que Woerner y Egger introdujeron QAE para riesgo financiero [1]. Usar CVaR de cartera conecta en cambio el algoritmo cuántico con una obligación regulatoria real (Basilea IV): una vara de medir externa –un test estadístico definido en la Sección 1.4, no una métrica *ad hoc*–.
- **Casi ningún trabajo previo ejecuta la parte cuántica en hardware real, ni con una etapa de selección de universo también cuántica.** La mayoría de los estudios comparables (Sección 1.6) son teóricos o ejecutan en hardware un circuito cuyo ancho crece con N , limitando el tamaño de cartera abordable en NISQ (§1.2). Este trabajo usa en cambio un circuito IQAE de *cuatro qubits de distribución más un ancilla (cinco físicos), de anchura fija, independiente de N* para carteras de hasta 100 activos –aunque ese circuito completo corre **en simulación**: en hardware IQM real se ejecuta solo una *sonda* mínima de dos qubits que reproduce la misma

dinámica de rondas de Grover (Contribución C1)–. Añade además una etapa previa de selección de universo formulada como QUBO para hardware D-Wave (*annealing*, no basado en puertas), ejecutada en la práctica sobre *annealing* simulado clásico (Sección 1.6): dos paradigmas cuánticos distintos en un mismo pipeline.

- **El hardware NISQ concreto (IQM) carecía de caracterización sistemática de sus límites operativos para IQAE.** No basta con saber que el algoritmo funciona en el límite ideal sin ruido: hace falta saber cuántas rondas aguanta un procesador físico concreto antes de que el ruido destruya la señal, y si conviene aplicar mitigación como ZNE –preguntas de *ingeniería de hardware*, no de teoría de algoritmos, donde se concentran las contribuciones originales de este trabajo (Sección 1.5)–.

En síntesis, el TFM no defiende que la computación cuántica ya resuelva mejor un problema financiero real que un ordenador clásico (Sección 1.7): usa ese problema, con vara de medir regulatoria externa, para *caracterizar experimentalmente* en qué condiciones un procesador cuántico ruidoso actual ejecuta IQAE de forma fiable y hasta qué precisión –el resultado central, que la aplicación financiera solo hace verificable de forma no arbitraria.

1.2. Glosario de siglas

El artículo de investigación del que parte esta memoria da por conocidas desde la primera línea las siglas especializadas que se usan aquí sin expandir; este TFM no puede permitirse esa suposición. Nueve de ellas –NISQ, QAE, IQAE, CVaR, VaR, QUBO, ZNE, MST y TMFG– están recogidas en el apartado *Acrónimos* del final del documento, con su expansión en inglés y la página donde aparecen; basta volver allí si una sigla reaparece capítulos después. Las cuatro que sostienen el argumento central se definen además en su sitio: CVaR y VaR en la Sección 1.4 de este capítulo y, paso a paso vía Rockafellar–Uryasev, en el Capítulo 2; QAE e IQAE en ese mismo capítulo. QUBO –el formato binario en el que se plantea la selección de activos antes de tocar ningún qubit de puertas, nativo del *annealing* de D-Wave y por tanto de un paradigma físico distinto del basado en puertas que usa IQAE– se formaliza en el Capítulo 2 y se aplica en el Capítulo 3.

1.3. Tres preguntas abiertas en hardware IQM

La ganancia asintótica de IQAE frente a Monte Carlo es una afirmación sobre el límite ideal, sin ruido; no dice nada sobre un procesador físico concreto, con puertas imperfectas y tiempos de coherencia finitos. La familia de procesadores superconductores de IQM –Emerald y Garnet, ambos usados aquí– carecía de caracterización sistemática publicada para tres preguntas, que estructuran el Capítulo 4 completo.

Q1: ¿cuántas rondas de Grover aguanta el hardware antes de que el ruido acumulado corrompa la estimación? Está descrito *cualitativamente* en hardware de IBM y Quantinuum [2], y hay trabajo reciente sobre variantes de menor profundidad a costa de peor convergencia [3], pero ninguno

reportaba el umbral *cuantitativo* en procesadores IQM. Responderla exige la fórmula de crecimiento de la profundidad con cada ronda y medir en cuál el hardware deja de responder con utilidad.

Q2: al estimar todo un vector de valores esperados condicionados –uno por activo–, ¿se mantiene estable el orden relativo entre componentes bajo el ruido NISQ? Es estabilidad de *ranking*, no precisión por componente: la necesita el optimizador (Capítulo 3) para decidir la dirección de actualización. El ruido de IQAE es *multiplicativo*, no aditivo, con consecuencias no triviales para cuándo se preserva el orden [4]; la condición exacta y el margen medido están en la Sección 1.5.

Q3: ¿mejora ZNE la precisión de un circuito de estimación de amplitud en NISQ, o el coste de profundidad que añade anula la ganancia? Trabajo previo cuantificó la ganancia de volumen cuántico efectivo bajo distintos esquemas de *folding* [5], pero no la interacción entre ZNE y el nivel de optimización del transpilador de Qiskit en hardware IQM. La respuesta no es un sí/no: depende de si se está por encima o por debajo del umbral de Q1.

1.4. Vehículo de validación: optimización de cartera bajo riesgo CVaR

Las tres preguntas anteriores son preguntas de física e ingeniería de hardware cuántico, respondibles en principio sobre un circuito de IQAE completamente abstracto sin aplicación financiera de por medio; este trabajo las *instancia*, sin embargo, sobre un problema concreto de optimización de cartera bajo riesgo CVaR, porque solo una aplicación externa regulada permite comprobar de forma independiente que el resultado cuántico “tiene sentido”.

Formalmente, si L denota la pérdida de una cartera (una variable aleatoria) y $\alpha \in (0, 1)$ es un nivel de confianza (por ejemplo, $\alpha = 5\%$), el valor en riesgo se define como el cuantil

$$\text{VaR}_\alpha(L) = \inf\{\ell : \mathbb{P}(L > \ell) \leq \alpha\},$$

es decir, la pérdida que solo se excede con probabilidad α . El valor en riesgo condicional es la pérdida media *más allá* de ese umbral:

$$\text{CVaR}_\alpha(L) = \mathbb{E}[L \mid L \geq \text{VaR}_\alpha(L)].$$

Nótese la diferencia conceptual: VaR solo dice *dónde* está el umbral de la cola; CVaR dice *cuánto se pierde en promedio una vez dentro de ella*, sensible por tanto a la forma de la cola más allá del umbral –una propiedad que VaR ignora por construcción, razón por la que Basilea IV migró de exigir VaR a CVaR (bajo el nombre de *Expected Shortfall*) como métrica de riesgo de mercado [6]–. El Capítulo 2 deriva la formulación de Rockafellar–Uryasev, que permite optimizar CVaR_α mediante un problema convexo en lugar de estimar cuantiles directamente: el gradiente de CVaR_α respecto a los pesos no exige diferenciar un cuantil (mal definido en general, porque VaR_α no es diferenciable), sino que se

reduce exactamente a una esperanza condicionada sobre los activos en la cola de pérdidas –el puente que conecta $CVaR_\alpha$ con IQAE–.

Calcular ese gradiente, lo que el optimizador necesita en cada iteración, exige estimar exactamente el tipo de valor esperado condicionado que IQAE está diseñada para estimar [1, 7]. La cartera óptima admite además una comprobación externa independiente de este trabajo: el test de especificación de Basilea IV de Acerbi y Székely [6, 8] (Capítulo 5), aplicado igual que a una cartera optimizada por cualquier método clásico. Así, la aplicación financiera convierte las preguntas de hardware cuántico de la Sección 1.3 en preguntas verificables, aunque las contribuciones se sitúen al nivel del procesamiento de información cuántica, no de la gestión de carteras.

1.5. Contribuciones

El paper original organiza sus contribuciones en C1 y C2 –que responden a Q1 y Q3– más una observación estructural, no numerada aparte, que da soporte a la aplicación y responde a Q2; esta sección conserva esa organización y cifras, y solo cambia el nivel de explicación técnica.

C1. Dinámica de rondas de Grover en IQM (Q1). Caracterización en hardware de IQAE en IQM Emerald y Garnet mediante una sonda mínima de codificación directa –un circuito de dos qubits, $R_y(2\theta)$ más un CNOT, que reproduce la amplitud de cola del pipeline ($\theta \approx 0,2265$) pero no su preparación de estado de cinco qubits–. La sonda obedece la fórmula de profundidad post-transpilación $\text{depth}(m) = 4 + 6m$ en puertas nativas de IQM, y el número de rondas utilizable está acotado por un umbral de *aliasing* que se *predice analíticamente* desde la geometría de la cola ($(2m+1)\theta \approx \pi/2$ en $m \approx 3$), no solo de forma empírica; como depende únicamente de la amplitud codificada, se aplica sin cambios al circuito real. Medidas en ambos dispositivos confirman el umbral predicho y cuantifican el suelo de precisión limitado por fidelidad por debajo de él, con una diferencia tentativa, de muestra pequeña, entre dispositivos en $m = 1$. El circuito propio del pipeline **no se ha ejecutado en hardware**; se reporta su profundidad tras transpilación *offline* (1,476 en $m = 0$, frente a 4 para la sonda) para dejar explícita esa brecha.

C2. Extrapolación a ruido cero sobre IQAE (Q3). Estudio comparativo ejecutado en IQM Emerald. Se comparan tres variantes de ZNE –*folding* global, *folding* de puerta, e inserción de identidad libre de *folding* [9]–, tanto en el modelo de ruido de `qiskit-aer` como en hardware real, documentando el mecanismo del fallo a nivel de transpilador `opt_level=0` y el resultado de que la estimación cruda de IQAE supera a cualquier extrapolación por debajo del umbral de *aliasing*. El análisis complementa los resultados de *folding* sensible a *layout* de [5].

Como soporte a la validación de la aplicación –no como una contribución numerada aparte– este trabajo también identifica (respondiendo a Q2) una desigualdad explícita de *brecha de ranking*: consecuencia de la desigualdad triangular aplicada a la estructura de error *multiplicativo* de IQAE, bajo la cual las estimaciones ruidosas preservan el orden entre activos que determina la dirección de actualización del optimizador Adam (Capítulos 3 y 4). Es una propiedad comprobable de una cartera concreta, no un teorema nuevo de validez general, y no la aporta la selección de la Etapa 1: la condición se cumple en el punto de partida equiponderado con el mismo margen ($\approx 1,83 \times \alpha \varepsilon = 0,005$) tanto para la cartera periférica seleccionada vía QUBO como para la cartera central de comparación. Lo que la selección QUBO sí aporta de forma verificada es diversificación estructural del universo y de su cola de pérdidas (Capítulos 5 y 6). En los pesos convergidos la condición se viola necesariamente, porque en un óptimo sobre el simplex los CVaR marginales se igualan: es una condición suficiente de *inicio de trayectoria*, no una garantía de diseño. Consistentemente, en los 46 experimentos de simulación de ruido la similitud coseno entre el subgradiente ruidoso y el ideal en ese punto de partida supera 0,9998.

El pipeline completo se valida de extremo a extremo **en simulación** sobre 374 constituyentes del S&P 500, con el circuito IQAE de anchura fija ya descrito, independiente de N . El hardware IQM real, en cambio, caracteriza de forma aislada la dinámica de rondas de Grover de IQAE (Contribución C1) vía una sonda mínima de dos qubits, no el circuito completo del pipeline ni el bucle de optimización –cuyo objetivo de precisión queda dos órdenes de magnitud por debajo del suelo del *aliasing* en dispositivo real (Sección 1.7)–. La batería de Basilea IV se usa únicamente como comprobación externa de sensatez, y el resultado es mixto: el test *unilateral* de Acerbi–Székely (estadístico Z_2) no se rechaza en ninguno de los cuatro niveles de confianza probados, pero el modelo **no** supera la batería completa de tres tests –Kupiec falla de forma conservadora salvo en $\alpha = 10\%$, Christoffersen falla en $\alpha \in \{2,5\%, 5\%\}$, y el semáforo regulatorio pasa a Amarillo/Rojo para $\alpha \geq 5\%$ –. El desglose test a test está en el Capítulo 5. No se reclama superioridad regulatoria sobre esta base.

1.6. Relación con el trabajo previo: cómo situar esta tesis

La Tabla 1.1 posiciona este trabajo frente a la literatura más cercana (traducción de la Tabla 1 del paper arXiv/TQC); sus notas a pie explican qué mide cada columna, decisiones de diseño no siempre obvias fuera del contexto NISQ.

Dos comparaciones adicionales sitúan mejor este trabajo. Vu *et al.* [3] proponen una variante de baja profundidad que evita la escala ε^{-1} a costa de garantías de sesgo más débiles –complementaria, no competidora, de la caracterización de hardware que aporta este trabajo para IQAE estándar–. El trabajo teórico más cercano, Skarlatos y Konofaos [7], estableció la complejidad de subgradiente $O(N\varepsilon^{-1})$ de forma analítica; este TFM aporta el complemento empírico –un régimen de validez medido, no derivado– sobre hardware IQM real. Contexto adicional, fuera de la comparación directa de

Trabajo	Venue	Hardware real ^a	N activos ^b	Anchura circuito ^c	Selección de red ^d	Backtest regulatorio ^e
Braine <i>et al.</i> [10]	TQE 2021	✓	pequeño	escala con N	—	—
Carrascal <i>et al.</i> [11]	TQE 2024	✓	≤ 19	escala con N	—	○
Herbert <i>et al.</i> [2]	TQE 2024	✓	n/a	1–2 qubits	—	—
Vu <i>et al.</i> [3]	ACM TQC 2025	—	n/a	variable	—	—
Agliardi <i>et al.</i> [12]	arXiv 2025 (IBM–Vanguard)	✓	109	escala con N	—	—
Skarlatos y Konofaos [7]	arXiv 2025	— (teórico)	cualquiera	4 qubits	—	—
Este trabajo	TFM (este documento)	✓ (sonda de 2 qubits)	hasta 100	4+1 qubits, fijo	○	○

^aSimula el circuito frente a ejecutarlo en un procesador real; crítico en NISQ, porque el ruido es justo lo que caracterizan Q1–Q3. ^{b,c}En casi todas las propuestas previas el ancho del circuito crece con N , lo que limita cuántos activos caben en NISQ; aquí son 4+1 qubits fijos –una única cantidad agregada por consulta– hasta $N = 100$ en simulación, y “hardware real” se refiere a la sonda de 2 qubits, no a ese circuito (C1). ^dSi la selección de activos también se resuelve en hardware cuántico: el ○ refleja que el QUBO de la Etapa 1 está formulado para D-Wave pero se ejecutó como *annealing* simulado clásico (Sección 3.2); ningún otro trabajo de la tabla aborda esta etapa. ^eBatería Basilea IV completa (Kupiec, Christoffersen, Acerbi–Székely unilateral): el ○ de Carrascal *et al.* indica Sharpe y VaR sin la batería; el de este trabajo, los resultados mixtos ya citados. Ningún trabajo alcanza ✓ en esta columna.

Tabla 1.1: Comparación con el trabajo previo más cercano (traducción directa de la Tabla 1 del paper). ✓: presente; ○: parcial; —: ausente.

la Tabla 1.1: revisiones de aplicaciones cuántico-financieras [13] y de *benchmarks* de opciones en hardware IBM [14].

1.7. Alcance y advertencias

Este trabajo *no* reclama una ventaja computacional cuántica a las escalas probadas ($N \leq 100$). El valor $p \approx 0,12$ (*bootstrap*, a una cola) al comparar la cartera PA frente al optimizador clásico SLSQP no es estadísticamente significativo ni a una cola; la tasa de victorias del 54.3 % observada a través de todos los experimentos tampoco lo es; y la ventana fuera de muestra usada para la validación cubre un único régimen de mercado alcista. Estas limitaciones se retoman con más detalle, y sin suavizarlas, en el Capítulo 6.

La ventaja práctica del pipeline está condicionada a que se cumplan *simultáneamente* tres condiciones:

- (i) $\varepsilon < 0,05$, para que la escala $O(\varepsilon^{-1})$ de IQAE efectivamente supere a Monte Carlo clásico en número de consultas necesarias;
- (ii) una *brecha mínima entre pares* de componentes del vector objetivo mayor que el término de ruido $2\varepsilon \max_k |g_k^*|$ de la cota de inversión de ranking (derivación, margen medido y su pérdida necesaria cerca del óptimo: Sección 1.5);
- (iii) un número de rondas de Grover que se mantenga por debajo del umbral de *aliasing* de IQAE (Contribución C1): ese umbral lo fija la geometría de la cola, no la fidelidad de puerta, que solo determina la precisión por debajo de él; alcanzar rondas más profundas exigiría diferir algorítmicamente hacia una amplitud codificada más pequeña, junto con la fidelidad necesaria para ejecutar los circuitos resultantes (Capítulos 4 y 6).

El hardware IQM actual satisface (i) y (ii) solo en simulación; la condición (iii) limita la precisión alcanzable en dispositivo real a $\varepsilon_{\text{hw}} \approx 0,20$ —muy por encima del umbral $\varepsilon < 0,05$ que (i) exige—. Las contribuciones de este TFM deben leerse, por tanto, como una *caracterización de régimen de validez* —dónde funciona el algoritmo, dónde deja de funcionar, y por qué— y no como una demostración de que la ventaja cuántica ya está aquí: esa distinción, más que ninguna cifra concreta, es el resultado que este TFM pide al lector que se lleve de este primer capítulo.

ESTADO DEL ARTE

Este capítulo desarrolla las cuatro piezas teóricas sobre las que se apoya todo el resto del TFM: la estimación cuántica de amplitud, el riesgo de cartera que este trabajo optimiza y por qué admite un problema convexo, la obligación regulatoria que motiva esa medida de riesgo, y la optimización binaria que decide –antes de cualquier circuito cuántico de puertas– qué activos entran en la cartera. El artículo del que parte esta memoria da por conocidas las cuatro piezas y se limita a citarlas; aquí se derivan paso a paso, con referencias al final de cada sección a los capítulos donde reaparece cada pieza.

2.1. De la estimación de amplitud clásica a la cuántica

Antes de hablar de circuitos conviene fijar el problema clásico: *estimar el valor esperado de una variable aleatoria*. Si X es acotada y se dispone de N muestras independientes, el estimador de Monte Carlo $\hat{X} = \frac{1}{N} \sum_k X^{(k)}$ es insesgado y su error típico decae como $\sigma(X)/\sqrt{N}$: fijar un error objetivo ε exige $N = O(\varepsilon^{-2})$ muestras. El caso que interesa aquí es X de Bernoulli: la cantidad que el circuito cuántico estima es la probabilidad de cola $a_\tau = \mathbb{P}(L > \tau)$, y el CVaR se recupera de a_τ y de las esperanzas condicionadas asociadas mediante un estimador de razón (Capítulo 3). Estimar a por muestreo directo hereda la escala $O(\varepsilon^{-2})$; la *estimación cuántica de amplitud* resuelve el mismo problema con escala $O(\varepsilon^{-1})$ por una razón geométrica, no estadística: en vez de repetir un experimento aleatorio, amplifica una amplitud de probabilidad mediante rotaciones deterministas y solo mide al final.

El operador $\mathcal{A}$ y el ángulo θ

Sea $\mathcal{A}$ un operador unitario sobre un registro cuántico –el sistema más un qubit auxiliar– tal que, partiendo de $|0\rangle$,

$$\mathcal{A}|0\rangle = \sin \theta |\psi_1\rangle + \cos \theta |\psi_0\rangle,$$

donde $|\psi_0\rangle$ y $|\psi_1\rangle$ son estados ortogonales que distinguen el resultado “malo” del “bueno” de un evento de interés y $\theta \in [0, \pi/2]$ codifica la amplitud a estimar, $a = \sin^2 \theta$. $\mathcal{A}$ es el análogo cuántico de generar

una muestra clásica; construirlo para el caso de CVaR es tarea del Capítulo 3. Medir $\mathcal{A}|0\rangle$ en ese punto daría el resultado “bueno” con probabilidad exactamente a : un solo bit de Bernoulli sesgado. La ganancia cuántica no está en $\mathcal{A}$, sino en lo que se puede hacer con él *antes* de medir.

El operador de Grover como rotación

Defínanse dos reflexiones: $S_\chi = I - 2|\psi_1\rangle\langle\psi_1|$, que invierte el signo de la componente “buena”, y $S_0 = I - 2|0\rangle\langle 0|$, que invierte el de la componente a lo largo del estado inicial. El *operador de Grover* compone ambas, conjugando la segunda por $\mathcal{A}$: $\mathcal{Q} = \mathcal{A} S_0 \mathcal{A}^\dagger S_\chi$. Dentro del subespacio real de dimensión dos generado por $\{|\psi_0\rangle, |\psi_1\rangle\}$ —el único donde vive el estado relevante— S_χ es la reflexión respecto al eje $|\psi_0\rangle$ y $\mathcal{A} S_0 \mathcal{A}^\dagger$ la reflexión respecto al eje que marca $\mathcal{A}|0\rangle$, a ángulo θ . Componer dos reflexiones separadas un ángulo θ es una rotación de 2θ , así que aplicar $\mathcal{Q}$ un total de m veces lleva el estado a $(2m + 1)\theta$:

$$\langle\psi_1|\mathcal{Q}^m\mathcal{A}|0\rangle = \sin((2m + 1)\theta). \quad (2.1)$$

Esta es la identidad clave de la estimación de amplitud [15]: la probabilidad de obtener el resultado “bueno” tras $\mathcal{Q}^m$ es $\sin^2((2m + 1)\theta)$, que oscila $(2m + 1)$ veces más deprisa en θ que la de una sola muestra —la señal que permite localizar θ , y con ella a , con menos repeticiones que el muestreo directo—.

La versión canónica de QAE [15] explota esa oscilación con estimación de fase: un registro auxiliar de $t = O(\log(1/\varepsilon))$ qubits, puertas $\mathcal{Q}^{2^j}$ controladas y una transformada de Fourier inversa bastan para leer θ con precisión ε usando $O(\varepsilon^{-1})$ aplicaciones de $\mathcal{Q}$. El problema práctico es que esas puertas controladas tienen profundidad exponencial en j y van entrelazadas con el registro auxiliar: el hardware NISQ actual no puede ejecutarlas de forma fiable. Evitar ese cuello de botella sin renunciar a la escala $O(\varepsilon^{-1})$ es exactamente lo que resuelve la variante iterativa.

2.2. Estimación Iterativa de Amplitud

La *estimación iterativa de amplitud cuántica* (IQAE, la variante de QAE introducida por Grinko *et al.* [16] que sustituye el circuito de estimación de fase por rondas de Grover y estadística clásica) parte de una observación simple: la identidad (2.1) ya contiene toda la información necesaria sobre θ sin registro auxiliar ni puerta controlada. Basta preparar, para una ronda m fija, $\mathcal{Q}^m\mathcal{A}|0\rangle$, medirlo N_m veces, y usar la fracción empírica de resultados “buenos” $\hat{p}_m$ como estimador de $\sin^2((2m + 1)\theta)$.

De una proporción muestral a un intervalo de confianza en θ

Esto convierte el problema en una estimación clásica de un parámetro de Bernoulli, salvo que el parámetro de interés no es $\hat{p}_m$ sino θ , relacionado con $\hat{p}_m$ mediante una función no lineal ($\arcsin$) y $2m + 1$ posibles ramas (porque $\sin^2$ no es inyectiva). Grinko *et al.* [16] resuelven esto con un procedimiento adaptativo: en cada ronda calculan un intervalo de confianza de Clopper–Pearson (exacto, no asintótico, para una proporción binomial) sobre $\hat{p}_m$, lo invierten sobre θ eligiendo la rama compatible con la ronda anterior –lo que evita la ambigüedad de fase–, y solo aumentan m cuando la nueva ronda no introduce una ambigüedad que el intervalo actual no pueda resolver. El procedimiento se detiene en cuanto la semianchura del intervalo acumulado sobre θ satisface $\hat{\varepsilon} \leq \varepsilon_{\text{objetivo}}$.

Proposición (complejidad de IQAE, [16]).

IQAE alcanza precisión ε con confianza $1 - \delta$ usando

$$O(\varepsilon^{-1} \log(1/\delta) \log \log(1/\varepsilon))$$

consultas al oráculo $\mathcal{A}$; el factor doblemente logarítmico es pequeño en la práctica y a menudo se omite en enunciados informales, dando la forma habitualmente citada $\frac{\pi}{2\varepsilon} \log(2/\delta)$.

Sin reproducir la prueba completa (extensa en el original), la intuición es que, si la ronda m se va doblando aproximadamente de una iteración a la siguiente ($m_k \sim 2^k$), bastan $O(\log(1/\varepsilon))$ rondas para que m_k alcance $O(1/\varepsilon)$, cada una con $N_{m_k} = O(\log(1/\delta))$ repeticiones prácticamente constantes; el coste telescopa a $O(\varepsilon^{-1} \log(1/\delta))$ más el factor doblemente logarítmico de la gestión de ramas. Esta escala $O(\varepsilon^{-1})$ supera a los $O(\varepsilon^{-2})$ de Monte Carlo (Sección 2.1) cuando $\varepsilon < 0,05$ (Sección 1.7); para el $\varepsilon = 0,005$ de este trabajo, la ganancia teórica es de aproximadamente $200\times$ en consultas por estimación –cifra ya adelantada, sin repetir su derivación, en el Capítulo 1–.

Herbert *et al.* [2] refinan la caracterización del ruido de estos procedimientos bajo hardware realista, validando un modelo de ruido gaussiano contra datos de IBM y Quantinuum –complementario, no sustituto, de la caracterización empírica en hardware IQM que aporta este TFM (Capítulo 4), retomada en la Sección 1.3 del Capítulo 1–.

2.3. Valor en riesgo condicional y la formulación de Rockafellar–Uryasev

Toca ahora fijar con precisión matemática la cantidad que este trabajo usa como vehículo de validación: el *valor en riesgo condicional* (CVaR, la pérdida media de una cartera condicionada a que esa pérdida esté entre las peores de una cola de probabilidad fijada de antemano), presentado de forma

intuitiva en la Sección 1.4. Aquí se repite la definición formal y, sobre todo, se demuestra –no solo se enuncia– la propiedad que hace de CVaR una cantidad optimizable.

Sea $\mathbf{w} \in \mathbb{R}^N$ un vector de pesos de cartera y $\mathbf{r} = (r_1, \dots, r_N)$ un vector aleatorio de retornos. La pérdida de la cartera es $L(\mathbf{w}) = -\mathbf{w}^\top \mathbf{r}$. Para un nivel de cola $\alpha \in (0, 1)$ (por ejemplo $\alpha = 0,05$),

$$\text{VaR}_\alpha(\mathbf{w}) = \inf\{z \in \mathbb{R} : \mathbb{P}(L(\mathbf{w}) > z) \leq \alpha\}, \quad \text{CVaR}_\alpha(\mathbf{w}) = \mathbb{E}[L(\mathbf{w}) \mid L(\mathbf{w}) \geq \text{VaR}_\alpha(\mathbf{w})].$$

La diferencia entre ambas cantidades importa para todo lo que sigue: VaR solo localiza la frontera de la cola, CVaR resume cuánto se pierde *dentro* de ella.

Por qué CVaR y no VaR: coherencia

Artzner *et al.* [17] llaman *coherente* a una medida de riesgo que cumple cuatro propiedades: subaditividad ($\rho(\mathbf{w}_1 + \mathbf{w}_2) \leq \rho(\mathbf{w}_1) + \rho(\mathbf{w}_2)$, la formalización de que diversificar nunca debería *aumentar* el riesgo reportado), monotonidad, homogeneidad positiva e invariancia por traslación. VaR_α *no* es subaditivo –pueden construirse dos posiciones cuyo VaR combinado supere la suma de los individuales, porque VaR solo mira dónde empieza la cola y es ciego a lo que hay más allá–, mientras que CVaR_α sí satisface las cuatro. Esta es la razón matemática, no solo regulatoria, por la que CVaR es la elección natural para *optimizar* el riesgo de una cartera; la razón regulatoria se desarrolla en la Sección 2.4.

La representación de Rockafellar–Uryasev

El problema práctico de optimizar $\text{CVaR}_\alpha(\mathbf{w})$ directamente sobre $\mathbf{w}$ es que su definición exige primero calcular $\text{VaR}_\alpha(\mathbf{w})$, un cuantil no diferenciable en general –diferenciar CVaR_α respecto a $\mathbf{w}$ parecería exigir diferenciar algo que no lo es–. Rockafellar y Uryasev [18] resuelven esto con una caracterización auxiliar que evita el problema. Fíjese $\mathbf{w}$ y defínase, para $\zeta \in \mathbb{R}$,

$$F(\zeta) = \zeta + \frac{1}{\alpha} \mathbb{E}[(L(\mathbf{w}) - \zeta)^+], \quad (x)^+ \triangleq \max(x, 0).$$

Teorema (Rockafellar–Uryasev, [18]).

Si $L(\mathbf{w})$ tiene función de distribución continua F_L , entonces $F(\zeta)$ es convexa en ζ , se minimiza exactamente en $\zeta^* = \text{VaR}_\alpha(\mathbf{w})$, y su valor mínimo es $\text{CVaR}_\alpha(\mathbf{w})$:

$$\text{CVaR}_\alpha(\mathbf{w}) = \min_{\zeta \in \mathbb{R}} \left\{ \zeta + \frac{1}{\alpha} \mathbb{E}[(L(\mathbf{w}) - \zeta)^+] \right\}. \quad (2.2)$$

Demostración.

Derivando bajo la integral (válido porque F_L es continua),

$$F'(\zeta) = 1 - \frac{1}{\alpha} \int_{\zeta}^{\infty} f_L(\ell) d\ell = 1 - \frac{1}{\alpha} \mathbb{P}(L(\mathbf{w}) > \zeta), \quad F''(\zeta) = \frac{1}{\alpha} f_L(\zeta) \geq 0,$$

así que F es convexa. Igualando $F'(\zeta) = 0$ se obtiene $\mathbb{P}(L(\mathbf{w}) > \zeta) = \alpha$, es decir $F_L(\zeta) = 1 - \alpha$, que es exactamente la definición de $\zeta^* = \text{VaR}_{\alpha}(\mathbf{w})$: el punto crítico único de una función convexa es su mínimo global. Queda comprobar que $F(\zeta^*) = \text{CVaR}_{\alpha}(\mathbf{w})$: usando $\mathbb{P}(L(\mathbf{w}) \geq \zeta^*) = \alpha$ (continuidad de F_L), separar el valor esperado en $F(\zeta^*)$ sobre el evento de cola y reagrupar términos deja $F(\zeta^*) = \mathbb{E}[L(\mathbf{w}) \mid L(\mathbf{w}) \geq \zeta^*] = \text{CVaR}_{\alpha}(\mathbf{w})$, que es la definición buscada. ■

El paper original cita la representación (2.2) sin demostración; aquí se reproduce la prueba completa porque es precisamente el resultado que, según se adelantó en la Sección 1.4, convierte un problema de cuantiles en un problema convexo estándar, resoluble por descenso de gradiente.

El subgradiente: por qué no hace falta diferenciar un cuantil

La ganancia práctica no es solo que $\text{CVaR}_{\alpha}(\mathbf{w})$ se pueda escribir como un mínimo: es que ese mínimo se puede *derivar* respecto a $\mathbf{w}$ sin volver a tocar ζ^* . Sea $g(\mathbf{w}) = \min_{\zeta} F(\mathbf{w}, \zeta)$ la función de valor de (2.2), ahora escrita explícitamente como función conjunta de $(\mathbf{w}, \zeta)$. Por el teorema de la envolvente (Danskin), cuando el minimizador $\zeta^*(\mathbf{w})$ es único, el gradiente de la función de valor respecto a $\mathbf{w}$ es el gradiente parcial de F respecto a $\mathbf{w}$ evaluado en $\zeta^*(\mathbf{w})$, sin término adicional por cómo ζ^* varía con $\mathbf{w}$: esa contribución se anula porque $\partial F / \partial \zeta = 0$ en el óptimo. Diferenciando $\max(L(\mathbf{w}) - \zeta, 0)$ respecto a w_i en los puntos donde $L(\mathbf{w}) \neq \zeta$ (en casi todo punto, dado que L es continua) se obtiene $\mathbf{1}_{L(\mathbf{w}) \geq \zeta} \cdot \partial L / \partial w_i$, y con $L(\mathbf{w}) = -\mathbf{w}^{\top} \mathbf{r}$ se tiene $\partial L / \partial w_i = -r_i$. Sustituyendo $\zeta = \zeta^*(\mathbf{w}) = \text{VaR}_{\alpha}(\mathbf{w})$:

$$g_i(\mathbf{w}) \triangleq \frac{\partial \text{CVaR}_{\alpha}}{\partial w_i} = \frac{1}{\alpha} \mathbb{E}[-r_i \mathbf{1}_{L(\mathbf{w}) \geq \text{VaR}_{\alpha}(\mathbf{w})}] = -\mathbb{E}[r_i \mid L(\mathbf{w}) \geq \text{VaR}_{\alpha}(\mathbf{w})], \quad (2.3)$$

usando $\mathbb{P}(L(\mathbf{w}) \geq \text{VaR}_{\alpha}(\mathbf{w})) = \alpha$ en el último paso. Este es el resultado que se anunció, sin derivar, en la Sección 1.4: el gradiente de CVaR_{α} respecto a un peso de cartera *no* exige diferenciar el cuantil $\text{VaR}_{\alpha}(\mathbf{w})$, sino que se reduce exactamente a (menos) el retorno esperado del activo i , condicionado a que la cartera esté en su propia cola de pérdidas. Esta esperanza condicionada —un valor esperado sobre un evento de cola, la misma estructura de $a = \sin^2 \theta$ de la Sección 2.1— es la que IQAE está diseñada para estimar [1, 7]: la conexión entre las dos primeras secciones de este capítulo es exactamente esta ecuación.

Estimador histórico y descomposición de Euler

Con T escenarios de retorno observados $\{\mathbf{r}^{(t)}\}_{t=1}^T$ y pérdidas realizadas $L^{(t)} = -\mathbf{w}^\top \mathbf{r}^{(t)}$, el estimador histórico de CVaR_α –versión discreta de la definición continua anterior– es la media de las peores pérdidas observadas, las que igualan o superan el VaR muestral:

$$\widehat{\text{CVaR}}_\alpha(\mathbf{w}) = \frac{1}{|\{t : L^{(t)} \geq \widehat{\text{VaR}}_\alpha\}|} \sum_{t : L^{(t)} \geq \widehat{\text{VaR}}_\alpha} L^{(t)}.$$

El denominador es el número *real* de observaciones en la cola, no $\lfloor T\alpha \rfloor$: el percentil que define $\widehat{\text{VaR}}_\alpha$ junto con una máscara $\geq$ produce sistemáticamente $\lfloor T\alpha \rfloor + 1$ elementos, convención consistente en las tres rutas de cómputo del código, incluida la descomposición de Euler de este capítulo. Un resultado adicional de Tasche [19] es que la contribución de riesgo de cada activo, $\text{RC}_i(\mathbf{w}) = w_i \cdot g_i(\mathbf{w})$, suma *exactamente* el total: $\sum_{i=1}^N \text{RC}_i(\mathbf{w}) = \widehat{\text{CVaR}}_\alpha(\mathbf{w})$. La razón es puramente matemática: $\widehat{\text{CVaR}}_\alpha(\mathbf{w})$ es homogénea de grado uno en $\mathbf{w}$ –escalar la cartera por $\lambda > 0$ escala la pérdida, y por tanto el riesgo, por el mismo λ , la propiedad (iii) de coherencia ya vista–, y el teorema de Euler para funciones homogéneas de grado uno da $\sum_i w_i \partial f / \partial w_i = f(\mathbf{w})$ para cualquier f con esa propiedad. Esta descomposición –base del índice de concentración de Herfindahl usado en la construcción de cartera– se retoma con cifras reales en el Capítulo 4.

2.4. Basilea IV, VaR y CVaR regulatorio desde cero

Las dos secciones anteriores han sido matemáticas; esta explica por qué la regulación bancaria obliga a usar CVaR –bajo el nombre de *expected shortfall*– para calcular el capital que un banco reserva frente a su cartera de negociación.

Basilea II (2004) usaba $\text{VaR}_{99\%}$ para el riesgo de mercado de la cartera de negociación. La crisis de 2008 expuso dos problemas: uno estadístico –dos carteras con idéntico $\text{VaR}_{99\%}$ pueden tener perfiles de pérdida catastrófica radicalmente distintos– y uno estructural, ya visto arriba: VaR no es coherente y puede penalizar la diversificación. La respuesta es la *Revisión Fundamental de la Cartera de Negociación* (FRTB), publicada en forma final por el BCBS en enero de 2019 [6]: sustituir $\text{VaR}_{99\%}$ por $\text{CVaR}_{97.5\%}$ –el *expected shortfall* regulatorio, matemáticamente idéntico al CVaR de la Sección 2.3– como métrica estándar de capital.

El cambio plantea un problema de verificación: para VaR basta contar cuántos días la pérdida real superó la predicha (test de Kupiec), mientras que CVaR, al ser ya un promedio sobre la cola, no admite un recuento de excepciones tan directo. Acerbi y Székely [8] cubren esa brecha con los estadísticos de verificación acumulada Z_1 y Z_2 , hoy referencia para el backtesting regulatorio de CVaR. Su desarrollo completo –junto con Kupiec y Christoffersen– y su aplicación a las carteras de este trabajo se reservan para el Capítulo 5. De ahí que optimizar CVaR sea un vehículo de validación no arbitrario: el resultado

cuántico se contrasta contra un estándar regulatorio que ninguna decisión de diseño de este trabajo puede sesgar a su favor.

2.5. QUBO y métodos de penalización

La última pieza de este capítulo es de optimización combinatoria clásica –resuelta, como se verá, en un tipo de hardware cuántico completamente distinto al de las Secciones 2.1–2.2– y precede en el pipeline (Capítulo 3) a todo lo anterior: *antes* de optimizar los pesos de una cartera de N activos bajo CVaR, hay que decidir *qué* N activos –de un universo mucho mayor– entran siquiera en ella, una decisión binaria que admite una formulación estándar de optimización combinatoria.

Optimización binaria cuadrática sin restricciones

Un problema de *optimización binaria cuadrática sin restricciones* (QUBO) tiene la forma

$$\min_{\mathbf{x} \in \{0,1\}^n} \mathbf{x}^\top Q \mathbf{x}$$

para una matriz real $n \times n$ Q , simétrica sin pérdida de generalidad. El formato es general por una identidad trivial: como $x_i \in \{0, 1\}$ se tiene $x_i^2 = x_i$, así que cualquier polinomio en variables booleanas se reescribe con grado a lo sumo dos –los términos de orden superior se reducen mediante variables auxiliares–. Es NP-difícil: equivale, salvo cambio de variable, a hallar el estado fundamental de un modelo de Ising clásico.

El formato no admite restricciones explícitas –cualquier $\mathbf{x} \in \{0, 1\}^n$ es ya “válida”–, así que toda restricción real (p. ej., “selecciona exactamente k de los n candidatos”) entra como *penalización*: si la restricción es $f(\mathbf{x}) = c$, se añade $\lambda (f(\mathbf{x}) - c)^2$ con $\lambda > 0$ calibrada para dominar el rango del objetivo sin distorsionar numéricamente el paisaje de optimización. Expandiendo el cuadrado con $x_i^2 = x_i$, la penalización se reduce a una forma cuadrática pura que se suma directamente a Q ; el Capítulo 3 desarrolla el cálculo con los coeficientes reales de este trabajo.

Hardware de *annealing* y precedentes de aplicación

El QUBO es también el formato de entrada nativo del hardware de *annealing* cuántico: en lugar de puertas discretas, el procesador evoluciona físicamente un sistema de espines a lo largo de un Hamiltoniano que interpola entre uno trivial y uno cuyo estado fundamental codifica la solución óptima. El solucionador híbrido de D-Wave [20] combina *annealing* cuántico con heurísticas clásicas y acepta instancias de hasta 10^4 variables binarias con soluciones casi óptimas en segundos [21] –un paradigma distinto del basado en puertas de IQAE; este trabajo usa ambos en el mismo pipeline–.

Hay precedentes de QUBO sobre decisiones financieras reales en IEEE *Transactions on Quantum Engineering*: Braine *et al.* [10] lo aplican a la liquidación de transacciones –intercambio de valores y efectivo entre contrapartes– en simulación y en hardware de IBM, y Harwood *et al.* [22] a rutas de vehículos con ventanas de tiempo, comparando formulaciones QUBO e Ising. Ninguno resuelve el problema de este TFM –selección de universo por periferalidad de red–, pero ambos validan el QUBO como vía práctica para llevar una decisión financiera a hardware cuántico. La formulación concreta que usa este trabajo, con el ejemplo numérico completo, está en el Capítulo 3.

2.6. Selección de universo: el árbol de expansión mínima de Mantegna y la periferalidad de Pozzi *et al.*

Queda abierta una pregunta de la Sección 2.5: ¿qué *criterio* decide qué activos seleccionar, es decir, qué codifica la matriz Q ? La respuesta viene de la econofísica de redes de mercado, que trata la correlación entre activos como un grafo.

Una distancia, no solo una correlación

Dadas las series de retornos (estandarizadas) de dos activos i, j con correlación $\rho_{ij} \in [-1, 1]$, defínase

$$d_{ij} = \sqrt{2(1 - \rho_{ij})}.$$

Esta cantidad no es una elección arbitraria: tratando X_i, X_j (normalizadas, media cero y varianza uno) como vectores en el espacio de escenarios, $\|X_i - X_j\|^2 = \|X_i\|^2 + \|X_j\|^2 - 2\langle X_i, X_j \rangle = 2 - 2\rho_{ij}$, de modo que d_{ij} es literalmente la distancia euclídea entre ambos vectores y, por tanto, cumple la desigualdad triangular y las propiedades de una métrica –algo que ρ_{ij} por sí sola no ofrece–, lo que permite tratar los activos como puntos de un espacio métrico.

El árbol de expansión mínima y la periferalidad de red

Mantegna [23] propuso construir, a partir de esta distancia, el grafo completo de N activos y extraer su *árbol de expansión mínima* (MST): el subgrafo de $N - 1$ aristas, sin ciclos, que conecta todos los activos minimizando la suma de pesos, y que revela una estructura jerárquica del mercado –sectores agrupados en ramas, unidos por “puentes”–. El pipeline usa, en vez del MST, su refinamiento planar, el grafo filtrado triangulado maximal (TMFG): parte del MST como esqueleto y añade de forma voraz, preservando la planaridad, las aristas de correlación más fuerte hasta $3(M - 2)$ para M activos (Capítulo 3).

Dentro de esa red hay nodos *centrales* –alta centralidad de intermediación– y *periféricos* –hojas o ramas terminales–. Pozzi *et al.* [24] mostraron que la distinción tiene consecuencias financieras medibles: los activos periféricos generan sistemáticamente carteras con mejor rendimiento ajustado por riesgo y menores pérdidas de cola, y cuantifican la perifericidad con un índice $X_i + Y_i$ que combina los rangos de un nodo en cinco medidas de grafo. (El símbolo π_i se reserva aquí para el índice que el pipeline implementa de verdad –uno menos la centralidad de autovector normalizada–, definido en el Capítulo 3.) La intuición financiera es que los activos periféricos, al pertenecer a sectores distintos, conservan independencia estadística incluso cuando las correlaciones entre activos centrales convergen a uno en pánico de mercado, lo que dispersa las componentes de $\{g_i(\mathbf{w})\}$ y las hace más informativas para el optimizador –propiedad emparentada con la condición (ii) de la Sección 1.7, aunque la cantidad que gobierna la preservación del ranking sea la brecha *mínima entre pares*, no la dispersión global–.

Estos son los precedentes sobre los que se apoya la selección de universo: el MST de Mantegna como andamio de conectividad del TMFG; el índice de Pozzi *et al.* como inspiración –no copia– de π_i ; la maximización voraz de modularidad de Newman [25] como método de respaldo para detectar comunidades, con Louvain por defecto; y Grande y Borondo [26], que aplican PMFG/TMFG sobre redes de cointegración de criptomonedas para *pairs trading* y validan de forma clásica e independiente el mismo principio de econofísica de redes. Frente a este último, el precedente más cercano, la contribución de este TFM es formular la selección de red –que ellos resuelven con métodos puramente clásicos– como un problema QUBO (Sección 2.5) construido a partir de π_i , en el formato nativo del *annealing* cuántico –resuelto por defecto con *annealing* simulado clásico, con el solver híbrido de D-Wave como *backend* opcional–, dentro de un pipeline híbrido de extremo a extremo, sin que ninguno de los dos enfoques reclame superioridad.

Con las cuatro piezas ya derivadas, el Capítulo 3 las ensambla en el pipeline completo, con un ejemplo numérico trabajado paso a paso.

METODOLOGÍA

3.1. El pipeline híbrido de tres etapas

El Capítulo 1 dejó planteado el problema y el Capítulo 2 deriva las piezas matemáticas que lo resuelven (la reducción de Rockafellar–Uryasev del gradiente del valor en riesgo condicional, CVaR, a una esperanza condicionada; la escala $O(\varepsilon^{-1})$ de la estimación cuántica de amplitud, QAE; el formato de optimización binaria cuadrática sin restricciones, QUBO). Este capítulo responde a una pregunta distinta y más operativa: dado todo eso, ¿qué hace exactamente el programa, paso a paso, para pasar de “374 activos del S&P 500” a “una cartera concreta de pesos w^* ”?

La respuesta tiene tres partes —la Figura 3.1 las resume de un vistazo—, porque hay tres preguntas distintas que responder en orden. **¿Qué activos entran en la cartera?** En vez de optimizar los 374 pesos a la vez, se preselecciona un subconjunto pequeño y estructuralmente diverso de N activos antes de tocar ningún qubit de puertas: la **Etapas 1** (Sección 3.2), formulada como QUBO —el formato nativo del *annealing* de D-Wave— y resuelta en los universos reportados con *annealing* simulado clásico. **Dada una cartera candidata, ¿en qué dirección hay que mover cada peso para reducir $\text{CVaR}_\alpha(w)$?** Ese subgradiente es, gracias a la reducción de Rockafellar–Uryasev, exactamente una esperanza condicionada sobre los activos en la cola de pérdidas: el tipo de cantidad que IQAE estima. Es la **Etapas 2** (Sección 3.3), resuelta en las 46 corridas en simulación (*qiskit-aer*) con un circuito de cuatro qubits de distribución más un ancilla, de anchura fija e independiente de N . **Dada esa dirección, ¿cómo se actualizan los pesos sin violar las restricciones?** Es la **Etapas 3** (Sección 3.4): un Adam puramente clásico que consume la estimación ruidosa de la Etapa 2, con los pesos *reparametrizados* vía softmax de parámetros libres, de modo que cualquier actualización produce automáticamente una cartera válida —y no por tecnicismo, sino porque el enfoque de proyección tiene un punto fijo espurio real (Sección 3.4.1)—.

Las Etapas 1 y 2 son las únicas con contacto cuántico —y de paradigmas físicos distintos, *annealing* frente a circuitos de puertas—, mientras que la Etapa 3 es enteramente clásica. En las ejecuciones reportadas ese contacto es efectivo solo en la Etapa 2 (circuitos IQAE simulados) y meramente *potencial* en la Etapa 1 (*annealing* simulado por defecto, D-Wave híbrido como *backend* opcional): el pipeline es

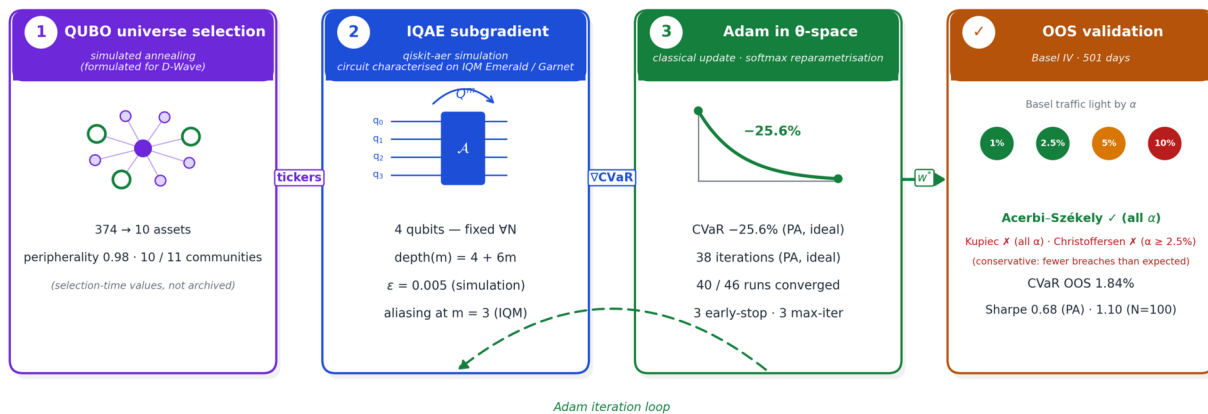


Figura 3.1: El pipeline híbrido de tres etapas (Figura 1 del paper). Etapa 1: selección de universo vía QUBO, $374 \rightarrow 10$ activos. Etapa 2: subgradiente de CVaR vía IQAE con circuito de anchura fija, simulado en qiskit-aer en las 46 corridas —las ejecuciones en IQM caracterizan una sonda de dos qubits, no este circuito—. Etapa 3: Adam clásico sobre la reparametrización softmax. El cuarto bloque es la validación fuera de muestra de 501 días. La flecha discontinua marca el bucle entre las Etapas 2 y 3.

una arquitectura híbrida con dos puntos de contacto cuántico independientes, no “un algoritmo cuántico” en sentido único. La distinción importa para el Capítulo 4: la caracterización en IQM (C1 y C2) informa sobre la dinámica de Grover medida en una sonda de dos qubits, no sobre el circuito propio de la Etapa 2 —que nunca se ejecutó en hardware— ni sobre el pipeline completo, que solo se valida en simulación.

3.2. Etapa 1: selección de universo vía QUBO

Por qué hace falta esta etapa

Dos activos con una correlación muy alta —dos bancos grandes, por ejemplo— aportan casi la misma información de riesgo, y un universo de partida de $M = 374$ activos contiene muchísima redundancia de este tipo. Antes de optimizar pesos conviene preguntarse qué subconjunto de N activos, mucho menor que M , captura la mayor diversidad estructural del mercado completo —en lugar de escoger los más líquidos o de mayor capitalización, que es justo lo que hace la cartera de comparación PC de la Sección 3.5, un *mega-cap basket* manual sin ningún criterio de red—.

La idea de este trabajo, formalizada como el QUBO que se resuelve más abajo (*optimización binaria cuadrática sin restricciones*: variables binarias, objetivo cuadrático, restricciones codificadas como penalizaciones; Capítulo 2 para la formalización general), es preferir activos *periféricos*: débilmente co-

rrelacionados con el grueso del mercado y, por tanto, portadores de información no redundante. Medir “perifericidad” con rigor requiere una *red de correlación filtrada*; la familia de filtros arranca en el MST (*árbol de expansión mínima*, introducida por Mantegna [23] y retomada en el Capítulo 2), y el pipeline real usa su refinamiento TMFG (*Triangulated Maximally Filtered Graph*), detallado en la siguiente subsección.

3.2.1. De la correlación a la red filtrada

Dada la matriz de correlación de Pearson ρ del periodo de entrenamiento se define primero una distancia entre cada par de activos:

$$d_{ij} = \sqrt{2(1 - \rho_{ij})}. \quad (3.1)$$

Dos activos perfectamente correlacionados tienen distancia cero; dos perfectamente anticorrelacionados, la máxima $d_{ij} = 2$. El grafo completo con estas $\binom{M}{2}$ distancias tiene demasiadas aristas para revelar estructura, así que hay que filtrarlo. El filtro real del pipeline no es el MST sino su refinamiento planar, el TMFG: partiendo del MST como esqueleto de conectividad se añaden de forma voraz las aristas de correlación más fuerte que no rompan la planaridad, hasta $3(M - 2)$ aristas. El MST es un andamio interno –garantiza que la red quede conectada–, no la red final.

Sobre esa red, el pipeline detecta K comunidades mediante Louvain (semilla fija 42); la maximización voraz de modularidad de Newman [25] existe en el código solo como *fallback*. Sobre los 374 activos, en la vintage congelada al seleccionar, produjo $K = 11$ comunidades de tamaños anotados $\{57, 49, 44, 39, 36, 33, 33, 28, 23, 16, 16\}$ –sin artefacto archivado; la reejecución sobre la vintage actual da $K = 13$ (Sección 3.2.3)–. Finalmente, para cada activo se calcula un índice de perifericidad

$$\pi_i = 1 - \tilde{c}_i, \quad (3.2)$$

con $\tilde{c}_i$ la centralidad de autovector en la red filtrada –correlaciones como pesos de arista– normalizada al rango $[0, 1]$. Esa centralidad mide cuánto participa un nodo en el núcleo correlacionado de la red, así que su complemento es máximo para los activos más desacoplados: los que aportan información no redundante. El índice implementa el *principio* de Pozzi *et al.* [24], no su fórmula $X + Y$ de dos componentes. Nótese que “periférico” aquí no equivale a “lejos del centro del grafo”: un nodo puente entre dos comunidades puede estar débilmente acoplado en correlación y ser a la vez crítico para el contagio entre sectores, dimensión que captura un término separado del QUBO –la penalización por intermediación, con signo opuesto a π_i –, porque ninguna de las dos por sí sola basta para excluir un puente crítico.

3.2.2. La formulación QUBO de cinco términos

Con la red filtrada, las comunidades $\{\mathcal{C}_k\}$, las periferalidades π_i y las intermediaciones b_i en la mano, la selección de N_{sel} activos se formula como el siguiente problema QUBO sobre variables binarias $x_i \in \{0, 1\}$ (“el activo i entra o no entra”), exactamente el que implementa `PortfolioSelectionQUBO.build`:

$$\begin{aligned} \text{QUBO}(\mathbf{x}) = & \lambda_\rho \sum_{1 \leq i < j \leq M} \rho_{ij} x_i x_j - \lambda_\pi \sum_{i=1}^M \pi_i x_i + \lambda_b \sum_{i=1}^M b_i x_i \\ & + P_{\text{card}} \left(\sum_{i=1}^M x_i - N_{\text{sel}} \right)^2 + P_{\text{comm}} \sum_{k: |\mathcal{C}_k| > M_c} \left(\sum_{i \in \mathcal{C}_k} x_i - M_c \right)^2. \end{aligned} \quad (3.3)$$

La Tabla 3.1 resume los cinco términos y su rol. El término (v) es un *tope máximo*, no una exigencia de representación mínima: al expandir el cuadrado también penaliza quedarse por debajo de M_c en una comunidad grande, pero su calibración (Sección 3.6) está pensada para disuadir la sobreconcentración, el modo de fallo relevante. Expandiendo los cuadrados y usando $x_i^2 = x_i$ para variables binarias, todo se reduce a un polinomio cuadrático $\mathbf{x}^\top \mathbf{Q} \mathbf{x} + \text{const}$, el formato que aceptan tanto el *annealing* simulado como los solvers de D-Wave (Capítulo 2).

Término	Parámetro	Rol
(i) Correlación intra-cartera	$\lambda_\rho = 1,0$	Desincentiva pares correlacionados (diversificación)
(ii) Periferalidad	$\lambda_\pi = 0,5$ (signo $-$)	Premia periferalidad agregada alta al minimizar
(iii) Intermediación (puentes)	$\lambda_b = 0,3$	Penaliza nodos de alta intermediación (contagio)
(iv) Cardinalidad	$P_{\text{card}} = 500$	Cuadrado penalizado forzando $\sum_i x_i = N_{\text{sel}}$
(v) Concentración por comunidad	$P_{\text{comm}} = 100, M_c = 3$	Tope por comunidad, solo si $ \mathcal{C}_k > M_c$

Tabla 3.1: Rol de los cinco términos de la Ec. (3.3), $N_{\text{sel}} = 10$.

3.2.3. El resultado real: la cartera PA

Sobre el universo real de $M = 374$ activos, con $N_{\text{sel}} = 10$, $K = 11$ y los pesos de (3.3) en sus valores reales ($\lambda_\rho = 1,0$, $\lambda_\pi = 0,5$, $\lambda_b = 0,3$, $P_{\text{card}} = 500$, $P_{\text{comm}} = 100$, $M_c = 3$), resuelto con *annealing* simulado (`neal`, semilla 42; el solver híbrido de D-Wave existe como alternativa opcional y no se usó), el procedimiento produjo la cartera **PA**: AAL, AAP, CBOE, CHRW, CLX, DXCM, FMC, NEM, NFLX, OXY, con periferalidad media anotada 0,980 y 10 de las 11 comunidades cubiertas. Las carteras de comparación son **PB** –los diez activos de mayor centralidad, justo el objetivo opuesto: AIG, AMP, BAC, BRK-B, ITW, IVZ, L, MCO, PRU, TROW; periferalidad 0,474, 4/11 comunidades– y **PC** –*mega-cap* manual, sin criterio de red: AAPL, MSFT, GOOGL, JPM, UNH, PG, XOM, CAT, NEE, LIN; periferalidad 0,720, 8/11–.

El artefacto de solver de esa selección **no se conservó**: ningún fichero de `results/` registra periferalidad, cobertura ni energía QUBO, así que esas cifras se reportan tal como quedaron anotadas,

sin poder verificarse. Una reejecución determinista sobre la vintage actual –364 activos, TMFG de 1 086 aristas, $K = 13$ – no reproduce los universos congelados, porque los retornos proceden de una fuente no versionada que recalcula retroactivamente los precios ajustados en cada descarga: la selección periférica regenerada comparte solo 4 de 10 tickers con PA, y la de centralidad 5 de 10 con PB. **Ninguno de los 46 experimentos depende de ese artefacto**, porque sus listas están congeladas en los `config/*.yaml`; todos los resultados están condicionados a esas listas y no a un procedimiento reejecutable –salvedad material, dado que el Capítulo 6 identifica la selección QUBO como el motor dominante del rendimiento–.

3.3. Etapa 2: estimación del subgradiente de CVaR vía IQAE

Qué necesita exactamente la Etapa 3

Fijados los N activos de la Etapa 1, el optimizador de la Etapa 3 necesita, en cada iteración t , el vector $\hat{g}(w_t) = [\hat{g}_1, \dots, \hat{g}_N]^\top$ de esperanzas condicionadas sobre la cola de pérdidas al que se reduce el subgradiente de $\text{CVaR}_\alpha(w_t)$ (Capítulo 2). Esta sección explica cómo se traduce esa esperanza condicionada en un circuito cuántico concreto, de solo cuatro qubits de distribución más un ancilla, que no crece con N .

Discretización de la distribución de pérdidas

Un ordenador cuántico de puertas no puede cargar una distribución continua de pérdidas directamente: hace falta discretizarla primero en 2^n niveles, donde n es el número de qubits del registro. Este trabajo usa $n = 4$ qubits, es decir, $2^4 = 16$ niveles de pérdida discretizados. A partir de los retornos de entrenamiento se calculan las probabilidades empíricas p_ℓ de que la pérdida de la cartera caiga en cada uno de esos 16 contenedores, y se construye un circuito de preparación de estado $\mathcal{A}$ tal que

$$\mathcal{A}|0\rangle^{\otimes n} = \sum_{\ell=0}^{2^n-1} \sqrt{p_\ell} |\ell\rangle. \quad (3.4)$$

Este es el mismo tipo de circuito de carga de distribución que el Capítulo 2 introduce al presentar QAE (la familia de algoritmos cuánticos de la que IQAE es la variante iterativa); aquí se particulariza a la distribución empírica de pérdidas de la cartera concreta bajo evaluación en la iteración t .

Por qué $n = 4$ y no más se justifica con detalle en la Sección 3.6. En síntesis: el circuito no codifica un qubit por activo, sino un único registro para la pérdida *total* de la cartera, agregada según w_t ; el número de activos entra en la *construcción clásica* de p_ℓ , no en el número de qubits, lo que mantiene el circuito de anchura fija incluso cuando N crece de 10 a 100.

El oráculo de cola y la probabilidad de cola

El primer objeto que hace falta estimar es la probabilidad de cola $a_\tau = \mathbb{P}(L_{\text{disc}} > \hat{\tau})$, donde $\hat{\tau}$ es el umbral que aproxima el valor en riesgo, VaR (el cuantil de pérdida que CVaR promedia más allá de sí mismo, definido con precisión en el Capítulo 1), al nivel α dentro de la discretización de 16 niveles (el mayor τ tal que la masa de probabilidad por encima de τ no exceda α). Un oráculo de Grover $\mathcal{O}_\tau$ marca con una fase (o, equivalentemente, activa un qubit ancilla) los estados $|\ell\rangle$ con $\ell > \hat{\tau}$; ejecutar IQAE —la variante *iterativa* de QAE introducida por Grinko *et al.* [16] y derivada con detalle en el Capítulo 2— sobre el circuito compuesto $\mathcal{AO}_\tau$ produce una estimación $\hat{a}_\tau$ de esa probabilidad de cola, con precisión objetivo ε y probabilidad de fallo δ acotadas de antemano.

3.3.1. CVaR marginal por activo: esperanza condicionada codificada directamente

La probabilidad de cola por sí sola no basta: hace falta, para cada activo i , su pérdida esperada condicionada a estar en la cola, $\mathbb{E}[r_i \mid L > \hat{\tau}]$. La forma más directa de plantear esto sobre papel sería construir un circuito de carga *conjunta* $\mathcal{A}_i$ que codificara $\mathbb{E}[r_i \mathbf{1}_{L>\tau}]$ sobre el mismo oráculo de cola $\mathcal{O}_\tau$ de la subsección anterior, y dividir la estimación resultante $\hat{c}_i$ por $\hat{a}_\tau$ para recuperar el valor condicionado —un *estimador de razón*, $\hat{g}_i = -\hat{c}_i/\hat{a}_\tau$ —. Ese no es, sin embargo, el mecanismo que implementa el código: en la ruta que realmente se ejecuta, la cola se selecciona *clásicamente* antes de construir ningún circuito cuántico. Dado el umbral $\hat{\tau}$ (en la práctica, el VaR clásico de la Sección 1.4), se filtran las observaciones históricas de pérdida para quedarse solo con los días de cola, $\{t : L^{(t)} \geq \hat{\tau}\}$, y se discretiza —con el mismo esquema de 16 niveles— únicamente el retorno del activo i *en esos días*, no en la muestra completa. El circuito de carga $\mathcal{A}_i$ resultante codifica, por construcción, la distribución ya *condicionada* a la cola; ejecutar IQAE sobre $\mathcal{A}_i$ solo —sin oráculo $\mathcal{O}_\tau$ en este segundo circuito, porque el condicionamiento ya sucedió en el preprocesado clásico— produce directamente una estimación $\hat{e}_i \approx \mathbb{E}[r_i \mid L \geq \hat{\tau}]$, y el subgradiente por activo se recupera sin ninguna división:

$$\hat{g}_i = -\hat{e}_i. \quad (3.5)$$

(`build_conditional_expectation_circuit` en `qae_circuits.py` construye este circuito directamente a partir de la máscara de cola clásica: la esperanza condicionada sale ya condicionada del circuito, así que la fórmula del subgradiente no necesita dividir por $P(\text{cola})$ en ningún punto de la ruta que produjo las 46 corridas archivadas.) Esto cuesta $N + 1$ llamadas a IQAE por cada evaluación completa del subgradiente —una para $\hat{a}_\tau$ y una por cada activo—, coste que escala linealmente con N a través del número de *llamadas* al mismo circuito de cuatro qubits más un ancilla, no a través de su *tamaño*. Esta distinción, coste lineal en invocaciones frente a coste lineal en anchura, es la que permite llegar a carteras de hasta 100 activos (familia S1, Sección 3.5) sin rediseñar el circuito cuántico.

3.3.2. Profundidad de circuito: la sonda de hardware frente al circuito real del pipeline

Cada llamada a IQAE ejecuta un número variable de rondas m del operador de Grover, y cada ronda añade profundidad al circuito; conviene separar dos profundidades fáciles de confundir. El circuito que se acaba de describir –preparación de estado más oráculo de umbral, sin ninguna ronda de amplificación todavía– ya es bastante profundo: su profundidad pre-transpilación archivada en $m = 0$ es 44 puertas para el circuito de probabilidad de cola y 47 para los de esperanza condicionada (75–84 en total). Transpilado *offline* a puertas nativas de IQM (`transpile(..., optimization_level=1)`, la misma llamada que usan los guiones de hardware), alcanza una profundidad de 1,476 con 782 puertas CZ en $m = 0$, y 5,395 en $m = 1$: una brecha de más de dos órdenes de magnitud frente a los valores que siguen. Este circuito **nunca se ha ejecutado en hardware IQM, a ningún m** : las 46 corridas de optimización lo ejecutan siempre en el simulador `qiskit-aer`.

La caracterización en hardware IQM del Capítulo 4 usa, en su lugar, una *sonda* mínima de dos qubits que reproduce la misma amplitud de cola codificada –la Sección 4.4.1 de ese capítulo justifica esa sustitución–. Para esa sonda, tras compilar a puertas nativas de IQM (`opt_level=1`), la profundidad observada sí sigue la fórmula empírica

$$\text{depth}(m) = D_{\mathcal{A}} + 6m, \quad (3.6)$$

con $D_{\mathcal{A}} = 4$. La Ecuación (3.6) es la ley de la *sonda*, no la del circuito de preparación de estado de cuatro qubits de distribución más un ancilla que se acaba de describir. El Capítulo 4 deriva esta fórmula con más detalle y la usa como pieza central del análisis del umbral de *aliasing* en hardware IQM (Contribución C1, Sección 1.5); aquí basta con retener que cada ronda adicional de Grover cuesta, en la sonda, seis puertas nativas más.

3.3.3. Ejemplo numérico trabajado: la cartera PA real

A diferencia del QUBO de la Etapa 1, para la Etapa 2 sí existe un artefacto real y trazable en el repositorio: `results/exp_A1_PA/tfm_comprehensive_metrics_latest.json`, que registra una ejecución completa del pipeline sobre la cartera PA (los diez tickers de la Sección 3.2.3) con pesos iniciales iguales $w_0 = (0,1,\dots,0,1)$. En w_0 , la estimación de la probabilidad de cola es exactamente el paso “*Tail probability*” del algoritmo de la Sección 3.3, con un número concreto de consultas en lugar de la cota asintótica, y la norma cuántica del subgradiente completo –única evaluación del coseno que el pipeline persiste por experimento– se contrasta directamente contra la descomposición de Euler clásica sin ruido (Ec. (3.5)), cuyo vector por activo recoge la Tabla A.2.

Sus cifras clave son: $\alpha = 0,05$, $n = 4$ qubits, $\varepsilon = 0,005$ y $T = 1,005$ días; una cola al 5 % de 51 observaciones ($\approx 0,05075$) con $\text{VaR}_{0,05}(\mathbf{w}_0) = 0,02205$; una estimación $\hat{a}_\tau = 0,049565$ frente al valor clásico exacto $0,050746$ —un error del 2,33 %— obtenida con 11,264 consultas al oráculo; y una norma cuántica del subgradiente $\|\hat{\mathbf{g}}(\mathbf{w}_0)\|_2 = 0,13039$ con coseno 0,99991 frente a la referencia, contra la norma clásica sin ruido de la descomposición de Euler, $\approx 0,1285$. Las 38 iteraciones de Adam suman 4,436,992 consultas, frente a las 2,432 de la referencia no iterativa de Woerner [1].

Ambas normas son del mismo orden de magnitud (diferencia relativa $\approx 1,4$ %, consistente con el 2,33 % de error por componente). Lo que importa para la Etapa 3 es el *ranking*, no el valor exacto: bajo pesos iguales OXY (0,0647) y AAL (0,0578) tienen la mayor contribución marginal y CLX (0,0117) la menor —vector completo en la Tabla A.2—, y ese orden gobierna hacia dónde mueve los pesos Adam.

3.4. Etapa 3: optimización con Adam

Por qué Adam y no descenso de subgradiente simple

Un descenso de subgradiente clásico actualizaría $\mathbf{w}_{t+1} = \mathbf{w}_t - \eta \hat{\mathbf{g}}_t$, usando directamente la magnitud estimada. Pero el ruido de $\hat{\mathbf{g}}_t$ es *multiplicativo*, no aditivo —el error crece proporcionalmente a la cantidad estimada, derivación en el Capítulo 4—, y bajo ruido multiplicativo un método que use la magnitud cruda lo amplifica justo donde el gradiente es mayor, que es donde más importa que la dirección sea fiable. Adam [27] normaliza cada componente por una estimación de su propia desviación reciente, de modo que la inflación de magnitud queda absorbida por el denominador y la dirección del paso la gobiernan el signo y el orden relativo de las componentes, no su escala: exactamente la propiedad que cuantifica la desigualdad de brecha de ranking del Capítulo 4. No es, por tanto, una elección de “optimizador popular”.

3.4.1. La reparametrización softmax: por qué no se proyecta

Antes de escribir la actualización de Adam hay que decidir *sobre qué variables* se optimiza. La opción aparentemente natural —dar el paso de gradiente en el espacio de pesos $\mathbf{w}$ y devolver el resultado al conjunto admisible $\mathcal{W} = \{\mathbf{w} \geq \mathbf{0}, \mathbf{w}^\top \mathbf{1} = 1, w_i \leq 0,3 \forall i\}$ mediante alguna operación de proyección o renormalización— esconde una trampa que merece entenderse, porque es la razón de ser del diseño real.

Considérese el punto de partida de todas las ejecuciones: pesos iguales, $\mathbf{w}_0 = (1/N, \dots, 1/N)$. El subgradiente de CVaR tiene ahí todas sus componentes positivas —cada activo *añade* riesgo de cola—, y la normalización de magnitud de Adam convierte cada componente del paso en ± 1 casi exacto, porque la corrección de sesgo cancela la escala (Sección 3.4.2). Con todas las componentes del mismo signo el paso crudo es *uniforme*, $\tilde{w}_i = 1/N - \eta$ para todo i , y cualquier operación que lo

devuelva al símplex renormalizando lo lleva de vuelta exactamente a w_0 . El punto de partida es un **punto fijo espurio** de la dinámica proyectada —espurio porque no es un óptimo: el gradiente contiene información direccional real que la combinación “normalización de Adam + renormalización” destruye—. Y no es un accidente de la primera iteración: como el iterante no se mueve, Adam ve el mismo gradiente en todas, el cociente vale $\text{sign}(\mathbf{g})$ siempre y el punto fijo lo es de toda la trayectoria.

La solución implementada es no proyectar en absoluto: *reparametrizar*. Los pesos se expresan como la función softmax de un vector de parámetros libres $\theta \in \mathbb{R}^N$ (SoftmaxParameterization, activa en todas las ejecuciones reportadas vía el valor por defecto `use_softmax=True`, que ninguna configuración ni cuaderno de reproducción modifica):

$$w_i = \frac{e^{\theta_i/\tau}}{\sum_{j=1}^N e^{\theta_j/\tau}}, \quad \tau = 1. \quad (3.7)$$

Sea cual sea θ , el vector w resultante es automáticamente positivo y suma uno: la optimización sobre θ es *sin restricciones*, y las restricciones de cartera se cumplen por construcción, no por corrección a posteriori. El gradiente respecto a θ se obtiene por regla de la cadena:

$$\mathbf{u}_t = \nabla_{\theta} \text{CVaR} = \frac{1}{\tau} (\mathbf{w} \odot \hat{\mathbf{g}}_t - (\mathbf{w}^\top \hat{\mathbf{g}}_t) \mathbf{w}), \quad \text{es decir,} \quad u_{t,i} = \frac{w_i}{\tau} (\hat{g}_{t,i} - \mathbf{w}^\top \hat{\mathbf{g}}_t). \quad (3.8)$$

La forma por componentes es reveladora: lo que mueve a θ_i no es el valor absoluto de $\hat{g}_{t,i}$, sino su *desviación respecto a la media ponderada de la cartera* $\mathbf{w}^\top \hat{\mathbf{g}}_t$. Una componente uniforme del gradiente —la que el enfoque proyectado convertía en un paso inútil— aquí se anula sola en la resta, y lo que queda es exactamente la información diferencial entre activos. El punto fijo espurio desaparece: en w_0 con gradiente no uniforme, (3.8) es distinto de cero y la primera iteración ya diferencia activos, como se comprueba con datos reales en la Sección 3.4.2. (Un gradiente *exactamente* uniforme sí daría $\mathbf{u}_t = \mathbf{0}$ —y debe darlo: sobre el símplex, un gradiente uniforme no contiene ninguna preferencia entre activos, y la softmax es invariante a desplazamientos uniformes de θ —.)

Dos notas de honestidad sobre el alcance de este diseño: la cota superior $w_i \leq 0,3$ no la garantiza la softmax, sino un único paso de recorte-y-renormalización aplicado tras cada actualización que resincroniza θ cuando algún peso excede el tope; y el código conserva una ruta alternativa con actualización proyectada (`use_softmax=False`), cuyo operador `_project_weights` es igualmente un recorte-y-renormalización, no la proyección euclídea al símplex, pero ninguna ejecución reportada la usa: el mecanismo real es, en todos los casos, (3.7)–(3.8).

La actualización de Adam paso a paso

En cada iteración t , tras obtener $\hat{\mathbf{g}}_t \leftarrow \text{IQAE}(\mathbf{w}_{t-1})$ mediante el procedimiento de la Sección 3.3 y transformarlo a $\mathbf{u}_t$ vía (3.8), Adam mantiene dos promedios móviles *en el espacio de parámetros* —el primer momento $\mathbf{m}_t$ (una media exponencial del gradiente) y el segundo momento $\mathbf{v}_t$ (una media

exponencial del gradiente al cuadrado, componente a componente)–:

$$\mathbf{m}_t = \beta_1 \mathbf{m}_{t-1} + (1 - \beta_1) \mathbf{u}_t, \quad (3.9)$$

$$\mathbf{v}_t = \beta_2 \mathbf{v}_{t-1} + (1 - \beta_2) \mathbf{u}_t^{\odot 2}. \quad (3.10)$$

Como $\mathbf{m}_0 = \mathbf{v}_0 = \mathbf{0}$, ambos promedios están sesgados hacia cero en las primeras iteraciones; Adam corrige ese sesgo dividiendo por $(1 - \beta_1^t)$ y $(1 - \beta_2^t)$ respectivamente, y el paso final actualiza los parámetros y recupera los pesos con la softmax (líneas 726–747):

$$\boldsymbol{\theta}_t = \boldsymbol{\theta}_{t-1} - \eta_t \frac{\mathbf{m}_t / (1 - \beta_1^t)}{\sqrt{\mathbf{v}_t / (1 - \beta_2^t) + \varepsilon_A}}, \quad \mathbf{w}_t = \text{softmax}(\boldsymbol{\theta}_t / \tau), \quad (3.11)$$

con tasa de aprendizaje decreciente $\eta_{t+1} = 0,995 \eta_t$ y, si hace falta, el paso de tope $w_i \leq 0,3$ ya descrito en la Sección 3.4.1. El denominador $\sqrt{\mathbf{v}_t / (1 - \beta_2^t) + \varepsilon_A}$ es la estimación de la desviación estándar reciente de cada componente del gradiente: dividir por ella normaliza la magnitud y deja, aproximadamente, solo el signo de cada componente gobernando el paso –el mecanismo que hace a Adam robusto frente al ruido multiplicativo de IQAE, y que opera aquí sobre $\mathbf{u}_t$, donde ese ruido entra como el mismo factor común porque (3.8) es lineal en $\hat{\mathbf{g}}_t$ –.

3.4.2. El primer paso real de PA, y una convención de contabilidad

Retomando el ejemplo de la Sección 3.3.3, con $\mathbf{w}_0 = (0,1, \dots, 0,1)$ y $\boldsymbol{\theta}_0 = \mathbf{0}$: la media ponderada de (3.8) es la media aritmética del vector de CVaR marginal, $\mathbf{w}_0^\top \mathbf{g} = 0,03726$, y cada componente de $\mathbf{u}_1$ vale $0,1 (g_i - 0,03726)$, cuyo signo parte el universo en cuatro activos por encima de la media –el mayor, OXY– y seis por debajo. La componente uniforme ha desaparecido en la resta. Con $\mathbf{m}_0 = \mathbf{v}_0 = \mathbf{0}$ la corrección de sesgo cancela exactamente la escala ($\hat{m}_1 = u_1$, $\sqrt{\hat{v}_1} = |u_1|$), así que el cociente de Adam vale $\text{sign}(u_1)$ y cada coordenada recibe un paso de $\pm \eta = \pm 0,05$; la softmax devuelve $w_{1,i} = 0,1040$ para los seis defensivos y $0,0941$ para los cuatro arriesgados. La primera iteración ya diferencia activos, en la dirección correcta, y bajo la ruta proyectada habría sido exactamente nula.

El historial archivado lo confirma: `cvar_history` registra $0,037232, 0,037232, 0,036538, \dots$ –las dos primeras entradas coinciden por convención de registro, no porque la cartera no se moviera–, y la expansión de primer orden $\text{CVaR}(\mathbf{w}_1) \approx \text{CVaR}(\mathbf{w}_0) + \mathbf{g}^\top (\mathbf{w}_1 - \mathbf{w}_0)$ predice $0,03653$ frente al $0,036538$ registrado. Desde ahí la cartera converge en 38 iteraciones a $\text{CVaR} = 0,027694$ –una reducción del 25,6%–, concentrándose en los activos de menor CVaR marginal inicial (CBOE $0,1 \rightarrow 0,290$, CLX $0,1 \rightarrow 0,279$, CHRW $0,1 \rightarrow 0,178$) y vaciando los de mayor (OXY $0,1 \rightarrow 0,016$, AAL y NFLX $0,1 \rightarrow 0,019$).

Un matiz de contabilidad explica la pareja $0,02769/0,02768$ que reaparece en los Capítulos 4, 5 y A: como el CVaR se evalúa *antes* de la actualización y los pesos se sobrescriben *después*, la comprobación de mejora guarda el par $(\text{CVaR}(\mathbf{w}_{t-1}), \mathbf{w}_t)$, de modo que el valor reportado ($0,027694$) es el del penúltimo iterado y los pesos son los del último, cuyo CVaR recalculado es $0,027681$. El desfase –como

mucho $7,4 \times 10^{-5}$ en toda la batería— tiene signo definido: negativo cuando la última iteración fue la de mejor CVaR, positivo cuando el bucle siguió sin mejorar, y exactamente cero en B2_PA/B2_n6_PA. El valor recalculado a partir de los pesos reportados es el canónico. El código se corrigió el 24 de agosto de 2026, después de archivar los 46 experimentos, así que la corrección **no es retroactiva**.

3.4.3. Criterios de convergencia

El bucle tiene *tres* vías de terminación, no una, porque cubren regímenes de fallo distintos. La primera es el criterio de optimalidad de primer orden: $\|\nabla_{\theta} \text{CVaR}\|_2 < \delta_c = 10^{-5}$, que detecta la optimalidad “limpia” pero que bajo ruido de IQAE puede no dispararse nunca, porque el gradiente estimado rara vez se anula del todo. La segunda es una ventana relativa: cambio máximo de CVaR por debajo de 10^{-3} en cinco iteraciones consecutivas, que detecta la meseta práctica con independencia del gradiente. La tercera es la paciencia, $P = 20$ iteraciones sin mejorar el mejor valor visto —no “pasos pequeños consecutivos”—, red de seguridad frente a dinámica no productiva. A ellas se suma un límite duro de iteraciones, variable por familia (50 con $N_{\text{sel}} = 10, 80$ y 100 en las corridas de escalabilidad), que no es meramente teórico: S1_PB_100 y S1_PC_100 lo agotan sin declarar convergencia.

3.5. Diseño experimental: qué varía entre las 46 configuraciones

Con las tres etapas fijadas, lo único que varía entre experimentos son cuatro ejes: el universo de partida (PA, PB o PC), el nivel de ruido simulado, los factores de escala de ZNE y la precisión objetivo ε , más ablaciones del número de qubits n_q en tres familias. La taxonomía completa de las nueve familias resultantes, con el desglose que suma las 46 corridas y el racional de qué aísla cada una, se presenta junto a sus resultados en la Tabla 4.1 del Capítulo 4, para no separar el diseño de las cifras que produce.

3.6. Justificación de hiperparámetros

La configuración completa de los 46 experimentos —datos, selección de universo vía QUBO, circuito IQAE, optimizador Adam, simulación de ruido y *bootstrap* de *backtesting*— se recoge en la Tabla A.3 del Capítulo A, verificada fila a fila contra los `config/*.yaml` del repositorio y contra `network_portfolio_selector.py` y `hybridoptimizer.py`, con una corrección respecto a la Tabla 4 del Apéndice F del paper: esa tabla omite la calibración de penalización de comunidad para $N_{\text{sel}} = 100$ ($P_{\text{comm}} = 10,000$, tope $M_c = 15$) que el código sí aplica. La Tabla 3.2 razona las decisiones de diseño más importantes de esa configuración.

Decisión	Justificación
$\varepsilon = 0,005$ (A1, B2, C2, E1–E3)	Estricto para que la escala $O(\varepsilon^{-1})$ suponga ventaja teórica frente al $O(\varepsilon^{-2})$ de Monte Carlo, sin disparar las consultas al oráculo hasta hacer inviable simular 46 experimentos. Los demás valores aíslan efectos: $\varepsilon = 0,002$ (A3) estrecha la precisión sin otro cambio; $\varepsilon = 0,010$ (D1/D2) comprueba si relajarla compensa el coste de ZNE agresivo; $\varepsilon = 0,001$ se reserva a A4_PA, fuera de la barrida y sobre una ventana más larga ($T = 1,508$ frente a 1,005 días), por lo que no es una comparación controlada de “solo ε ” (Sección 6.3).
$n = 4$ qubits, fijo	Fija $2^n = 16$ niveles de discretización de la pérdida: equilibra resolución y profundidad, pues cada qubit extra encarece la preparación de estado $\mathcal{A}$ (Ec. (3.6)). Mantenerlo fijo para $N \in \{10, 30, 100\}$ es la decisión de diseño más importante del pipeline –evita rediseñar el circuito al cambiar el tamaño de cartera, a diferencia de la literatura comparada (Tabla 1.1)–. El barrido $n_q \in \{3, 5, 6\}$ de A1 mide el coste de apartarse de ella.
Adam frente a SGD	La normalización de magnitud por componente es lo que lo hace preferible bajo el ruido <i>multiplicativo</i> de IQAE (Sección 3.4), no una preferencia por el “optimizador popular”.
$\beta_1=0,9$, $\beta_2=0,999$, $\varepsilon_A=10^{-8}$, $\eta=0,05$	Los tres primeros son los valores por defecto de Adam [27]; solo η se aparta del suyo (0,001) hacia un paso mayor, razonable con 50–100 iteraciones por corrida y un subgradiente ya acotado al simplex. No reajustar el resto evita un grado de libertad que enturbiaría atribuir efectos al ruido cuántico.
Cota de peso $u_i = 0,3$	Sin ella nada impide concentrar toda la cartera en el activo de menor CVaR marginal. Tres veces $1/N_{\text{sel}}$ para $N_{\text{sel}} = 10$: laxa para no forzar una solución casi uniforme, estricta frente a concentración extrema. En PA no llega a activarse (peso máximo 0,290 en CBOE, a 0,010 del límite): actúa cerca de su régimen activo, no como salvaguarda inerte.
$\delta_c = 10^{-5}$, ventana relativa, paciencia $P = 20$	Cubren regímenes de fallo distintos, por eso coexisten: bajo ruido de IQAE el gradiente estimado rara vez se anula, así que δ_c por sí solo podría no dispararse nunca; la ventana relativa detecta la meseta práctica; la paciencia es la red de seguridad frente a dinámica no productiva. El límite de iteraciones no es teórico: S1_PB_100 y S1_PC_100 lo agotan sin converger, indicio de que para N_{sel} grande el presupuesto de iteraciones puede ser el factor limitante.
$P_{\text{card}} = 500$, $P_{\text{comm}} = 100$	La cardinalidad debe dominar el rango completo del objetivo blando: la contribución máxima del término de correlación es $\lambda_\rho N_{\text{sel}}(N_{\text{sel}} - 1)/2 \approx 45$ para $N_{\text{sel}} = 10$, así que 500 da margen de $11\times$; para $N_{\text{sel}} = 100$ (rango $\approx 4,950$) se recalibra a 100,000 y $P_{\text{comm}} = 10,000$ con $M_c = 15$. La penalización de comunidad es deliberadamente menor: preferencia estructural, no igualdad dura. Si aun así el <i>annealer</i> devolviera $K \pm \delta$ activos, el guion aplica reparación voraz por rango de perifericidad.
4,096 disparos; $p_{1q} = p_{2q}/10$	4,096 es el compromiso estándar en QAE: el ruido de muestreo no domina sobre el error sistemático que se quiere medir, sin encarecer 46 simulaciones. En hardware los disparos varían por bloque (2,048 en el barrido de Grover y en la oleada 2; 4,096 en preparación de estado, bloques ZNE y bloque G; 8,192 en IQAE oleada 1 y Garnet tras la oleada 3; detalle en la Sección 4.4.3). La proporción 10:1 refleja que las puertas de un qubit son típicamente un orden de magnitud más fiables que las de dos, y son estas las que crecen con cada ronda de Grover.
<i>Bootstrap</i> de bloque circular, $B = 5,000$, bloques de 21 días	El remuestreo observación a observación rompería la autocorrelación de los retornos y daría intervalos artificialmente estrechos. El bloque circular [28] preserva bloques contiguos de un mes bursátil, y $B = 5,000$ es convencional para estabilizar percentiles. Es el procedimiento del $p \approx 0,12$ del Capítulo 5.

Tabla 3.2: Razón de ser de las decisiones de diseño más importantes de la batería de 46 experimentos.

RESULTADOS EXPERIMENTALES

Este capítulo responde con datos a las tres preguntas de la Sección 1.3: cuántas rondas de Grover aguanta el hardware IQM antes de que el ruido corrompa la estimación (Q1), si el orden relativo entre componentes del subgradiente sobrevive al ruido (Q2), y si ZNE mejora la precisión de IQAE o su coste de profundidad anula la ganancia (Q3). Las secciones de simulación (4.2 y 4.3) preceden a las de hardware físico (4.4) porque las primeras fijan el universo y el punto de operación sobre los que las segundas miden.

4.1. Diseño experimental: la batería de 46 experimentos

4.1.1. Universos de cartera y datos

Todos los experimentos parten de la misma base de datos: precios de cierre ajustados diarios de $M = 374$ componentes del S&P 500, de enero de 2020 a diciembre de 2025 (1 506 sesiones), excluyendo de antemano los activos con algún dato faltante.¹ El periodo se divide en entrenamiento (2020-01-02 a 2023-12-29, $T = 1\,005$ sesiones) y fuera de muestra (2024-01-02 a 2025-12-30, $T_{\text{oos}} = 501$). El nivel de CVaR es $\alpha = 0,05$ en todo el capítulo, salvo donde se indique.

Sobre esa base, la Etapa 1 construye tres universos de diez activos pensados para aislar el efecto de la *estructura* del universo, no solo su tamaño:

PA (periférico vía QUBO). Solución del problema QUBO de selección (Capítulo 3) sobre la red de correlaciones del periodo de entrenamiento filtrada mediante un grafo filtrado triangulado maximal (TMFG; el árbol de expansión mínima, MST, es solo el andamio de esa construcción, Sección 3.2.1). Activos: AAL, AAP, CBOE, CHRW, CLX, DXCM, FMC, NEM, NFLX, OXY; cubre 10 de las 11 comunidades de mercado del momento de la selección.

¹ $M = 374$ (y $K = 11$ comunidades) refleja la topología vigente cuando se congelaron los universos. El feed de precios ajustados de Yahoo Finance no está versionado `-auto_adjust=True` recalcula toda la serie en cada descarga-, así que repetir hoy la descarga da $M = 364$ y $K = 13$. Ambas topologías son internamente consistentes sobre su propia vintage; la discrepancia es de procedencia del feed, no de cómputo, y no afecta a las listas congeladas ni a ningún resultado. Su efecto sobre la pareja 0,02768/0,02769 de la Sección 4.2.2 es de orden 10^{-8} .

PB (centralidad máxima). Los diez activos con mayor centralidad de red —el objetivo opuesto al de QUBO—. Activos: AIG, AMP, BAC, BRK-B, ITW, IVZ, L, MCO, PRU, TROW; 4 de 11 comunidades.

PC (mega-cap manual). Cesta representativa de grandes capitalizaciones, construida sin ningún criterio de red: AAPL, MSFT, GOOGL, JPM, UNH, PG, XOM, CAT, NEE, LIN; 8 de 11 comunidades.

Procedencia de la selección de la Etapa 1.

Los recuentos de cobertura de comunidades anteriores, y las periferalidades medias registradas al seleccionar los universos (0,980 para PA, 0,474 para PB, 0,720 para PC), son valores anotados en su momento cuyo artefacto de solver no se archivó —ningún fichero de `results/` los contiene, y cada JSON registra la etapa de selección como desactivada— y no pueden volver a verificarse. El único artefacto archivado de la Etapa 1 es una reejecución sobre la vintage de datos actual (`selection_report_20260822_123534.json`, *annealing* simulado): 364 activos, TMFG con 1 086 aristas, $K = 13$ comunidades. Su selección periférica tiene periferalidad media 0,990 y cubre 10 de 13 comunidades, pero comparte solo 4 de 10 tickers con la PA congelada (y 5 de 10 en el caso de PB): los universos congelados *no* son reproducibles desde la vintage actual. Es una salvedad material, porque el Capítulo 6 identifica la selección QUBO como el motor dominante del rendimiento: todos los resultados de este capítulo y del siguiente están condicionados a las listas de tickers congeladas —especificadas arriba y en `config/config*_{PA,PB,PC}.yaml`—, no a un procedimiento de selección reejecutable con el mismo resultado.

PA y PB son los extremos deliberados de un mismo eje —periferalidad máxima frente a centralidad máxima, ambos vía la misma metodología de red—; PC es un control externo, construido sin la maquinaria de QUBO sobre la red TMFG, que comprueba que las propiedades de PA no son un artefacto de compararlas solo contra su opuesto metodológico. Esta terna es el eje sobre el que se organiza casi toda la tabla de experimentos de la Sección 4.1.2.

4.1.2. Taxonomía de familias: qué aísla cada una

La Tabla 4.1 resume las nueve familias que, combinadas con los tres universos y, en tres de ellas, con variaciones de n_q , producen las 46 corridas, y explicita qué aísla cada una.

Además de los 46 experimentos de simulador, todo experimento se contrasta contra seis optimizadores clásicos ejecutados sobre el mismo universo y los mismos datos de entrenamiento: programación secuencial por mínimos cuadrados (SLSQP), optimización restringida por aproximación lineal (COBYLA), Markowitz de varianza mínima, paridad de riesgo, descenso de subgradiente clásico, y CVaR de Monte Carlo con 10 000 escenarios. Estas líneas base clásicas son las que permiten interpretar

Familia	N	Ruido	Factores de escala ZNE λ	ε	Carteras evaluadas
A1	10	ninguno	no	0.005	PA, PB, PC
A3	10	ninguno	no	0.002	PA, PB, PC
B2	10	moderado ($p_{2q} = 10^{-3}$)	[1, 1, 5, 2]	0.005	PA, PB, PC
C2	10	alto ($p_{2q} = 5 \cdot 10^{-3}$)	[1, 1, 5, 2]	0.005	PA, PB, PC
D1	10	muy alto ($p_{2q} = 10^{-2}$)	[1, 1, 5, 2]	0.010	PA, PB, PC
D2	10	muy alto ($p_{2q} = 10^{-2}$)	[1, 1, 5, 2]	0.005	PA, PB, PC
E1	10	$p_{2q} = 10^{-3}$	[1, 3, 5]	0.005	PA, PB, PC
E2	10	$p_{2q} = 5 \cdot 10^{-3}$	[1, 3, 5]	0.005	PA, PB, PC
E3	10	$p_{2q} = 10^{-2}$	[1, 3, 5]	0.005	PA, PB, PC
S1	{10, 30, 100*}	ninguno	no	0.005	PA, PB, PC

*S1_PC_100 usa en realidad $N = 96$ activos, no 100: su fichero de configuración (`config_S1_PC_100.yaml`) lista 96 tickers, verificado en `n_assets` del JSON archivado; S1_PA_100 y S1_PB_100 sí usan 100.

Tabla 4.1: Taxonomía de los 46 experimentos de simulador (Tabla 1 del paper). Todos usan Adam con $\varepsilon = 0,005$ salvo indicación contraria; el ruido se refiere a p_{2q} . A1, B2 y D1 incluyen ablaciones de n_q sobre PA (y A1 también sobre PB/PC), que completan el recuento: A1 = 12, A3 = 3, B2 = 5, C2 = 3, D1 = 5, D2 = 3, E1/E2/E3 = 9, S1 = 6. Los bloques de hardware y de ZNE se ejecutan aparte y no cuentan aquí. Fuente: `config/config_<FAMILIA>_<CARTERA>.yaml`.

cualquier resultado cuántico en términos relativos —“¿mejora sobre SLSQP?”— en lugar de en términos absolutos sin punto de comparación; su papel como vara de medir para la validación regulatoria se retoma con detalle en el Capítulo 5.

Todos los experimentos usan una semilla fija (42), lo que hace que las diferencias entre configuraciones sean atribuibles al diseño experimental y no a variabilidad de inicialización aleatoria. El código fuente completo, los datos de retornos, las formulaciones QUBO para D-Wave, los modelos de ruido y los registros de optimización de cada experimento están disponibles en el repositorio público del proyecto (Sección 7.3), organizados en `results/exp_<EXP>/`.

4.2. Geometría del QUBO y calidad del universo seleccionado

4.2.1. Calidad del universo seleccionado

Ocho métricas resumen la calidad de cada universo. Las cuatro primeras —media y máximo de las correlaciones por pares, porcentaje de los $\binom{10}{2} = 45$ pares con $\rho > 0,5$, y cuota de varianza del primer autovalor— miden hasta qué punto el universo se comporta como un único activo replicado diez veces. La dimensión efectiva $PR = (\sum_k \lambda_k)^2 / \sum_k \lambda_k^2 \in [1, N]$ confirma esa lectura desde todo el espectro y no desde un único valor. La cobertura de comunidades mide diversificación sectorial y es el único dato

no recalculable, anotado al seleccionar los universos. La dispersión de CVaR –cociente entre la CVaR marginal máxima y mínima– es un indicador real de diversificación de cola, pero **no** la cantidad que gobierna la condición de brecha de ranking (Sección 4.3.2: esa usa la brecha *mínima* entre pares). Y el porcentaje de días de cola en que *todos* los activos caen a la vez indica si queda variación explotable por el subgradiente incluso en los peores días.

La Tabla 4.2 reúne los valores numéricos de estas ocho métricas para los tres universos.

Universo	Media $\bar{\rho}$	Máx. ρ	Pares $\rho > 0,5$	Autoval. dom.	Dim. efect.	Comunidades	Dispersión CVaR	Días cola neg.
PA (QUBO)	0.207	0.447	0%	29.6%	6.78	10/11	5.54×	25.5%
PB (Central)	0.710	0.871	100%	74.2%	1.78	4/11	2.17×	84.3%
PC (Manual)	0.477	0.780	40%	53.3%	3.14	8/11	1.96×	82.4%

Qué aísla cada familia. **A1**: línea base sin ruido ni ZNE, referencia frente a la que se mide cualquier degradación posterior. **A3**: precisión más fina ($\varepsilon = 0,002$) sobre A1, para ver si estrechar ε cambia la cartera óptima o solo la refina. **B2**: ruido moderado con factor de ZNE fraccionario; uno de sus tres experimentos, B2_PA, resulta un control negativo (Sección 4.3.3). **C2**: más ruido que B2 con factores enteros, para aislar la subida de ruido sin el factor fraccionario. **D1/D2**: ruido muy alto con el mismo factor fraccionario que B2 y ε distinto; sus ocho corridas convergen con normalidad, lo que descarta el factor fraccionario como causa del fallo de B2_PA. **E1/E2/E3**: barrido puro de p_{2q} sin ZNE, familia central para Q2. **S1**: escalabilidad ($N \in \{10, 30, 100\}$) sin ruido, única familia que varía el tamaño de cartera.

Tabla 4.2: Métricas de calidad de universo (entrenamiento 2020–2023, Tabla 2 del paper). “Comunidades” es la cobertura anotada al seleccionar, único dato no recalculable; las demás columnas se recomputan desde los retornos de las listas de tickers congeladas.

En síntesis, la selección QUBO produce una estructura genuinamente multifactorial y una cola de pérdidas donde la mayoría de los peores días no son todo-negativos –un margen de diversificación que el subgradiente de IQAE puede detectar–, pero la dispersión de CVaR que resume esa diversificación no es, como ya se ha señalado, la cantidad que controla el umbral de inversión de ranking de la Sección 4.3: esa cantidad es la brecha mínima entre pares de activos, prácticamente idéntica en ambos universos (verificado allí con cifras reales).

4.2.2. Descomposición de Euler: convergencia de cuántico y clásico

La descomposición de Euler reparte la CVaR total en contribuciones aditivas por activo, lo que permite comparar *cómo* concentran el riesgo distintos métodos, no solo cuánta CVaR alcanzan. La Tabla 4.3 resume, para PA en entrenamiento, la CVaR total, el índice de Herfindahl–Hirschman (HHI: suma de los cuadrados de las contribuciones porcentuales; vale 1 si todo el riesgo lo aporta un activo y $1/N$ si se reparte por igual), el número efectivo de activos que contribuyen ($1/\text{HHI}$) y la cuota de los tres más concentrados.

Método	CVaR	HHI	eff N	Cuota Top-3
Híbrido cuántico	0.02768	0.210	4.76	74.9%
SLSQP	0.02717	0.223	4.48	77.2%
Markowitz	0.02790	0.201	4.97	69.5%
Paridad de riesgo	0.03276	0.102	9.82	35.1%
Peso igual	0.03723	0.119	8.39	46.5%

Tabla 4.3: Métricas de concentración de riesgo de Euler para PA en entrenamiento (Tabla 3 del paper). $\text{eff } N = 1/\text{HHI}$. Fuente: `results/euler_decomposition/A1_PA_euler.json`.

La columna de CVaR del método híbrido cuántico es **0.02768** (campo `quantum_optimal.total_cvar` de `A1_PA_euler.json`), la cifra canónica de todos los análisis de descomposición de esta memoria. La métrica *in-sample* del propio optimizador, `cvar_loss= 0,02769`, es el desfase de contabilidad del “mejor-hasta-ahora” descrito en la Sección 3.4.2 –no un error de cálculo ni un efecto de la deriva de datos: el CVaR inicial de pesos iguales coincide bit a bit entre ficheros en 45 de las 46 corridas–. El desfase, como mucho $7,4 \times 10^{-5}$ (0,29 %, S1_PC_100) y nulo en B2_PA/B2_n6_PA, no afecta a ninguna conclusión; esta memoria cita 0,02768 salvo en la Tabla 5.1, que reporta 0,02769 por coherencia con las demás cifras de su fila.

Tanto el optimizador híbrido como SLSQP concentran entre el 74 % y el 77 % de la CVaR en los mismos tres activos (CLX, CBOE, CHRW), pese a partir de fuentes de gradiente y reglas de actualización distintas: dos algoritmos llegando a la misma estructura de riesgo con el mismo estimador subyacente refleja la geometría del paisaje de subgradiente de CVaR, no una coincidencia. La paridad de riesgo logra diversificación casi máxima ($\text{eff } N = 9,82$) pero paga una prima de CVaR del 18 %, lo que ilustra la tensión entre repartir el riesgo por igual y minimizar la pérdida de cola bajo sensibilidades heterogéneas.

4.3. Robustez al ruido y estabilidad del gradiente

Esta sección responde a Q2 (Sección 1.3): si el vector de subgradientes de CVaR estimado por IQAE mantiene su *dirección* bajo ruido de hardware, aunque su *magnitud* se degrade. Es, junto con la Sección 4.4, el núcleo técnico de este capítulo.

4.3.1. Similitud coseno en los 46 experimentos

La similitud coseno entre el subgradiente estimado por IQAE y el clásico exacto cuantifica directamente si el ruido invierte o preserva la dirección de actualización. Se evalúa como **una única medición por experimento en el punto de partida equiponderado** $w_0 = 1/N$ –la llamada de validación previa a la optimización–, no como promedio sobre las iteraciones: el coseno que el optimizador calcula en cada paso no se persiste, así que la preservación de la dirección *a lo largo* de la trayectoria se infiere de la convergencia y del argumento de la Sección 4.3.2, no se mide. La Tabla 4.4 resume, para cinco configuraciones representativas sobre PA, esa similitud en w_0 , el error relativo de CVaR y el HHI de la solución convergida.

El hecho llamativo es la disociación entre columnas: la similitud coseno en w_0 se mantiene por encima de 0,9998 en las cinco filas –incluida la de mayor ruido simulado, diez veces la tasa nominal de IQM–, mientras el error de CVaR crece hasta el 232,6 % en ese mismo punto. Bajo ruido alto la *magnitud* del subgradiente se distorsiona enormemente, pero su *dirección* en el punto de partida apenas se

Config.	Ruido (p_{2q})	$\cos \angle(\hat{g}, g^*)$	Error CVaR	HHI
A1 (ideal)	0	0.999915	10.7 %	0.2102
E1 (bajo)	10^{-3}	0.999963	14.9 %	0.1858
E2 (medio)	$5 \cdot 10^{-3}$	0.999954	109.8 %	0.1638
E3 (alto)	10^{-2}	0.999940	232.6 %	0.1643
D1 (ZNE)	10^{-2}	0.999871	233.5 %	0.1656

Tabla 4.4: Similitud coseno, error de CVaR y HHI de Euler para PA bajo niveles crecientes de ruido (Tabla 4 del paper). El coseno se evalúa en el punto de partida equiponderado y supera 0,9998 en los 46 experimentos.

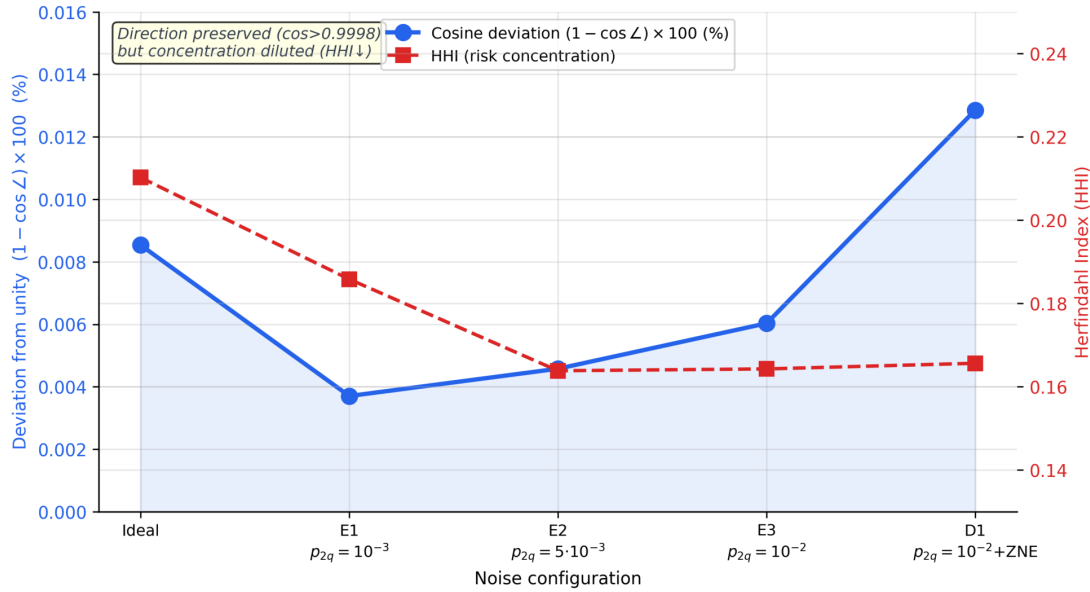


Figura 4.1: Estabilidad de la dirección del gradiente y concentración de riesgo bajo ruido creciente para las cinco configuraciones de la Tabla 4.4. Eje izquierdo (círculos): desviación coseno $(1 - \cos \angle) \times 100 \%$ en w_0 . Eje derecho (cuadrados): índice de Herfindahl de las contribuciones de riesgo. Pese a la caída del HHI, las 46 corridas satisfacen $\cos \angle > 0,9998$: la disociación entre dirección y magnitud que sustenta el argumento de brecha de ranking. Figura 5 del paper.

mueve: consecuencia estructural que desarrolla la siguiente subsección, no casualidad de estos cinco experimentos. La desviación respecto de la unidad va de $0,37 \times 10^{-2} \%$ (E1) a $1,29 \times 10^{-2} \%$ (D1) –no E3–, y recorriendo los 46 experimentos el peor caso real es todavía mayor: C2_PA alcanza $1,59 \times 10^{-2} \%$, tampoco la configuración con la tasa de error más alta. Aun así, los 46 satisfacen $\cos \angle > 0,9998$. Lo que verifican, exactamente, es que en w_0 la condición de brecha de ranking se cumple para ambos universos por un margen $\approx 1,83 \times$; en los pesos convergidos se viola necesariamente, así que este resultado es el desenlace predicho *allí donde se evalúa*, no una garantía sobre toda la trayectoria.

4.3.2. La condición de brecha de ranking

El paper resume este resultado en una única desigualdad compacta; aquí se reconstruye el argumento, distinguiendo qué parte es matemática pura –válida para cualquier implementación física de IQAE– y qué parte es un hecho empírico. La distinción importa: confundir un resultado matemático necesario con un hecho contingente del hardware sobregeneralizaría la conclusión.

Estructura del error, y cuándo invierte un orden.

IQAE, vía el estimador de razón del Capítulo 3, devuelve una estimación del subgradiente de la forma

$$\hat{g}_i = g_i^* (1 + \delta_i), \quad |\delta_i| \lesssim \varepsilon, \quad (4.1)$$

es decir, cada componente se perturba *multiplicativamente* en torno a su valor verdadero $g_i^* = -\mathbb{E}[r_i \mid L \geq \text{VaR}_\alpha]$, no aditivamente. Es la propiedad que lo distingue de Monte Carlo clásico –cuyo error es aditivo, $\hat{g}_i = g_i^* + \eta_i$, con η_i gaussiano y de magnitud independiente de g_i^* – y el origen de todo lo que sigue: (4.1) se deriva de cómo el estimador de razón convierte una probabilidad medida en una estimación de amplitud, no del hardware que la ejecute.

El optimizador necesita que el orden relativo entre g_i^* y g_j^* se preserve en la versión ruidosa; si el ruido lo invierte, se mueve en la dirección equivocada. Una inversión exige $|\delta_i g_i^* - \delta_j g_j^*| > |g_i^* - g_j^*|$, y por la desigualdad triangular junto con $|\delta_i|, |\delta_j| \leq \varepsilon$,

$$|\delta_i g_i^* - \delta_j g_j^*| \leq |\delta_i| |g_i^*| + |\delta_j| |g_j^*| \leq 2\varepsilon \max_k |g_k^*|,$$

de modo que *ningún* par puede invertir su orden siempre que

$$\min_{i \neq j} |g_i^* - g_j^*| > 2\varepsilon \max_k |g_k^*|. \quad (4.2)$$

Es una condición *suficiente*, no necesaria –puede haber realizaciones del ruido que preserven el orden sin cumplirla–, y *conservadora*, porque la cota triangular solo se alcanza con igualdad si el peor par (δ_i, δ_j) coincide con el par de menor brecha.

Por qué el ranking basta: Adam y la softmax.

Adam opera sobre el gradiente transformado por la regla de la cadena (Ec. (3.8)), lineal en $\hat{g}$ para w fijo, y divide un primer momento por la raíz de un segundo: si el gradiente se reescala por un factor común, numerador y denominador se reescalan igual y se cancelan a primer orden, así que la dirección del paso depende del signo y del orden relativo de las componentes, no de su magnitud. La cantidad cuyo orden usa Adam es $u_i = \tau^{-1} w_i (\hat{g}_i - w^\top \hat{g})$, no $\hat{g}_i$, y ambos órdenes coinciden

exactamente solo con pesos iguales, donde u es una imagen afín de $\hat{g}$; lejos de w_0 el factor w_i repondera cada componente, y la preservación de la dirección pasa de exacta a de primer orden. El mapa final a pesos tampoco introduce inversiones: la softmax preserva el orden de forma estricta y exacta a cualquier temperatura, pues $w_i/w_j = e^{(\theta_i - \theta_j)/\tau}$; el tope $\min(\cdot, 0,3)$ es débilmente monótono —puede crear empates, nunca inversiones—, y el reescalado positivo y la resincronización logarítmica son estrictamente monótonos.

Recapitulando: la desigualdad (4.2), la invariancia de Adam frente a un reescalado común y la monotonía de la softmax son **matemática pura**. La única pieza empírica es la forma multiplicativa del error, (4.1), y es también la bisagra: con un estimador de cola aditivo la cadena se rompería, porque un error aditivo no se cancela en el cociente de Adam. La robustez observada no surge de que el ruido se promedie hasta desaparecer, sino de que IQAE distorsiona g^* mediante un reescalado diagonal —que preserva el orden si la brecha lo permite— en lugar de una rotación aleatoria.

Cuándo se cumple, en números: en el punto de partida.

En el punto de partida equiponderado la condición se cumple para **ambas** carteras bajo $\varepsilon = 0,005$, aunque no por la razón que sugeriría la dispersión de CVaR (Tabla 4.2): lo que entra en (4.2) es la *brecha mínima entre pares*, no el cociente máximo/mínimo. Calculada sobre el vector de CVaR marginal de pesos iguales (`results/euler_decomposition/A1_P*_euler.json`), esa brecha es $1,18 \times 10^{-3}$ para PA (NEM—CHRW) y $1,17 \times 10^{-3}$ para PB (L—BAC), prácticamente idéntica pese a que las dispersiones son muy distintas ($5,54 \times$ frente a $2,17 \times$). Frente al lado derecho, $2\varepsilon \max_k |g_k^*| = 6,47 \times 10^{-4}$ (PA) y $6,40 \times 10^{-4}$ (PB), ambas superan el umbral por un factor casi idéntico, $\approx 1,83 \times$. Tratar el cociente de dispersión como si fuera ese margen llevaría a la conclusión contraria —margen reducido o inexistente para PB—, y los datos la descartan. La selección QUBO sigue aportando diversificación real (Sección 4.2.2) y sigue ampliando la dispersión de CVaR marginal, pero esa ampliación no es el mecanismo que mantiene la condición satisfecha: no es una garantía de diseño.

Dónde tiene que fallar: en los pesos convergidos.

En el otro extremo de la trayectoria el cuadro se invierte necesariamente. El CVaR marginal es la derivada de Euler $\partial \text{CVaR} / \partial w_i$, y en un mínimo sobre el simplex las condiciones de Karush—Kuhn—Tucker igualan los marginales de todos los activos mantenidos en el interior; la brecha mínima tiene, por tanto, que colapsar a lo largo de cualquier trayectoria descendente. Los datos muestran exactamente eso: en los pesos reportados la brecha mínima es $8,9 \times 10^{-5}$ para PA (AAP—NFLX) frente a $2\varepsilon \max_k |g_k^*| = 4,6 \times 10^{-4}$ —violación por $\approx 5 \times$ — y $1,5 \times 10^{-4}$ para PB (TROW—L) frente a $5,6 \times 10^{-4}$ —por $\approx 3,7 \times$ —. Equivalentemente, la precisión crítica $\varepsilon_c = \min_{i \neq j} |g_i^* - g_j^*| / (2 \max_k |g_k^*|)$ cae de $\approx 0,0091$ en w_0 a $\approx 0,00096$ (PA) y $\approx 0,0014$ (PB), por debajo del punto de operación $\varepsilon = 0,005$.² La condición es, en

²En el óptimo reportado de PA los marginales van de 0,0234 (NEM) a 0,0464 (OXY) pese a que la softmax mantiene todos los pesos positivos: el criterio de parada pondera por w_i , y para OXY, AAL, AAP y NFLX ($w_i \approx 0,016$ – $0,019$) ese factor suprime

consecuencia, suficiente *de inicio de trayectoria*: garantiza el ranking de las primeras actualizaciones, no el de las últimas. Cerca del óptimo, donde falla, las componentes cuyo orden puede invertirse son las de marginal casi igual, e intercambiar dos de ellas cambia el CVaR solo en una cantidad de primer orden en su brecha –inocuo para el objetivo–; su efecto sobre la *dirección* no se mide aquí, porque el coseno exportado se evalúa únicamente en w_0 .

Régimen de validez.

El argumento presupone que $|\delta_i|$ se mantiene por debajo del umbral de *aliasing* del dispositivo, situado en la Sección 4.4 en la ronda $m = 3$ (profundidad 22, en Emerald y en Garnet). Por encima, la actualización de verosimilitud de IQAE colapsa sobre la rama simétrica $\theta \leftrightarrow \pi/2 - \theta$ y produce $\hat{a} \approx 1 - a$ en vez de $(1 + \delta)a$: un plegado de fase que ningún esquema de optimización puede compensar, porque ya no es un error multiplicativo pequeño sino un salto a otra rama de la solución. Como ese cruce es geométrico a esta amplitud de cola (Sección 4.4.3), ninguna mejora de fidelidad restaura el argumento por encima de él: lo restaura restringirse a $m \leq 2$, o diferir algorítmicamente la rama simétrica con una amplitud codificada más pequeña. Con la misma honestidad que exige el Capítulo 6: la condición se verifica aquí en el punto de partida bajo ruido depolarizante *simulado*, no en hardware físico, de modo que si sobrevive a los errores de puerta *coherentes* responsables del *aliasing* en $m = 3$ sigue sin comprobarse.

En síntesis: en el punto de partida el ruido actúa como un reescalado, no como una rotación, del subgradiente estimado –lo infla hacia afuera sin cambiar su dirección–, y la normalización de segundo momento de Adam absorbe esa inflación. El optimizador converge bajo todos los niveles de ruido probados; la meseta de CVaR alcanzada se desplaza hacia arriba y la solución queda menos concentrada, como cuantifica la siguiente subsección.

Efecto del ruido sobre la concentración de la cartera

Aunque la dirección del gradiente en el punto de partida es invariante al ruido, el *contraste de magnitud* entre activos sí se erosiona: las estimaciones de CVaR marginal infladas por el ruido se parecen más entre sí, y Adam da pasos menos decisivos. El HHI cae de 0,210 (ideal) a 0,164 (E3, Tabla 4.4), la CVaR resultante es un 6,8 % más alta (0,02955 frente a 0,02768, la cifra canónica de la Sección 4.2.2), y la cuota de riesgo de los tres activos más concentrados cae del 74,9 % ideal (Tabla 4.3) al 60,7 % bajo E3, con el activo antes decisivo CLX desplazado del podio hasta el quinto puesto. El ruido, en otras palabras, no cambia *qué* activos son de riesgo pero sí *cuánto* se diferencian entre sí a ojos del optimizador, aplanando ligeramente –sin invertir– la jerarquía de riesgo.

u_i y la softmax los empuja a $w_i \rightarrow 0$ en vez de igualar su marginal. Son, a efectos prácticos, los activos *no mantenidos* de la condición KKT; la igualación solo afecta a los seis restantes, cuyos marginales se estrechan de 0,0117–0,0367 a 0,0234–0,0302.

4.3.3. Convergencia bajo ruido y el estancamiento de B2_PA y B2_n6_PA

La terminación de los 46 experimentos se reparte en tres grupos, verificados uno a uno sobre `metrics.optimization` de cada JSON: **40 convergen con normalidad** vía los criterios de norma de gradiente o de variación relativa de CVaR; **3 activan la parada temprana** –B2_PA y B2_n6_PA, ambas estancadas en el punto inicial con mejora del 0,0 %, y C2_PA, que *no* es un estancamiento porque mejora un 18,87 % en 44 iteraciones antes de agotar la paciencia–; y **3 agotan** `max_iterations` sin cumplir ningún criterio: B2_PB, S1_PB_100 y S1_PC_100, con mejoras del 12,9 % al 26,8 %. Las dos de escalabilidad siguen descendiendo al terminar –su última entrada de `cvar_history` es también la mínima–; B2_PB no, de modo que su cifra reportada es de nuevo el mejor-hasta-ahora. La forma de la curva –descenso rápido seguido de meseta– se preserva en todos los niveles de ruido (Figura 4.2): solo sube la altura de la meseta. Ninguno de los 46 oscila ni diverge.

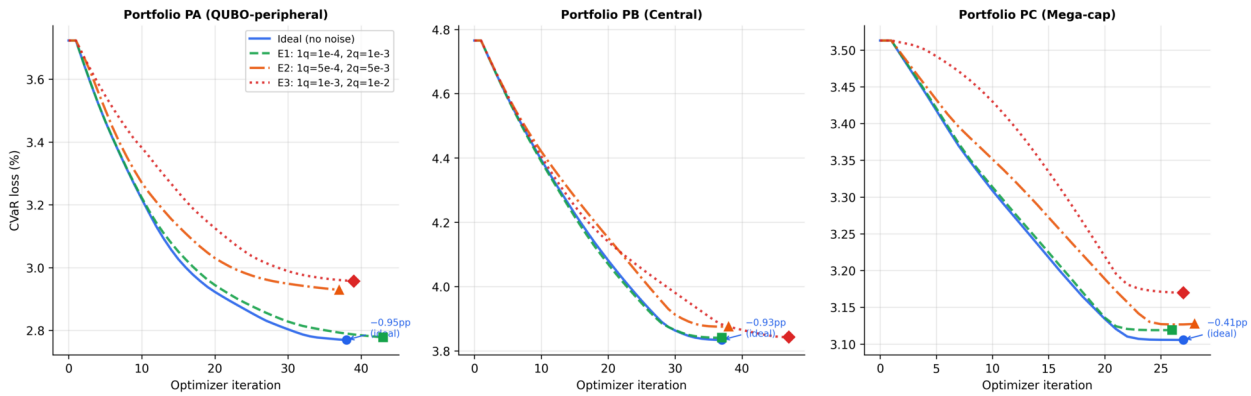


Figura 4.2: Pérdida CVaR (%) frente a la iteración para PA, PB y PC bajo la configuración ideal A1 y las tres de ruido E1–E3. Las doce curvas convergen con normalidad: el ruido desplaza la meseta hacia arriba sin alterar la forma de la trayectoria. Los marcadores finales señalan la parada de cada corrida y la anotación junto a cada curva ideal da la reducción en puntos porcentuales (–0,95 PA, –0,93 PB, –0,41 PC). En la batería completa 40/46 convergen así, y ninguna de las doce representadas está entre las seis restantes. Figura 6 del paper.

Dentro de la familia B2 (ruido $p_{2q} = 10^{-3}$, la tasa nominal de IQM; $\lambda \in \{1, 1,5, 2\}$; $\varepsilon = 0,005$), B2_PA y B2_n6_PA son las dos únicas corridas de la batería con una mejora reportada de 0,0 %: los pesos finales coinciden con el punto inicial y sus registros `convergence_history` son idénticos bit a bit. No es coincidencia entre dos ablaciones: comparando sus `config/*.yaml` línea a línea, la única diferencia es una clave `phase2.qae.n_qubits: 6` que el pipeline nunca lee para el ancho del circuito –el campo que se consume de verdad es `num_distribution_qubits`, que sigue valiendo 4 en ambos–. B2_n6_PA es, por tanto, una repetición idéntica de B2_PA bajo otra etiqueta: la batería contiene una única configuración verificada de estancamiento, registrada dos veces.

El JSON de B2_PA descarta que ese 0,0 % sea la firma de un paso de Adam casi nulo: el CVaR *in-sample* empeora monótonicamente un 12,9 % a lo largo de 20 iteraciones mientras gradient_norm_w se mantiene entre $6\,950\times$ y $9\,040\times$ el umbral de convergencia $\tau = 10^{-5}$ –pasos sustanciales en todo momento–. Lo que termina la corrida es la parada temprana sobre el par “mejor-hasta-ahora”, que nunca se actualiza tras la iteración 0 porque el CVaR no vuelve a mejorar: el aparente estancamiento en el punto de partida es un artefacto de la convención de reporte, no evidencia de que el optimizador no se moviera.

Ese modo de fallo, más informativo que un paso aleatorio, apunta a un problema sistemático de sesgo o signo propio de esta configuración. Se han comprobado y descartado tres hipótesis contra los ficheros de configuración reales: el factor de escala fraccionario $\lambda = 1,5$ –las ocho configuraciones D1/D2 lo comparten y convergen con normalidad–, el nivel de ruido –D2_PA, idéntica salvo un error de dos qubits $10\times$ mayor, mejora un 20,7 % en 40 iteraciones, lo contrario de lo que predeciría el ruido como causa– y el número de qubits –B2_n3_PA converge con normalidad, +10,9 % en 24 iteraciones–. No queda una explicación verificada de por qué B2_PA se estanca mientras sus vecinos más cercanos en el espacio de configuración convergen, y se reporta como pregunta abierta, no como un mecanismo sin verificar. Las similitudes coseno de ambas corridas en el punto de partida permanecen dentro de $\cos \angle > 0,9998$ (Tabla 4.4): la dirección inicial del gradiente no es lo que falla, y el estancamiento no contradice el resultado de similitud coseno, sea cual sea su causa raíz. Se conserva en la batería como resultado negativo documentado, sin excluirlo ni reponderar por su causa las estadísticas de convergencia.

4.4. Caracterización en hardware IQM: Emerald y Garnet

Esta sección responde a Q1 y Q3 (Sección 1.3) con datos de hardware físico real –los procesadores superconductores IQM Emerald (Crystal 54) y Garnet (Crystal 20), accedidos vía Amazon Braket y, para la comparación de ZNE, vía IQM Resonance–, no de simulación; es, junto con la derivación de la Sección 4.3.2, la parte más técnica del capítulo, y corresponde a la Contribución C1 del Capítulo 1 (caracterización en hardware de IQAE) y a la Contribución C2 (comparativa de ZNE). Dicho desde antes de la primera cifra: ninguna de las ejecuciones de esta sección corre el circuito IQAE propio del pipeline (cuatro qubits de distribución más un ancilla, Sección 3.3 del Capítulo 3), sino una *sonda* mínima de dos qubits que reproduce la misma amplitud de cola codificada; la Sección 4.4.1 explica qué es esa sonda y qué se puede concluir legítimamente de ella sobre el circuito real, y qué no.

Qué es el *aliasing* en estimación de amplitud, y por qué aparece

IQAE estima una amplitud $a \in [0, 1]$ –aquí, la probabilidad de que la pérdida caiga en la cola de CVaR– codificándola como un ángulo θ con $a = \sin^2 \theta$. Cada ronda de Grover no mide θ directamente,

sino la probabilidad $p(m) = \sin^2((2m+1)\theta)$, que es periódica y, sobre todo, simétrica alrededor de $\pi/2$: $\sin^2(\pi/2+x) = \sin^2(\pi/2-x)$. Si $(2m+1)\theta$ se acerca demasiado a $\pi/2$, $p(m)$ ya no distingue si el ángulo verdadero era θ o su reflejo $\pi/2 - \theta$. En el caso ideal IQAE resuelve la ambigüedad combinando rondas con distinto m ; en un dispositivo real, el ruido acumulado puede hacer que la estimación converja a la rama equivocada y el dispositivo reporte $\hat{a} \approx 1 - a$. Eso es el *aliasing* de amplitud.

La condición geométrica se sigue directamente: la ambigüedad aparece cuando $(2m+1)\theta \approx \pi/2$, es decir, en la ronda

$$m \approx \frac{1}{2} \left(\frac{\pi}{2\theta} - 1 \right), \quad (4.3)$$

matemática pura, independiente del hardware, que predice *antes de ejecutar nada* dónde entrará el algoritmo en la zona de ambigüedad. Para la cola al 5 % de PA, $\theta = \arcsin \sqrt{0,0504} = 0,2265$, y sustituyendo sale $m \approx 3$. El valor de la caracterización que sigue está en que *confirma* con mediciones físicas una predicción hecha por adelantado, no en descubrir empíricamente un umbral desconocido.

4.4.1. ¿Qué se ejecutó realmente en hardware? La sonda de dos qubits

Antes de presentar ninguna medida hay que ser explícitos sobre qué circuito se ejecutó de verdad en los procesadores IQM Emerald y Garnet, porque **no** es el circuito de cuatro qubits de distribución más un ancilla que el Capítulo 3 describe como el corazón de la Etapa 2 del pipeline. Es una práctica reconocida en la literatura de caracterización de hardware NISQ, no un atajo improvisado, y es el contenido de esta sección.

Por qué no se ejecutó el circuito real del pipeline.

El circuito de la Etapa 2, ya en $m = 0$ –preparación de estado más oráculo de umbral, sin ninguna ronda de amplificación–, tiene una profundidad pre-transpilación archivada de 44 puertas para el circuito de probabilidad de cola y 47 para los de esperanza condicionada, con 75–84 en total sobre 5 qubits. Transpilado *offline* a puertas nativas de IQM –misma llamada `transpile(..., optimization_level=1)` que usan los guiones de hardware– alcanza 1,476 de profundidad con 782 puertas CZ ya en $m = 0$, y 5,395 en $m = 1$: dos órdenes de magnitud por encima de lo que las fidelidades de dos qubits del hardware NISQ actual sostienen de forma coherente. Ejecutarlo habría producido solo ruido, y de hecho **nunca se ha ejecutado en hardware IQM, a ningún m** : las 46 corridas lo ejecutan siempre en `qiskit-aer`.

Qué se ejecutó en su lugar: la sonda.

Los siete guiones de hardware construyen todos el mismo circuito mínimo de dos qubits, $\mathcal{A}_{\text{sonda}} = CX_{0 \rightarrow 1} R_y(2\theta)_0$, que prepara $\cos \theta |00\rangle + \sin \theta |11\rangle$ con $\sin^2 \theta = a$ sobre un qubit de datos y un ancilla,

seguido del operador de Grover $\mathcal{Q} = \mathcal{A} S_0 \mathcal{A}^\dagger S_\chi$ con $S_0 = X^{\otimes 2} CZ X^{\otimes 2}$ y $S_\chi = Z$ sobre el ancilla. Preparar la misma amplitud codificada con un qubit de datos y un ancilla, en lugar del registro completo de 16 amplitudes, es la construcción que Woerner y Egger [1] usaron en sus demostraciones de hardware: una estrategia reconocida para aislar la dinámica de rondas de Grover de la parte que solo añade profundidad sin cambiarla, no una idea nueva de este trabajo. Los metadatos archivados documentan la elección sin ambigüedad (`phase1_direct_results_latest.json`, campo `encoding`: rotación directa `Ry(2*arcsin(sqrt(a_ideal)))` seguida de `CX`), y la cabecera del guion da el motivo: el circuito cargado por el método de Möttönen transpila a ≈ 52 de profundidad antes de ninguna amplificación.

Por qué es legítima –y en qué se queda corta.

El umbral de *aliasing* de la Sección 4.4 depende *únicamente* de la amplitud codificada $a = \sin^2 \theta$, no de qué circuito la codifique. La sonda codifica exactamente la misma a que el circuito real (mismo $\theta \approx 0,2265$) y reproduce por tanto la misma dinámica $p(m) = \sin^2((2m+1)\theta)$: la conclusión del umbral $m \approx 3$ *sí* se transfiere sin cambios. Lo que **no** se transfiere es el suelo de precisión sub-umbral limitado por fidelidad, que escala con la profundidad: ahí sonda y circuito real difieren en más de dos órdenes de magnitud en todo m (1,476 frente a 4 en $m = 0$). Por tanto $\varepsilon_{\text{hw}} \approx 0,20$ es el suelo de la sonda y una *cota inferior* del suelo del circuito real –que no se ha medido–, lo que refuerza, en lugar de debilitar, la conclusión negativa de la Amenaza 2 (Sección 6.3): la brecha entre el hardware actual y el objetivo $\varepsilon = 0,005$ es, como poco, tan grande como la que allí se reporta. Un detalle de trazabilidad, divulgado por transparencia: la amplitud de referencia de los guiones, $a_{\text{ideal}} = 0,050442$, es una constante fija en el código –intentan leer un campo que no existe en ningún JSON de `results/` y recaen en el valor por defecto–, frente a los valores archivados reales 0,050482 (A1_n3_PA) y 0,049565 (A1_PA); los tres están a menos del 1 % del 0,05 nominal, así que ninguna conclusión depende de esta elección.

4.4.2. Fórmula de profundidad de circuito y umbral de *aliasing*

Para hablar de “ronda m ” en términos de hardware hace falta saber cuánto crece la profundidad física con cada ronda. Para la sonda de dos qubits –no el circuito del pipeline–, la profundidad tras compilar a puertas nativas de IQM (`opt_level=1`) sigue

$$\text{profundidad}(m) = D_{\mathcal{A}} + 6m, \quad (4.4)$$

con $D_{\mathcal{A}} = 4$. Verificada en $m \in \{0, \dots, 6\}$ en Emerald y en $m \in \{0, 1, 3\}$ en Garnet, la pendiente +6 refleja tres puertas CZ por aplicación de $\mathcal{Q}$ –una de $\mathcal{A}^\dagger$, una de S_0 , una de $\mathcal{A}$ – más sus capas de rotación de un qubit; el recuento archivado de CZ, $1 + 3m$, lo confirma. La pendiente es específica de esta sonda y del transpilador de IQM: para el circuito real del pipeline $D_{\mathcal{A}}$ queda sustituida por 1,476.

Ronda m	Profundidad medida	Fórmula $4 + 6m$
0	4	4
1	10	10
2	16*	16
3	22	22
4	28	28
5	34	34
6	40	40

*Derivada de la fórmula: ningún registro de trabajo archivado cubre $m = 2$.

Tabla 4.5: Profundidad post-transpilación en puertas nativas de IQM Emerald para la **sonda de dos qubits** –no el circuito del pipeline–, verificando profundidad(m) = $4 + 6m$ en todo el rango probado (Tabla 5 del paper). El registro que habría cubierto $m = 2$ no está archivado, así que esa fila se deriva de la fórmula*.

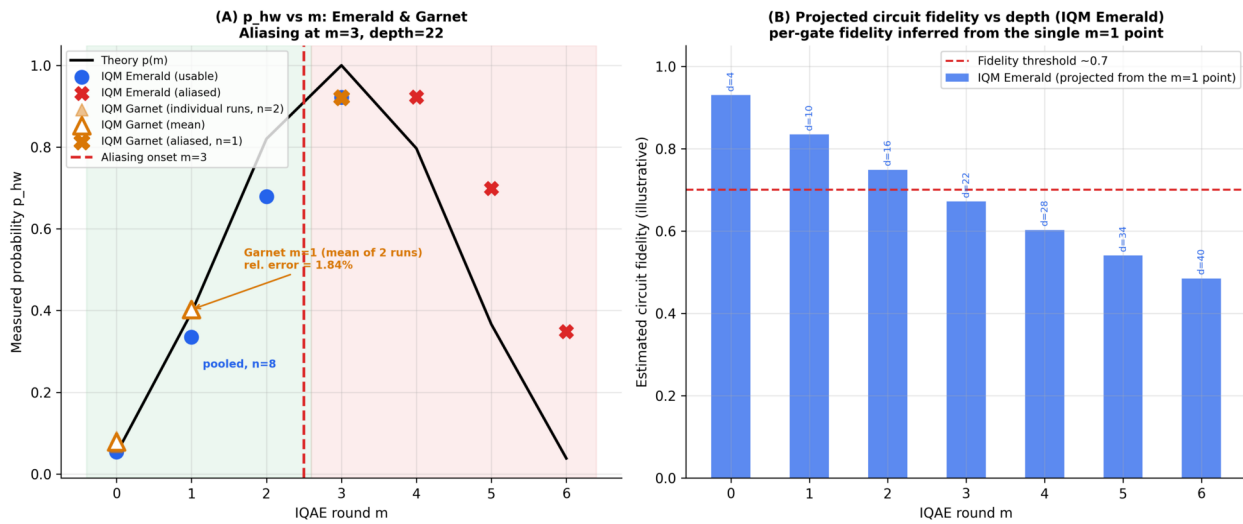


Figura 4.3: Probabilidad medida p_{hw} frente a la ronda m para la **sonda de dos qubits** en Emerald (círculos) y Garnet (triángulos), contra la predicción $p_{th}(m) = \sin^2((2m+1)\theta)$ con $\theta = 0,2265$. **(A):** curva completa $m = 0, \dots, 6$ con el colapso de rama predicho en $m \approx 3$; verde, rango utilizable ($m \leq 2$); rojo, aliased. Garnet muestra sus dos ejecuciones y su media. **(B),** solo ilustrativo: la fidelidad proyectada se infiere de un único punto $m = 1$ por dispositivo, ayuda visual sin peso probatorio. Figura 7 del paper.

Ambos dispositivos siguen la teoría hasta $m = 2$ (profundidad ≤ 16) y alcanzan el colapso geométrico de rama exactamente en $m = 3$ (profundidad 22), como predice (4.3). Allí la predicción es $p_{th} = 0,9998$ –casi la unidad, porque $(2m+1)\theta \approx \pi/2$ – y las medidas son 0,921 (Emerald) y 0,9205 (Garnet), mutuamente consistentes. El pequeño error nominal es engañoso: refleja la saturación de $\sin^2$ cerca de su máximo, no una estimación fiel. El síntoma real del *aliasing* es que más allá de $m = 3$ la verosimilitud de IQAE ya no separa las dos ramas –el intervalo se actualiza para $\theta \in [\pi/2 - \hat{\varepsilon}, \pi/2]$ en lugar de $[0, \hat{\varepsilon}]$, equivalente a estimar $a \approx 1 - a$ –, no el tamaño del error puntual. Esto limita la precisión fiable a $\varepsilon_{hw} \approx 0,20$, la semianchura alcanzable con la máxima ronda utilizable $m = 2$: la cifra que deja la condición (iii) de la ventaja práctica muy por encima del umbral $\varepsilon < 0,05$ de la condición (i).

4.4.3. Comparación entre dispositivos: Emerald frente a Garnet

La Tabla 4.6 reúne los resultados de hardware en ambos dispositivos para las rondas $m = 0, 1, 2, 3$.

Dispositivo	m	p_{hw}	p_{th}	Error	¿Aliasado?
IQM Emerald	0	0.0537	0.0504	6.5 %	No
IQM Emerald	1	0.335 (8 ejec.)	0.3950	15.2 %	No
IQM Emerald	2	0.679	0.8200	17.2 %	No
IQM Emerald	3	0.921	0.9998	7.8 % [†]	Sí
IQM Garnet	0	0.078 (2 ejec.)	0.0504	55.6 %	No
IQM Garnet	1	0.402 (2 ejec.)	0.3950	1.84 %	No
IQM Garnet	3	0.9205	0.9998	7.9 % [†]	Sí

Tabla 4.6: Resultados de IQAE en hardware para la sonda de dos qubits (Tabla 6 del paper). Predicciones teóricas: $p_{th} = 0,0504, 0,3950, 0,8200, 0,9998$ para $m = 0, 1, 2, 3$. Garnet en $m \in \{0, 1\}$ promedia dos ejecuciones; Emerald en $m = 1$ es la media agrupada de ocho. [†]En $m = 3$ el error nominal es pequeño porque la curva está cerca de su pico, pero la ejecución se marca aliasada porque IQAE ya no distingue las ramas θ y $\pi/2 - \theta$.

Fuente: los bloques `phase1.grover` y `pooled_stats` de los registros de Emerald, y `iqae_garnet_results_latest.json` junto con `garnet_new` para Garnet, todos bajo `results/hardware_validation/`; las cifras de Garnet promedian sus dos ejecuciones, cuyos identificadores recoge la Sección 4.4.3.

En $m = 0$ Garnet muestra un error relativo mayor que Emerald (55,6 % frente a 6,5 %), en parte porque la amplitud objetivo está cerca de cero: un sesgo aditivo pequeño –de lectura o de preparación de estado, el canal dominante antes de amplificar– produce ahí un error porcentual grande, mientras que la misma desviación absoluta en $m = 1$ da un error relativo mucho menor. $m = 0$ es, por tanto, el punto menos informativo y no se pondera en las conclusiones. En $m = 1$ el error medio de Garnet es 1,84 %, unas ocho veces menor que el 15,2 % agrupado de Emerald, con sus dos ejecuciones mutuamente consistentes; no es atribuible a la topología –ambos comparten la arquitectura de red cuadrada de IQM y difieren en número de qubits, no en conectividad– sino a factores no controlados: ciclos de calibración distintos y el par de qubits que el compilador de Braket selecciona en cada dispositivo. Dos advertencias la matizan –muestra pequeña en Garnet ($n = 2$ frente a $n = 8$) y fidelidades por qubit no emparejadas–, y se retoman como Amenaza 1 en el Capítulo 6. La diferencia *no* se extiende al umbral de *aliasing*: en $m = 3$ ambos devuelven $p_{hw} \approx 0,92$, consistente con el origen geométrico de (4.3). Los factores no emparejados entre ambos dispositivos quedan documentados para que el lector juzgue por sí mismo el alcance de la comparación: todas las ejecuciones se enviaron vía Amazon Braket el 13 de abril de 2026 en una ventana de 26 minutos, pero Emerald y Garnet corresponden a ciclos de calibración distintos y los pares de qubits que asigna el compilador para la sonda no están emparejados por fidelidad. Los disparos varían por bloque: 2048 en el barrido de Grover de la fase 1 y en las cinco repeticiones de la oleada 2, 4096 en la comprobación de preparación de estado, los bloques ZNE, las rondas extendidas y el bloque G de Garnet, y 8192 en las rondas de IQAE de la

oleada 1 (v2) y de Garnet tras la oleada 3. Aparte queda el bloque D de la oleada 3 (Emerald, $m = 1$), un estudio de disparos a 1 024, 4 096 y 8 192 ($p_{\text{hw}} = 0,326, 0,343$ y $0,328$) que es el origen del mínimo de 1 024 citado en las tablas de parámetros.

El análisis de repetibilidad de la oleada 2 en $m = 1$ agrupa 8 ejecuciones sobre Emerald (20 480 disparos): media $\bar{p}_{\text{hw}} = 0,335$, IC 95 % $[0,329, 0,341]$, $z = -20,9$ frente a la teoría con el error estándar entre ejecuciones ($s = 0,0081$). Un intervalo tan estrecho confirma que el error es *sistemático*, no estadístico: el ruido de disparo da una semianchura de $\approx 0,006$ frente a un sesgo observado de $-0,060$, diez veces mayor, así que repetir la medición no lo haría desaparecer.

Origen geométrico frente a origen de fidelidad del umbral

Cerramos la comparación distinguiendo dos preguntas fáciles de confundir: ¿por qué aparece el umbral de *aliasing* exactamente en $m = 3$?, y ¿por qué los dos dispositivos difieren por debajo de él? La primera respuesta es geométrica e independiente del dispositivo: la amplificación legítima alcanza la rama simétrica $(2m + 1)\theta \approx \pi/2$ en $m \approx 3$ para $\theta \approx 0,2265$ en cualquier plataforma, sin importar su fidelidad, y los datos lo respaldan –Emerald y Garnet devuelven $p_{\text{hw}}(3)$ casi idéntico (0,921 frente a 0,9205) pese a sus tamaños y ciclos de calibración distintos–. La segunda es distinta: los errores de puerta coherentes sí sesgan $\hat{\theta}$ por debajo del umbral –visible en $m = 1$ – y en principio podrían adelantar la rama simétrica si la fidelidad se degradara lo suficiente, pero a las fidelidades actuales el origen geométrico se alcanza primero y ambos entran en *aliasing* juntos en $m = 3$. El suelo de precisión sub-umbral sí lo gobierna la fidelidad, con el sesgo escalando aproximadamente como $(1 - f_{2q}) \cdot \text{profundidad}(m)$, pero no mueve el cruce: la ronda en que $(2m+1)\theta$ alcanza $\pi/2$ la fija la amplitud codificada, es decir el nivel de cola $\alpha \approx 0,05$, no el dispositivo ni el tamaño del registro (más qubits de discretización refinan la masa de cola y dejan θ intacto). Alcanzar $m \geq 4$ a este nivel de cola exige, por tanto, una modificación *algorítmica* que difiera la rama simétrica –codificar una amplitud reescalada a/c mediante una rotación sobre un qubit auxiliar, el recurso de la valoración cuántica de opciones [14], o apuntar a un nivel de cola más pequeño– junto con la fidelidad necesaria para circuitos más profundos; esta última no puede por sí sola desbloquear $m \geq 7$ a $\theta \approx 0,2265$, donde $(2m+1)\theta \approx 3,4 \text{ rad} > \pi$.

Esta generalización geométrica se ha verificado en hardware real *para un único valor de $\theta \approx 0,2265$* , el de la cola al 5 % de la cartera PA. Los otros niveles de la batería de Basilea IV ($\alpha = 0,01, 0,10$) nunca se ejecutaron en IQAE sobre hardware físico, así que si el umbral se desplaza a otro m para esos θ sigue sin comprobarse –séptima amenaza a la validez del Capítulo 6–.

Identificadores de trabajo en hardware IQM

Los identificadores de trabajo (UUID) de Amazon Braket e IQM Resonance que sustentan los resultados de esta sección –trazables, cada uno, a un dispositivo y una fecha concretos– se re-

cogen como muestra representativa en el Apéndice A.4; el listado completo por bloque está en `results/hardware_validation/` en el repositorio del proyecto (Sección 7.3).

4.4.4. Mitigación de errores con ZNE: comparativa y fallo del *folding*

La extrapolación a ruido cero (ZNE) es una familia de técnicas de *mitigación* —distinta de la corrección de errores en sentido estricto, que requeriría codificación redundante— que estima el resultado sin ruido extrapolando mediciones bajo ruido *amplificado artificialmente* a varios factores $\lambda = 1, 2, 3, \dots$, conocida (o supuesta) la forma funcional de la degradación, hasta $\lambda = 0$ —un punto que nunca se mide directamente—. Esta sección responde a Q3: ¿mejora ZNE la precisión de IQAE en hardware IQM, o el coste adicional de profundidad que introduce anula la ganancia?

El fallo del *folding* global en hardware

La forma más directa de amplificar el ruido artificialmente es el *folding* (pliegue) de circuito: sustituir el circuito U por $U(U^{-1}U)^n$, que es matemáticamente idéntico a U en ausencia de ruido pero que, al ejecutar n veces la secuencia $U^{-1}U$, expone al dispositivo real a $2n + 1$ veces más ruido acumulado. Aplicado al circuito de IQAE en $m = 1$ (profundidad base 10) con factores de escala $\lambda \in \{1, 3, 5\}$, las profundidades post-transpilación *esperadas* son $\{10, 30, 50\}$; las profundidades *realmente* medidas a `opt_level=0` son $\{10, 136, 226\}$ —una sobrecarga de profundidad de $4,5\times$ por encima del factor de escala previsto—. La causa raíz: a `opt_level=0` el transpilador no cancela ninguna puerta entre U^{-1} y el siguiente U , así que la profundidad real crece como $\sim 4,5 \cdot \lambda \cdot \text{profundidad}_0$, no $\lambda \cdot \text{profundidad}_0$.

A $\lambda = 3$ (profundidad real 136), la fidelidad de circuito estimada es $F \approx (0,995)^{136} \approx 0,51$ —esencialmente un bit aleatorio—. La estimación cruda de IQAE en $\lambda = 1$, pese a su propio sesgo sistemático, se mantiene mucho más cerca de la amplitud teórica (error relativo agrupado 15,2 %, $\Delta p_{\text{hw}} \approx 0,060$ en la media de la oleada 2) que la estimación extrapolada por ZNE en $\lambda = 3$, que diverge hasta 0,271 de desviación absoluta: la extrapolación no solo no mejora la estimación cruda, la empeora drásticamente. ZNE vía *folding* global es, en estas condiciones, contraproducente en el hardware IQM actual.

Comparativa de métodos de mitigación en simulador

¿Es un problema específico del *folding* global, o genérico a cualquier variante de ZNE sobre IQAE? Para responderlo sin gastar más tiempo de hardware se ejecutó una comparación controlada en el modelo de ruido depolarizante de `qiskit-aer` sobre el circuito de IQAE en $m = 1$ ($p_{\text{th}} = 0,395$), bajo cuatro condiciones —cruda; *folding* global $U \rightarrow U(U^\dagger U)^n$; *folding* de puerta, $L_j \rightarrow L_j(L_j^\dagger L_j)^n$ aplicado in situ; e inserción de identidad de tipo FIIM, que amplifica solo las puertas de dos qubits autoinversas mediante $G \rightarrow G(GG)^n$, un esquema libre de inversión porque $G^{-1} = G$ — y tres niveles de ruido que

coinciden con las tasas de E1/E2/E3, con $\lambda \in \{1, 3, 5\}$, 4 096 disparos y tres réplicas por punto (semilla 42).

Método	Ruido	prof. _{1,3,5}	Δ_{cruda}	Δ_{ZNE}	¿cruda mejor?
<i>Folding</i> global y de puerta*	E1 ($p_{2q} = 0,001$)	9, 25, 41	0,0017	0,0023	✓
	E2 ($p_{2q} = 0,005$)	9, 25, 41	0,0000	0,0019	✓
	E3 ($p_{2q} = 0,01$)	9, 25, 41	0,0025	0,0011	—
Inserción de identidad (FIIM)	E1 ($p_{2q} = 0,001$)	9, 17, 25	0,0017	0,0023	✓
	E2 ($p_{2q} = 0,005$)	9, 17, 25	0,0000	0,0015	✓
	E3 ($p_{2q} = 0,01$)	9, 17, 25	0,0025	0,0012	—

*Bajo ruido depolarizante independiente de la puerta ambos métodos coinciden cifra a cifra, como predice la teoría, y se dan en una sola fila.

Tabla 4.7: Métodos de ZNE sobre IQAE en el modelo depolarizante de `qiskit-aer` ($m = 1$, $p_{\text{th}} = 0,395$, 4096 disparos, 3 réplicas, semilla 42). La profundidad es post-transpilación a $\text{opt_level}=0$. Δ_{cruda} es el sesgo de la estimación cruda en $\lambda = 1$; Δ_{ZNE} , el de la extrapolación de Richardson.

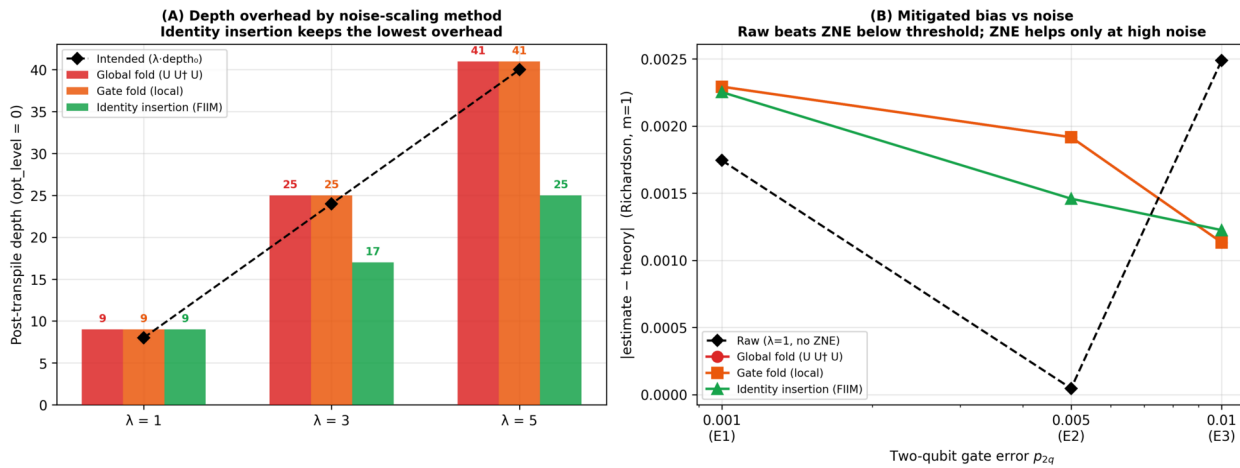


Figura 4.4: Comparativa de ZNE en el modelo de ruido de `qiskit-aer` (IQAE en $m = 1$). **(A):** profundidad post-transpilación a $\text{opt_level}=0$ frente al factor de escala; los dos *folding* la inflan por encima de la prevista $\lambda \cdot \text{prof}_0$ (rombos), y FIIM queda muy por debajo. **(B):** sesgo de la extrapolación de Richardson frente a p_{2q} , con la cruda en $\lambda = 1$ como referencia. Las curvas de *folding* global y de puerta se superponen exactamente en (B), como predice su equivalencia bajo ruido depolarizante, de ahí que solo una sea visible. Figura 9 del paper.

Dos hallazgos emergen (Figura 4.4). En profundidad, la inserción de identidad alcanza una sobrecarga menor que el *folding* (25 frente a 41 en $\lambda = 5$), lo que confirma que la inflación de la Sección 4.4.4 es un artefacto de amplificar el circuito completo y no solo el ruido dominante de las puertas de dos qubits. En precisión, para $m = 1$ por debajo del umbral de *aliasing*, **ningún** método de extrapolación mejora la estimación cruda al objetivo $\varepsilon = 0,005$ (sesgo crudo $\leq 2,5 \times 10^{-3}$, extrapolado comparable o mayor en E1/E2); solo en E3 ZNE reduce el sesgo, y allí la inserción de identidad es competitiva con el *folding* a la mitad de la profundidad. La recomendación accionable: por debajo del umbral, preferir la estimación cruda e invertir el presupuesto de disparos en repeticiones en lugar de en escalado de

ruido; donde ZNE sí está justificado, a ruido más alto, la inserción de identidad es la opción eficiente en profundidad.

Confirmación en hardware IQM Emerald

Para confirmar que las conclusiones de simulador sobreviven al ruido real, la comparación se repitió sobre IQM Emerald vía IQM Resonance (16 circuitos, 2048 disparos; dos réplicas por punto salvo FIIM en $\lambda \in \{1, 5\}$, con tres; el *folding* de puerta se omite porque coincide con el global bajo ruido depolarizante). La compilación a puertas nativas agrava la penalización de profundidad frente al recuento lógico: el *folding* global infla $m = 1$ de 10 a 112 ($\lambda = 3$) y 186 ($\lambda = 5$), mientras la inserción de identidad alcanza solo 70 y 102, aproximadamente la mitad a igual λ .

Método	prof. _{HW} ($\lambda=1, 3, 5$)	$p_{\text{hw}}(\lambda=1)$	Δ_{cruda}	Δ_{ZNE}
Cruda ($\lambda = 1$)	10	0,336	0,058	—
<i>Folding</i> global	10, 112, 186	0,338	0,057	0,217
Inserción de identidad (FIIM)	10, 70, 102	0,337	0,058	0,198

Tabla 4.8: Comparación de ZNE en IQM Emerald ($m = 1$, $p_{\text{th}} = 0,395$, 2048 disparos). prof._{HW} es la profundidad post-transpilación en la base nativa $\{r, \text{CZ}\}$. La estimación cruda es la más precisa en todos los casos.

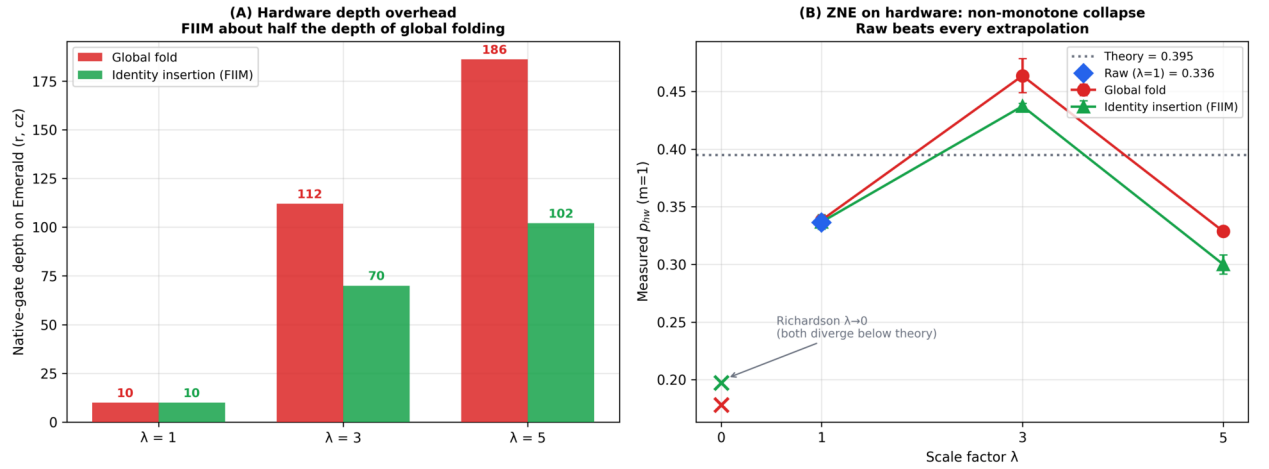


Figura 4.5: Comparativa de ZNE en IQM Emerald (IQAE en $m = 1$, 2048 disparos; dos réplicas por punto, tres para FIIM en $\lambda \in \{1, 5\}$). **(A):** profundidad en puertas nativas frente al factor de escala; FIIM alcanza la mitad que el *folding* global a igual λ . **(B):** p_{hw} frente al factor de escala; ambos suben en $\lambda = 3$ y colapsan hacia la mezcla máxima en $\lambda = 5$, de modo que Richardson diverge por debajo de la teoría y el punto crudo (rombo) es el más preciso, por un factor ~ 4 . Figura 10 del paper.

En precisión el resultado es inequívoco (Figura 4.5, panel B): la estimación cruda en $m = 1$ es $p_{\text{hw}} = 0,336 \pm 0,006$ (error del 14,8 %, consistente con el valor agrupado de la oleada 2), mientras que las curvas escaladas son marcadamente *no monótonas* —suben hasta $\approx 0,46$ (0,44 para FIIM) en $\lambda = 3$ antes de colapsar hacia la mezcla máxima en $\lambda = 5$ —, de modo que la extrapolación de Richardson

diverge *por debajo* de la teoría (0,178 y 0,197 frente a 0,395): un sesgo de $\sim 0,20$, unas cuatro veces el sesgo crudo de 0,058. El hardware corrobora el hallazgo de simulador en su forma más contundente: por debajo del umbral de *aliasing* ningún método digital de ZNE mejora la estimación cruda, y a las profundidades que exige $\lambda \geq 3$ la extrapolación es activamente perjudicial. La respuesta a Q3 es, por tanto, matizada: ZNE no es inútil en general —a ruido alto sí reduce el sesgo—, pero en el régimen de este trabajo su coste de profundidad supera cualquier ganancia.

4.5. Síntesis del capítulo

Los resultados cierran, con datos concretos, las tres preguntas de la introducción. Sobre Q1 (cuántas rondas de Grover aguanta el hardware): $m \leq 2$ en ambos procesadores IQM, con el umbral predicho de antemano por pura geometría y confirmado —no descubierto— por la medición. Sobre Q2 (si la dirección del subgradiente sobrevive al ruido): sí, en el punto de partida de cada trayectoria, por la condición de brecha de ranking derivada en la Sección 4.3.2 y verificada en los 46 experimentos; y es una garantía que se agota a lo largo de la trayectoria, porque en los pesos convergidos las condiciones KKT igualan los CVaR marginales y la brecha mínima colapsa, de modo que la preservación de la dirección *a lo largo* de la trayectoria se infiere de la convergencia, no se mide. Sobre Q3 (si ZNE mejora la precisión en este régimen): no, por debajo del umbral de *aliasing*, y el mecanismo del fallo —sobrecarga de profundidad del *folding* global— queda caracterizado en simulador y en hardware real.

Un matiz atraviesa las tres respuestas: la selección QUBO de la Etapa 1 aporta una mejora real y verificada de diversificación, pero *no* es lo que hace posible la respuesta afirmativa a Q2 —la condición se cumple con el mismo margen para PA y para PB, porque la gobierna la brecha mínima entre pares y no la dispersión que la selección amplía—. Lo que la selección sí explica es la brecha de rendimiento entre universos, más ancha que cualquier brecha inducida por el ruido dentro de un universo fijo: ese hilo abre el Capítulo 6. El Capítulo 5 cierra antes el círculo abierto en la Sección 1.4, usando la cartera PA como entrada a la batería de Basilea IV.

APLICACIÓN FINANCIERA

5.1. Backtesting regulatorio desde cero

Un modelo de riesgo de mercado –cuántico o clásico– no se valida comprobando que "parece razonable", sino contrastando la serie de pérdidas observadas en una ventana fuera de muestra (*out-of-sample*, en adelante OOS) contra lo que el modelo predijo, con las herramientas estadísticas formales que siguen; la regulación bancaria internacional (el marco de Basilea) exige superar esta comprobación antes de que un banco pueda usar un modelo interno para calcular el capital que debe reservar por riesgo de mercado [6].

5.1.1. La unidad básica y los dos tests de VaR

Dado un nivel α , el VaR_α es el umbral de pérdida que solo se supera con probabilidad α . En un backtest sobre T días ocurre una *excepción* el día t si la pérdida realizada supera el VaR_α predicho para ese día, $I_t = \mathbb{1}\{L_t \geq \text{VaR}_{\alpha,t}\} \in \{0, 1\}$. Si el modelo está bien calibrado la fracción de días con excepción debe converger a α : ni más –modelo demasiado optimista– ni menos –modelo excesivamente conservador, que exige más capital del necesario–.

Kupiec: cobertura incondicional.

La pregunta más simple es cuántas excepciones ocurrieron y si ese número es compatible con α . Con $x = \sum_t I_t$ y $\hat{p} = x/T$, bajo H_0 cada día es un ensayo de Bernoulli independiente de probabilidad α , de modo que $x \sim \text{Bin}(T, \alpha)$. Kupiec [29] contrasta $p = \alpha$ frente a $p = \hat{p}$ con el estadístico de razón de verosimilitudes

$$LR_{uc} = -2 \ln \left[\frac{(1 - \alpha)^{T-x} \alpha^x}{(1 - \hat{p})^{T-x} \hat{p}^x} \right] \sim \chi^2(1),$$

que cuenta excepciones sin mirar *cuándo* ocurrieron –limitación que motiva el siguiente test–.

Christoffersen: cobertura condicional.

Un modelo puede acertar la tasa media y aun así ser malo si las excepciones se agrupan en el tiempo, dentro de una misma crisis, en lugar de repartirse. Christoffersen [30] añade una comprobación de independencia –si las excepciones fueran independientes, la probabilidad de una hoy no dependería de si hubo una ayer– mediante un segundo estadístico $LR_{ind} \sim \chi^2(1)$, y combina ambas sumándolos: $LR_{cc} = LR_{uc} + LR_{ind} \sim \chi^2(2)$ bajo H_0 . Rechazar LR_{cc} puede deberse a una tasa incorrecta, a agrupación temporal o a ambas. Un modelo que agrupa sus fallos en los peores momentos es más peligroso que otro con el mismo total mejor repartido –el capital faltaría justo cuando más se necesita–, razón por la que Basilea exige esta comprobación además de la de Kupiec.

5.1.2. Los estadísticos Z_1 y Z_2 de Acerbi–Székely

Los dos tests anteriores solo miran el VaR: cuentan excepciones binarias sin preguntar *cuánto* se superó, precisamente lo que CVaR sí captura. Acerbi y Székely [8] cierran esa carencia con dos estadísticos que difieren solo en cómo normalizan la misma suma de pérdidas en exceso. Si X_t es el resultado del día t –negativo cuando hay pérdida–, el primero, que en la convención original de los autores se llama Z_2 (*incondicional*), normaliza por el número de excepciones *esperado*, $T\alpha$; el segundo, Z_1 (*condicional*), por el número *efectivamente observado*, $N_T = \sum_t \mathbb{1}\{X_t \leq -\text{VaR}_{\alpha,t}\}$:

$$Z_2 = \frac{1}{T\alpha} \sum_{t=1}^T \frac{X_t \mathbb{1}\{X_t \leq -\text{VaR}_{\alpha,t}\}}{\text{ES}_{\alpha,t}} + 1, \quad Z_1 = \frac{1}{N_T} \sum_{t=1}^T \frac{X_t \mathbb{1}\{X_t \leq -\text{VaR}_{\alpha,t}\}}{\text{ES}_{\alpha,t}} + 1.$$

Ambos promedian, sobre los días con excepción, las pérdidas realizadas normalizadas por el ES predicho, desplazadas en +1 para que su valor esperado sea cero bajo un modelo correctamente especificado. Un Z significativamente negativo indica pérdidas peores de lo predicho –modelo optimista, el caso que preocupa al regulador–; uno positivo, lo contrario. La diferencia práctica es de peso estadístico, no de interpretación: Z_2 normaliza por una constante y tiene varianza estable incluso con pocas excepciones, mientras que N_T puede bajar a 1 o 3 sobre 501 días en los niveles más exigentes. Por eso se reporta Z_2 como principal y Z_1 como complemento.

Ni Z_1 ni Z_2 tienen distribución asintótica sencilla bajo H_0 , así que su significación solo puede evaluarse remuestreando –con bloques de 22 días en el repositorio, distintos de los 21 de la Sección 5.4 por ser remuestreos independientes–. Pero a *qué* se aplica ese remuestreo es donde una versión anterior de este análisis se equivocó.

5.1.3. Un error heredado y su corrección

El código archivado etiqueta como “ z_1 ” la expresión de Z_2 con el signo invertido, de modo que el valor $-0,528$ que circulaba no era el Z_1 condicional sino $-Z_2$: corregido, $Z_2(\alpha = 5\%) = +0,528$

y $Z_1(\alpha = 5\%) = +0,261$, los valores ya usados en la Tabla 5.4. El “ $p = 0,47$ ” y el intervalo $\approx [-0,82, -0,24]$ archivados no proceden de ningún test de hipótesis, sino de un *bootstrap* de bloques sobre las propias pérdidas OOS observadas, centrado por construcción en el valor observado: de ahí que el “ p -valor” rondara 0,5 y el intervalo excluyera el cero con independencia de si el modelo acertaba.

La forma correcta es simular la ley predictiva del modelo y ver dónde cae el valor observado dentro de esa distribución. Como este pipeline predice VaR/ES por simulación histórica, su ley predictiva es la distribución empírica de las pérdidas de *entrenamiento*: se generan 20 000 trayectorias sintéticas de 501 días remuestreando esa distribución –con VaR/ES fijos en sus valores de entrenamiento, como en un backtest real– bajo dos esquemas, i.i.d. y bloques circulares de 22 días. Bajo el criterio unilateral propio de Acerbi–Székely (H_1 : las pérdidas de cola realizadas superan el ES predicho, la dirección que importa al regulador), el test **no se rechaza en ninguno de los cuatro niveles α bajo ninguno de los dos esquemas**: el p -valor unilateral nunca baja de 0,91 y vale, en $\alpha = 5\%$, 0,997 (i.i.d.) y 0,949 (bloques). La versión bilateral es menos uniformemente favorable –en $\alpha = 5\%$ rechaza bajo i.i.d. ($p = 0,006$) pero no bajo bloques ($p = 0,103$), el esquema más defendible aquí por preservar la agrupación de volatilidad–. Ambas lecturas apuntan en la misma dirección: el estadístico observado cae en la cola *buena* de la distribución nula –las pérdidas medias en los días de excepción son solo el 47% de la masa de ES presupuestada, con 16 excepciones frente a 25,05 esperadas–, es decir, un modelo que **sobreestima** las pérdidas de cola en vez de subestimarlas.

En síntesis, la batería comprueba tres cosas progresivamente más finas: Kupiec, ¿es correcto el número total de fallos?; Christoffersen, ¿están bien repartidos en el tiempo?; Acerbi–Székely, ¿es correcta la magnitud de la pérdida en esos fallos? Los tres se aplican después a la cartera cuántica igual que a una optimizada por cualquier método clásico.

5.2. Validación in-sample: ¿optimiza bien el pipeline cuántico?

La primera comprobación, previa a lo regulatorio, es la más directa: sobre el periodo de entrenamiento, ¿reduce el optimizador híbrido el CVaR de la cartera, y cómo se compara con los métodos clásicos estándar? La cartera de referencia es PA, con nivel de cola $\alpha = 5\%$.

La Tabla 5.1 admite tres lecturas. El optimizador híbrido **no** gana a los solvers clásicos: queda 1,4 puntos porcentuales por debajo de SLSQP y COBYLA –ambos alcanzan la misma solución, 27,0%– e iguala aproximadamente a Markowitz, brecha atribuible al suelo de precisión $\varepsilon = 0,005$ y a cierta no convexidad del paisaje del gradiente ruidoso. Sí supera con claridad a los métodos más simples –Monte Carlo, paridad de riesgo, subgradiente clásico–, señal de que resuelve el problema en vez

Método	CVaR	VaR	Mejora	Tiempo (s)
SLSQP	0.02717	0.01598	27.0 %	0.02
COBYLA	0.02717	0.01615	27.0 %	0.45
Markowitz	0.02790	0.01644	25.1 %	0.01
Monte Carlo	0.02891	0.01772	22.4 %	0.87
Paridad de riesgo	0.03276	0.01931	12.0 %	0.004
Subgradiente	0.03362	0.02014	9.7 %	0.03
Híbrido cuántico	0.02769	0.01620	25.6 %	2,157
<i>Pesos iguales</i>	<i>0.03723</i>	<i>0.02205</i>	—	—

Tabla 5.1: Comparación *in-sample* sobre el periodo de entrenamiento; la mejora se calcula respecto a pesos iguales (CVaR = 0,03723). El híbrido cuántico queda a 1,4 puntos porcentuales de SLSQP e iguala a Markowitz; el suelo de precisión es el objetivo $\varepsilon = 0,005$. Tabla C.1 del paper, verificada contra `exp_A1_PA`. El 0,02769 es el `cvar_loss` del penúltimo iterado por el desfase de contabilidad; el canónico recalculado es 0,02768, sin alterar el orden de los métodos.

de producir ruido. Y el coste en tiempo, 2,157 s frente a 0,02 s de SLSQP, es apropiado para una demostración en hardware NISQ, no para producción.

5.3. Escalabilidad a anchura de circuito fija

Una propiedad estructural del pipeline es que el circuito de IQAE usa siempre cuatro qubits de distribución más un ancilla con independencia de N : codifican la distribución de pérdida discretizada, no un activo por qubit. Eso permite preguntar cómo escala el resto del pipeline –tiempo clásico, calidad de la solución– cuando N crece sin que el circuito se ensanche.

N	Mejora CVaR	CVaR OOS	Sharpe OOS	HHI	Tiempo (s)	Consultas al oráculo (M, sim.)
10	25.6 %	0.01839	0.681	0.210	2,157	4.4
30	35.5 %	0.01927	0.979	0.118	5,173	24.1
100	37.9 %	0.01795	1.099	0.091	18,215	125.0

Tabla 5.2: Escalabilidad de PA para $N \in \{10, 30, 100\}$ con el mismo circuito de cuatro qubits de distribución más un ancilla (Tabla C.2 del paper). El HHI es el índice de concentración sobre contribuciones de Euler, no sobre pesos.

La columna HHI se calcula sobre las contribuciones porcentuales al riesgo de la descomposición de Euler, no sobre los pesos ($\text{HHI} = \sum_i (\text{CR}_i / \text{CVaR})^2$): vale 1 si todo el riesgo lo aporta un activo y $1/N$ si se reparte por igual –verificado contra `herfindahl_index`, mientras que la suma directa de los pesos al cuadrado daría 0,209/0,119/0,091–. Su descenso con N solo refleja más activos entre los que repartir el riesgo; más relevante es que la mejora de CVaR y el Sharpe OOS crecen de forma monótona con N , lo que sugiere –sin poder afirmarlo con tres puntos– más margen en carteras grandes.

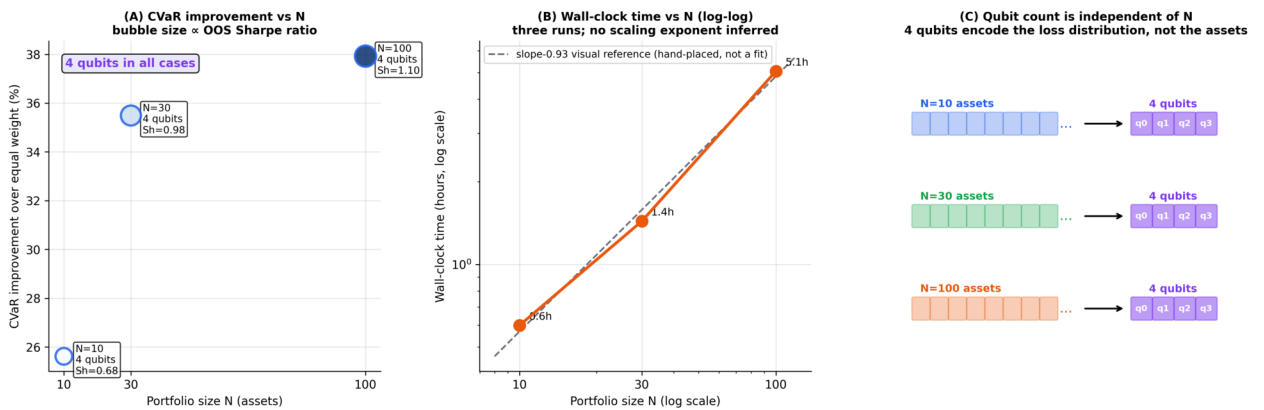


Figura 5.1: Escalabilidad de PA para $N \in \{10, 30, 100\}$ con circuito de anchura fija. Figura C.1 del paper.

Con solo tres valores de N , sin embargo, **no se ajusta ningún exponente de escalado**: la línea de pendiente 0,93 del panel central de la Figura 5.1 es una referencia visual colocada a mano, no un ajuste de regresión. Esta cautela se retoma en el Capítulo 6 como una de las amenazas a la validez del trabajo.

5.4. Validación fuera de muestra

Todo lo anterior se calcula sobre el periodo que el optimizador ya vio. La pregunta que importa es si esa ventaja se mantiene sobre datos nunca vistos. La ventana OOS cubre 501 días, del 2 de enero de 2024 al 30 de diciembre de 2025 –un único régimen alcista impulsado por el auge de la inteligencia artificial, no una muestra de regímenes distintos–. Sobre ella la cartera cuántica alcanza el CVaR OOS más bajo de todos los métodos comparados (0,01839, un 2,5 % menos que SLSQP), con una contrapartida clara en rentabilidad ajustada a riesgo.

Método	CVaR	Sharpe	Rent. anual.	Vol. anual.	Máx. drawdown
CVaR cuántico	0.01839	0.681	9.2 %	13.5 %	−14,4 %
SLSQP clásico	0.01887	0.783	11.1 %	14.1 %	−14,8 %
Inversa de volatilidad	0.02467	0.026	0.4 %	16.9 %	−21,4 %
Pesos iguales	0.02703	−0,076	−1,4 %	18.2 %	−24,9 %

Tabla 5.3: Rendimiento OOS de PA (501 días, 2024–2025). Estimación puntual $\Delta CVaR_{Q-C}^{obs} = -0,00047$ (−2,5 % relativo); *bootstrap* de bloques circulares ($B = 10,000$, bloque de 21 días, semilla 42): media −0,00049, IC 95 % = [−0,00138, +0,00032], $p = 0,121$ a una cola. Tabla D.1 del paper; Sharpe con tasa libre de riesgo = 0.

El Sharpe de la cartera cuántica (0,681) queda por debajo del de SLSQP (0,783): el optimizador cuántico minimiza pérdida de cola, no rentabilidad ajustada a riesgo, y en un mercado alcista sostenido los métodos con algo más de volatilidad capturaron más subida. No contradice el resultado de CVaR –son objetivos distintos–, pero sí es una limitación real.

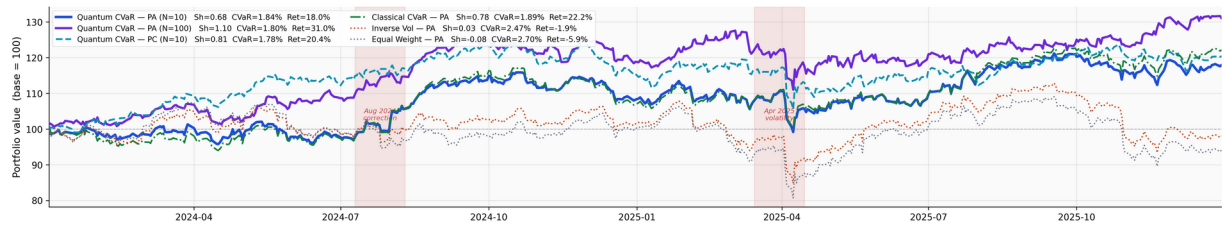


Figura 5.2: Riqueza acumulada fuera de muestra (base 100) sobre los 501 días de 2024–2025. Seis trayectorias: los cuatro métodos sobre PA de la Tabla 5.3 más dos ejecuciones cuánticas de contexto (S1_PA_100 y A1_PC), para leer el efecto del universo sobre los mismos ejes. Todos los métodos optimizados por CVaR superan a pesos iguales; la trayectoria cuántica con $N = 100$ termina en 131,0. Las bandas marcan la corrección de agosto de 2024 y el repunte de abril de 2025. Figura 13 del paper.

5.4.1. Qué dice, y qué no dice, el $p \approx 0,12$

El $p \approx 0,12$ de la Tabla 5.3¹ es un test *a una cola*: comprueba si el CVaR OOS de PA es *menor* que el de SLSQP, en el universo QUBO-periférico concreto de la Etapa 1. Un test a una cola es por construcción más fácil de superar que uno a dos colas, y aun así $p \approx 0,12$ queda muy por encima de 0,05: ausencia de significación, no significación marginal. Más relevante todavía, agregando los 46 experimentos el optimizador cuántico gana al mejor método clásico en 25 (54,3 %), tasa que un binomial a una cola sitúa en $p = 0,329$, tampoco significativo. La ventaja puntual de PA no se generaliza: ambos p -valores apuntan a que se concentra en su universo, así que el motor dominante es la *construcción del universo* y no la estimación del gradiente por IQAE *per se*. Este trabajo lee el resultado OOS como *sugerente, no concluyente*.

5.5. Resultado honesto: la batería de Basilea IV falla parcialmente

Con los tres tests ya definidos (Sección 5.1) y los datos OOS ya presentados (Sección 5.4), esta sección aplica la batería completa a la cartera PA optimizada cuánticamente sobre la misma ventana de 501 días. El resultado se presenta sin suavizar: **el modelo no supera la batería completa de tres tests**, y solo uno —el más directamente alineado con la especificación de CVaR/ES— pasa en los cuatro niveles de confianza probados.

α	Exc.	Tasa	Semáforo	Kupiec	Christ. CC	Acerbi–S.
1.0 %	1	0.20 %	Verde	Falla*	Pasa	Pasa
2.5 %	3	0.60 %	Verde	Falla*	Falla [†]	Pasa
5.0 %	16	3.19 %	Amarillo	Falla*	Falla [†]	Pasa
10.0 %	38	7.58 %	Rojo	Pasa	Pasa	Pasa

En $\alpha = 5\%$, signo corregido: $Z_2 = +0,528$ (el registro archivado citaba $-0,528$); test unilateral frente a hipótesis nula genuina simulada por remuestreo (20 000 réplicas): $p = 0,997$ (i.i.d.) / 0,949 (bloques de 22 días), no rechazado en los cuatro α (detalle completo

en la Tabla 5.5). Sustituye al “ $p = 0,47$ ” archivado, artefacto de un *bootstrap* centrado por construcción en el propio valor observado (Sección 5.1.3).

Tabla 5.4: Tests regulatorios de Basilea IV para la cartera PA optimizada cuánticamente (501 días OOS). *Fallo conservador: la tasa observada queda por debajo de α . [†]Agrupación temporal leve ($\hat{\pi}_{11} = 0,67$ en $\alpha = 2,5\%$, 0,19 en 5 %). Tabla C.3 del paper. La columna Acerbi–Székely usa el test *unilateral* bajo la hipótesis nula genuina de la Sección 5.1.3, no el “ $p = 0,47$ ” archivado.

Las cifras completas de Acerbi–Székely en los cuatro niveles de confianza — Z_2 , Z_1 y los p -valores unilaterales y bilaterales bajo ambos esquemas de remuestreo— se recogen en la Tabla 5.5.

¹El $p = 0,043$ citado en versiones anteriores de este TFM no resultó reproducible —ningún artefacto de *results/* lo sustenta— y se ha sustituido por el valor verificado tras corregir dos errores del script generador (nombres de campo de metadatos incorrectos y ausencia de subselección de *tickers*): con los parámetros por defecto ($B = 10,000$, bloque de 21 días, semilla 42) la ejecución exacta da $p = 0,1209$, con rango $[0,1096, 0,1266]$ en una rejilla equivalente.



Figura 5.3: Backtesting regulatorio de Basilea IV para la cartera cuántica PA. (A) excepciones de VaR a $\alpha = 5\%$: 16 (3,19 %, por debajo del 5 % nominal). (B) ES predicho frente a realizado por α , con el semáforo (clasificación heurística, no la regla binomial). (C) resumen de los tests en los cuatro niveles. (D) estadístico Z_2 con el signo corregido; no se rechaza bajo el test unilateral en ningún nivel. Figura 12 del paper.

α	Z_2	Z_1	p unilateral (i.i.d./bl. 22)	p bilateral (i.i.d./bl. 22)	Acerbi-S.
1.0 %	+0,816	+0,075	0,987 / 0,915	0,026 / 0,171	Pasa
2.5 %	+0,796	+0,146	0,999 / 0,983	0,001 / 0,035	Pasa
5.0 %	+0,528	+0,261	0,997 / 0,949	0,006 / 0,103	Pasa
10.0 %	+0,397	+0,204	0,997 / 0,956	0,005 / 0,088	Pasa

Tabla 5.5: Z_2 y Z_1 de Acerbi-Székely –signo corregido– y p -valores frente a 20 000 trayectorias sintéticas de 501 días remuestreadas de la distribución empírica de entrenamiento, bajo remuestreo i.i.d. y de bloques circulares de 22 días. El unilateral no se rechaza en ningún nivel; el bilateral rechaza bajo i.i.d. en los cuatro pero solo en $\alpha = 2,5\%$ bajo bloques.

Léase la Tabla 5.4 con calma: el patrón que muestra es el opuesto de lo esperable si el objetivo fuera “aprobar” la batería a toda costa. **Kupiec falla en tres de los cuatro niveles** ($\alpha = 1\%$, $2,5\%$, 5%) en la dirección *conservadora* –tasa observada por debajo de la nominal, $0,20\%$ frente a $1,0\%$ en el nivel más exigente–, y pasa en $\alpha = 10\%$ ($p = 0,061$) por margen estrecho, con $\alpha = 5\%$ fallando por muy poco ($p = 0,0475$, justo bajo el corte). Un modelo que falla por exceso de cautela no es peligroso para el banco –sobrestima el riesgo–, pero tampoco es correcto en el sentido estadístico que el test exige. **Christoffersen pasa en $\alpha = 1\%$ y $\alpha = 10\%$ y falla en $\alpha \in \{2,5\%, 5\%\}$** por agrupación temporal

leve; que pase en $\alpha = 10\%$ pese al semáforo Rojo de ese mismo nivel no es contradicción, porque evalúan cosas distintas –independencia temporal frente a recuento con umbrales heurísticos–. Y **solo Acerbi–Székely no se rechaza en los cuatro niveles**, en su versión unilateral, con $p \geq 0,91$ bajo ambos esquemas de remuestreo (Tabla 5.5); la bilateral es menos favorable –rechaza a $\alpha = 5\%$ con i.i.d. pero no con bloques de 22 días–, y ambas lecturas coinciden en que el modelo sobreestima las pérdidas de cola.

De este patrón se derivan varias consecuencias, ninguna suavizada respecto al paper original. Los fallos de Kupiec y Christoffersen, conservadores en dirección, se leen mejor como artefacto del régimen OOS –un único periodo alcista– que como mala especificación patológica, lectura consistente con que ambos pasen en $\alpha = 10\%$.

Además, el semáforo de la Tabla 5.4 no aplica la regla propia de Basilea: la implementación de este trabajo reescala linealmente a 501 días los umbrales de conteo que Basilea define para el VaR al 99 % sobre 250 días (Verde hasta 4 excepciones, Amarillo hasta 9, Rojo desde 10) y los aplica por igual a los cuatro niveles α –extensión heurística propia, no la regla binomial real (Amarillo si la probabilidad acumulada supera 95 %; Rojo si supera 99,99 %). Bajo esa regla binomial, con $T = 501$ días, los umbrales Amarillo/Rojo serían 9/15 excepciones en $\alpha = 1\%$, 19/27 en 2,5 %, 33/45 en 5 % y 61/77 en 10 %: con las 16 excepciones observadas en $\alpha = 5\%$ –*muy* por debajo de 33/45– **los cuatro niveles clasificarían como Verde**, porque los recuentos caen en la cola *inferior* de la binomial, el lado que no penaliza. Corrigiendo además la escala de multiplicadores a las cifras de Basilea IV/FRTB realmente en vigor (BCBS 2019 [6], MAR32.9/MAR99: Verde 1,50×, Amarillo 1,70–1,92×, Rojo 2,00× –**no** la escala de Basilea II de 1996 ya discontinuada, citada por error en una versión anterior de este TFM–), la clasificación heurística implicaría 1,70–1,92× en $\alpha = 5\%$ y 2,00× en $\alpha = 10\%$, mientras que bajo la regla binomial real ningún nivel supera Verde y no se aplicaría recargo más allá de 1,50×. Las entradas Amarilla y Roja son, por tanto, un artefacto de aplicar umbrales calibrados a la cola del 1 % también a las del 5 % y el 10 %.

En síntesis, la conclusión de este TFM es deliberadamente estrecha: el modelo optimizado **no** muestra mala especificación patológica en la métrica más directamente alineada con la calidad del riesgo de cola –el test unilateral de Acerbi–Székely–, pero **no supera la batería completa de tres tests**, y **no se reclama ninguna superioridad regulatoria sobre esta base**: ni frente a un optimizador clásico de CVaR, que se enfrentaría al mismo patrón de fallo conservador en esta misma ventana alcista, ni en ningún otro sentido.

DISCUSIÓN Y LIMITACIONES

Este capítulo reúne, sin diluirlas, las limitaciones que los capítulos anteriores han ido señalando donde aparecían: primero los dos resultados que conviene acotar antes de generalizarlos —el papel real de la selección QUBO y el alcance del resultado de similitud coseno—, y después las siete amenazas a la validez, cada una con qué es, por qué importa y qué haría falta para resolverla.

6.1. La selección QUBO como motor dominante del rendimiento

A lo largo de los 46 experimentos sistemáticos del Capítulo 4, la brecha de rendimiento entre la cartera PA (seleccionada vía QUBO, maximizando perifericidad de red) y las alternativas PB (selección por ranking central) y PC (selección manual de mega-capitalización) es más ancha que cualquier brecha inducida por ruido dentro de un universo fijo. La cadena causal es estructural, no cualitativa, y se sigue cifra a cifra.

El objetivo QUBO de la Etapa 1 —definido formalmente en el Capítulo 2— maximiza la perifericidad de cada activo en la red de correlación y exige cobertura de comunidades (que los activos no se concentren en un mismo clúster). El resultado es un universo diversificado en su cola de riesgo, no solo en varianza promedio, donde la mayoría de los peores días no son todo-negativos.

El contraste con el universo central PB es nítido: correlación media entre pares 0,207 frente a 0,710, peores días “todo-negativos” 25,5 % frente a 84,3 %, y *spread* de CVaR marginal 5,54× frente a 2,17×. Ese *spread* es un indicador real de diversificación pero, como precisa la Sección 6.2, no es la cantidad que determina el margen frente al ruido: esa es la brecha *mínima* entre pares, prácticamente idéntica en términos absolutos para ambos universos.

La concentración resultante en PA —tres activos (CLX, CBOE, CHRW) que acumulan el 74 % del riesgo de cola en muestra— no es un artefacto de sobreajuste: se preserva fuera de muestra, con el HHI sobre contribuciones de Euler pasando de 0,210 en muestra a 0,216 fuera. Un HHI estable entre ventanas señala un clúster defensivo genuino, no una coincidencia del periodo de entrenamiento. Y

dentro del mismo universo, el optimizador híbrido y SLSQP convergen –pese a partir de gradientes distintos– a soluciones que colocan entre el 74 % y el 77 % del CVaR en esos mismos tres activos: lo que determina la solución es el paisaje del CVaR marginal del universo seleccionado, no el detalle algorítmico de cómo se llega a él.

6.2. Similitud coseno: alcance del resultado

El Capítulo 4 reporta una similitud coseno superior a 0,9998 entre el gradiente de CVaR sin ruido y el calculado bajo ruido simulado, evaluada en w_0 de cada uno de los 46 experimentos (una única evaluación por corrida, previa a la optimización). Antes de interpretarla como propiedad del algoritmo cuántico: es una **consecuencia matemática** de la estructura del subgradiente histórico de CVaR –derivada en el Capítulo 2 vía Rockafellar–Uryasev– combinada con la geometría del universo PA, **no** un fenómeno específico de la computación cuántica.

La condición exacta –derivada en el Capítulo 4– establece que preservar el orden entre dos activos bajo ruido multiplicativo de nivel ε requiere $|g_i^* - g_j^*| > 2\varepsilon \max_k |g_k^*|$, y es **suficiente**, no necesaria. El umbral crítico $\varepsilon_c = \min_{i \neq j} |g_i^* - g_j^*| / (2 \max_k |g_k^*|)$ vale $\approx 0,0091$ en w_0 tanto para PA como para PB – $\approx 1,83 \times$ el objetivo $\varepsilon = 0,005$ en ambos casos– pese a *spreads* de CVaR muy distintos ($5,54 \times$ frente a $2,17 \times$): lo que gobierna ε_c es la brecha *mínima* entre pares, no el cociente de dispersión, y los datos descartan leer el *spread* más estrecho de PB como falta de margen. En los pesos convergidos ε_c cae a $\approx 0,00096$ (PA) y $\approx 0,0014$ (PB), por debajo del punto de operación, como exige la igualdad KKT de los marginales sobre el símplex.

Esto acota lo que la selección QUBO aporta: es necesaria por la diversificación de cola, pero **no** ensancha el margen frente al ruido –propiedad del paisaje de CVaR marginal compartida por PA y PB: la selección determina cuánto hay que ganar, no si el algoritmo tiene margen para operar–. El $> 0,9998$ confirma que la condición se cumple con margen en el punto de partida, sin delimitar la frontera exacta, matiz retomado como sexta amenaza.

6.3. Siete amenazas a la validez

Para cada amenaza se explica qué es la limitación, **por qué importa** y **qué haría falta** para resolverla, sin suavizar ni omitir ninguna.

Amenaza 1 — Hardware: tamaño de muestra pequeño en Garnet

La comparación entre dispositivos en $m = 1$ descansa sobre solo **dos** ejecuciones independientes en Garnet (error medio 1,84 %) frente a **ocho** agrupadas en Emerald (15,2 %): una diferencia de factor

~ 8 que es una observación de muestra pequeña, no un hecho causal. **Por qué importa:** con dos puntos no se puede distinguir una propiedad real y estable del dispositivo –mejor calibración de puertas de dos qubits, por ejemplo– de la varianza de una muestra diminuta, y como ambos comparten arquitectura la diferencia tampoco es atribuible a la topología. **Qué haría falta:** al menos cuatro ejecuciones en Garnet en $m \in \{0, 1, 2\}$ sobre un par de qubits seleccionado para igualar la fidelidad por qubit de Emerald –única forma de separar el efecto del dispositivo del de la muestra–.

Amenaza 2 — Hardware: suelo de precisión dos órdenes de magnitud por encima del objetivo

El hardware IQAE actual alcanza $\varepsilon_{\text{hw}} \approx 0,20$ frente al objetivo de simulación $\varepsilon = 0,005$. La ventaja $O(\varepsilon^{-1})$ solo se activa para $\varepsilon < 0,05$; a $\varepsilon = 0,20$ bastan ~ 25 muestras para que Monte Carlo iguale esa precisión, así que la ganancia asintótica es hoy irrelevante. **Por qué importa:** es probablemente la limitación más directa –el pipeline completo, no la caracterización aislada de C1, solo se ha validado de extremo a extremo *en simulación*–, de modo que ninguna cifra de mejora de CVaR debe leerse como evidencia de que se obtendría hoy en un procesador físico. **Qué haría falta:** la ganancia teórica completa exige $m = 7$ rondas, pero a esta amplitud de cola esa ronda cae *más allá* del umbral de *aliasing* ($(2m+1)\theta \approx 3,4 \text{ rad} > \pi$, Sección 4.4.3); el camino completo –diferimiento algorítmico de la rama simétrica más fidelidad de puerta superior al 99,9%– se detalla en el punto (a) de la Sección 6.4.

Amenaza 3 — Escalado: la precisión está limitada por discretización en $n = 4$

Se ejecutó en simulación un cuarto valor de precisión ($\varepsilon = 0,001$, A4_PA), pero **no** es comparación controlada frente a A3_PA: se entrenó sobre una ventana más larga y distinta ($T = 1,508$ días frente a 1,005) y convergió en 34 iteraciones frente a 38, de modo que el error de CVaR realizado queda confundido por el cambio de ventana. El argumento sobre las consultas al oráculo sí sobrevive si se elige la medida correcta, y expone un **suelo de discretización**: la media por iteración –la más fiable de las tres medidas archivadas, por no depender del tamaño de ventana– pasa de $1,0419 \times 10^6$ a $1,0428 \times 10^6$ al reducir ε a la mitad (+0,09%), muy por debajo de la duplicación que un escalado $O(\varepsilon^{-1})$ genuino predeciría. **Por qué importa:** el desglose del presupuesto de error revela que en ambas ejecuciones domina el error de **truncación** –la pérdida de resolución de discretizar en $n = 4$ qubits, 16 contenedores–, no el ruido de estimación de amplitud; por debajo de $\varepsilon \approx 0,002$ es la rejilla, no ε , la que fija la precisión alcanzable. **Qué haría falta:** bajar el suelo exige más qubits de registro, pero el coste no es el ingenuo $O(2^n)$ –cada iteración emite un número fijo de $(N+1) \cdot O(\varepsilon^{-1})$ llamadas a IQAE, independiente de n ; lo que escala es solo el coste *clásico* de simular cada llamada–. Las ablaciones A1_n5_PA y A1_n6_PA lo confirman: convergen en las mismas 38 iteraciones y dejan el CVaR prácticamente sin cambios ($0,02768 \rightarrow 0,02769 \rightarrow 0,02770$). Verificar el exponente de Heisenberg exige, por tanto, más qubits y promediar el conteo de consultas sobre semillas, porque el número de rondas es adaptativo y discreto; la misma advertencia aplica a los tres valores de N del análisis de escalado, que motivan pero no establecen una ley de potencia.

Amenaza 4 — Aplicación: el resultado fuera de muestra no es significativo

La tasa de victorias OOS frente a SLSQP a lo largo de todos los experimentos es 25/46 (54,3 %), con p -valor binomial a una cola de 0,329; el resultado específico de PA ($p \approx 0,12$ por *bootstrap*) tampoco lo es. **Por qué importa:** el punto estimado ($-2,5$ % relativo en CVaR OOS) podría sugerir una mejora sólida, y no lo es —ambos p -valores superan ampliamente 0,05—; presentarlo sin este matiz sería sobrevender el resultado. **Qué haría falta:** ventanas OOS más largas y múltiples universos independientes para plantear siquiera una ventaja sistemática, que —donde existe— se concentra en el universo QUBO-periférico (Sección 6.1): el resultado refleja la construcción del universo más que la estimación del gradiente *per se*.

Amenaza 5 — Aplicación: especificidad de régimen fuera de muestra

El periodo OOS, 2024–2025, es un mercado alcista sostenido impulsado por el entusiasmo en torno a la inteligencia artificial: precisamente el régimen en el que minimizar riesgo de cola **menos** se recompensa. **Por qué importa:** cualquier resultado de mejora se ha medido en un único régimen, y no se sabe si se sostiene, se invierte o desaparece en uno bajista o de crisis, que es donde la minimización de riesgo de cola debería demostrar más valor; afirmar que el pipeline “funciona fuera de muestra” sin esta salvedad sería engañoso. **Qué haría falta:** validación por ventana móvil a través de regímenes heterogéneos (COVID-19 de 2020, subida de tipos de 2022), de las siete amenazas el paso siguiente más importante para la aplicación financiera.

Amenaza 6 — Teoría: la condición de brecha de ranking es suficiente, no necesaria

La condición se obtiene vía la desigualdad triangular, conservadora por construcción, así que un universo que la viole podría aun así preservar el orden para las realizaciones de ruido encontradas en la práctica: garantiza el peor caso, no descarta que casos menos adversos funcionen. **Por qué importa:** es además una condición de *inicio de trayectoria*, no de convergencia —la igualdad KKT de los CVaR marginales en el óptimo la hace fallar para ambos universos, y el coseno exportado se evalúa solo en w_0 —, de modo que la preservación de la dirección a lo largo de la trayectoria se infiere, no se mide; la cifra de $> 0,9998$ confirma el régimen suficiente en w_0 , no la frontera exacta de robustez. **Qué haría falta:** un análisis específico de las realizaciones de ruido bajo las que el ranking se preserva incluso violando la condición —problema distinto, fuera del alcance— y persistir el coseno por iteración, que el optimizador calcula pero los artefactos archivados no guardan.

Amenaza 7 — Hardware: umbral de *aliasing* verificado a un solo θ

El Capítulo 4 predice desde la geometría de la cola que el *aliasing* aparece en $m \approx 3$, según $(2m + 1)\theta \approx \pi/2$, y afirma que ese umbral aparece en cualquier plataforma con independencia de la fidelidad. Esa generalidad solo se ha **verificado en hardware para un único punto de operación**, $\theta \approx$

0,2265. **Por qué importa:** la batería de Basilea IV evalúa la cartera a cuatro niveles de confianza y solo $\alpha = 0,05$ se ha ejecutado en hardware; los otros tres implican valores de θ distintos, nunca probados, y no se sabe si el umbral se desplaza a otra ronda ni si la concordancia entre Emerald y Garnet se mantendría. **Qué haría falta:** repetir la caracterización en hardware IQM para $\alpha = 0,01, 0,025, 0,10$ y comprobar en cada caso si el onset predicho se mantiene en $m \approx 3$ o se desplaza.

6.4. Implicaciones y camino hacia una ventaja práctica

Recogidas las siete amenazas, resta extraer qué se puede concluir a pesar de ellas y qué haría falta para que la ventaja práctica deje de ser condicional.

Del lado del hardware

La construcción del universo importa más que el estimador cuántico a las escalas actuales; las futuras mejoras de fidelidad mejorarán la precisión por debajo del umbral, pero solo *dentro* del régimen no aliasado $m \leq 2$, porque el umbral es geométrico y la fidelidad no lo desplaza. Sobre mitigación: el *folding* global resulta contraproducente en IQM a estas profundidades con `opt_level=0`; la inserción de identidad sin *folding* [9] logra la misma mitigación con la mitad de sobrecoste, pero **por debajo** del umbral ningún método de extrapolación mejora la estimación cruda: ahí el presupuesto de disparos rinde más repitiendo la medida que escalando el ruido [31,32].

Implicación regulatoria.

El test *unilateral* de Acerbi–Székely (Z_2 , convención de signo propia de sus autores) no se rechaza para la cartera PA en ninguno de los cuatro niveles de confianza de Basilea IV probados, bajo una hipótesis nula simulada genuinamente (Sección 5.1.3 del Capítulo 5): la especificación de CVaR del modelo no es patológica en la dirección que preocupa al regulador. Kupiec y Christoffersen fallan en varios niveles, siempre de forma conservadora, y el semáforo regulatorio –heurística propia, no la regla binomial de Basilea, que aquí clasificaría los cuatro niveles como Verde– empeora con α (detalle nivel a nivel en la Tabla 5.4 del Capítulo 5), efecto del régimen alcista del periodo OOS (Amenaza 5, Sección 6.3), no de un defecto de especificación. Bajo los multiplicadores de Basilea IV/FRTB en vigor (Verde $1,50\times$, Amarillo $1,70\text{--}1,92\times$, Rojo $2,00\times$), esas zonas implicarían un recargo de capital no trivial bajo la clasificación heurística, frente a ninguno bajo la regla binomial real. En cualquier caso, este trabajo **no reclama ninguna superioridad regulatoria** sobre optimizadores clásicos de CVaR.

Camino hacia una ventaja práctica.

La optimización cuántica de CVaR viable en producción –no solo demostrable en laboratorio– requiere, según el análisis de este trabajo, cuatro avances:

- (a) un diferimiento *algorítmico* del umbral de *aliasing* –amplitud codificada más pequeña, reescalada, o nivel de cola más pequeño (Amenaza 2, Sección 6.3)– junto con una fidelidad de puerta de dos qubits superior al 99,9 % para que los circuitos de $m = 7$ resultantes sigan siendo precisos: un hito que las hojas de ruta de los fabricantes sitúan a finales de esta década [33], pero que depende de factores de calibración y fabricación fuera del control de este trabajo (objetivo de la industria, no previsión propia);
- (b) ZNE eficiente en profundidad y libre de *folding* –inserción de identidad [9] u otros enfoques equivalentes [31], no *folding* global [32]– aplicado en el régimen de ruido *por encima* del umbral de *aliasing*, donde la mitigación realmente ayuda;
- (c) una verificación empírica del exponente $O(\varepsilon^{-1})$ que cubra el rango $\varepsilon \in [0,001, 0,010]$ en hardware por debajo del umbral de *aliasing* –lo que exige resolver primero el suelo de discretización de la Amenaza 3 (Sección 6.3);
- (d) validación fuera de muestra por ventana móvil a través de regímenes de mercado heterogéneos –directamente la Amenaza 5 (Sección 6.3).

La mitad de fidelidad de (a) depende de avances de hardware ajenos a este trabajo; la mitad de reescalado de (a) y los puntos (b)–(d) son experimentos ejecutables sobre el pipeline ya construido. Esa distinción –terceros frente a siguiente paso propio– estructura la agenda de trabajo futuro del Capítulo 7, que retoma los cuatro puntos como propuestas concretas.

CONCLUSIONES Y TRABAJO FUTURO

7.1. Síntesis de resultados

Este capítulo de cierre no repite la derivación de los capítulos anteriores ni añade cifras nuevas: reúne los resultados en torno a las tres preguntas de hardware Q1, Q2 y Q3 de la Sección 1.3.

7.1.1. Dos hallazgos de hardware

El primer hallazgo –Contribución C1 (Secciones 1.5, 4.4.1 y 4.4.2)– es un *umbral de aliasing* robusto en ambos procesadores IQM de este trabajo (Emerald y Garnet), medido en una sonda de dos qubits que reproduce la amplitud de cola del pipeline ($\theta \approx 0,2265$, la del CVaR al 5 % de la cartera PA) pero no su preparación de estado de cinco qubits vía el método de Möttönen, no ejecutada en hardware. Con la fórmula empírica de profundidad $\text{depth}(m) = D_A + 6m$ de la sonda ($D_A = 4$, tres CZ por ronda en puertas nativas de IQM), la Sección 4.4.2 muestra que el número de rondas de Grover utilizable dentro de IQAE no está limitado por la fidelidad del hardware, sino por una ambigüedad de fase, el *aliasing*, predicha analíticamente por la geometría de la cola ($(2m+1)\theta \approx \pi/2$): ocurre en $m \approx 3$ para $\theta \approx 0,2265$. Como el umbral depende solo de la amplitud codificada, la conclusión se transfiere al circuito real del pipeline –transpilado *offline*, más de dos órdenes de magnitud más profundo que la sonda en todo m (1,476 en $m = 0$, frente a 4)–. Medidas de una sola ejecución en ambos dispositivos confirman el umbral predicho ($p_{\text{hw}} \approx 0,92$ en $m = 3$), acotando la precisión fiable de la sonda en hardware actual a $\varepsilon_{\text{hw}} \approx 0,20$ –cota inferior no medida del suelo de precisión del circuito real. Por debajo de ese umbral se observa una diferencia entre dispositivos en $m = 1$ (error de Emerald 15,2 %, agregado sobre ocho ejecuciones, frente a Garnet 1,84 % sobre solo dos), que la Sección 6.3 trata como observación de muestra pequeña, no causal: ambos procesadores comparten arquitectura de red cuadrada de IQM, descartando un efecto de familia de topología.

El segundo hallazgo –Contribución C2– responde a Q3, sobre mitigación de errores. La Sección 4.4.4 muestra que el *folding* global de circuito ($\text{opt_level}=0$) –una de las variantes de ZNE comparadas– introduce un sobrecoste de profundidad lógica de $4,5\times$, frente a una alternativa libre de *folding* por inserción de identidad [9] que logra mitigación comparable con aproximadamente la mi-

tad de esa profundidad, en una comparación de hardware posterior. Ambas variantes, sin embargo, producen en hardware real una respuesta no monótona dominada por decoherencia cuya extrapolación diverge por debajo del valor verdadero: por debajo del umbral de *aliasing*, la estimación cruda de IQAE, sin extrapolación alguna, es más precisa que cualquiera de las dos —confirmado en el modelo de ruido de `qiskit-aer` y en hardware Emerald real (Sección 4.4.4).

Robustez al ruido: la brecha de ranking (Q2)

Como soporte de la aplicación, no contribución numerada aparte, la Sección 4.3.2 deriva y verifica una desigualdad de *brecha de ranking*, consecuencia de la desigualdad triangular sobre la estructura de error *multiplicativo* de IQAE: las estimaciones ruidosas preservan el orden entre activos que determina la dirección de actualización de Adam. La similitud coseno entre subgradientes ruidosos e ideales se mantiene por encima de 0,9998 en los 46 experimentos, incluso a un ruido un orden de magnitud por encima de la tasa nominal de IQM. Pero es solo una condición de *inicio de trayectoria* —se viola necesariamente en los pesos convergidos, donde KKT iguala los CVaR marginales— y no la aporta la selección QUBO: el margen es $\varepsilon_c \approx 0,0091$ para PA y para PB, $\approx 1,83 \times$ el objetivo en ambos casos, pese a dispersiones de riesgo marginal muy distintas. Lo que la selección sí aporta, verificado, es diversificación estructural del universo.

7.1.2. Resultados de la aplicación financiera

In-sample, el optimizador híbrido reduce la CVaR de PA un 25,6 % frente a pesos iguales, comparable a Markowitz (25,1 %) y a 1,4 puntos de SLSQP (27,0 %) —diferencia atribuible al suelo $\varepsilon = 0,005$ y a una ligera no convexidad del gradiente ruidoso—, a un coste de 2,157 s frente a 0,02 s, propio de una demostración NISQ y no de un sistema en producción. *Fuera de muestra* (501 días, 2024–2025) alcanza el CVaR más bajo de todos los métodos (0,01839, un 2,5 % menos que SLSQP), pero el respaldo estadístico es débil y eso es tan parte del resultado como la cifra: ni el $p \approx 0,12$ de *bootstrap* ni la tasa de victorias global 25/46 (54,3 %, binomial $p = 0,329$) son significativos. La ventaja, donde aparece, se concentra en el universo QUBO-periférico —el motor es la construcción del universo, no la estimación del gradiente—, y la ventana cubre un único régimen alcista en el que minimizar riesgo de cola está menos recompensado.

La batería de Basilea IV se usa solo como comprobación externa de sensatez. Únicamente el test unilateral de Acerbi–Székely se sostiene en los cuatro niveles de confianza; Kupiec y Christoffersen fallan de forma conservadora en la mayoría, aunque pasan en $\alpha = 10\%$ (Christoffersen también en 1 %), y el semáforo —clasificación heurística propia, no la regla binomial, bajo la cual los cuatro niveles serían Verde a este tamaño muestral— pasa a Amarillo/Rojo para $\alpha \geq 5\%$. No se reclama superioridad regulatoria sobre esta base.

7.2. Trabajo futuro: cuatro condiciones concretas para la ventaja práctica

La Sección 6.4 enumera cuatro condiciones cuya concurrencia *simultánea* –no bastaría con una– sería necesaria para que este pipeline pasara de caracterizar un régimen de validez a ofrecer ventaja práctica. Son las siguientes.

(a) Diferir el umbral geométrico de *aliasing*

El umbral en $m \approx 3$ impide alcanzar $\varepsilon < 0,05$, que exigiría $m \geq 7$ y profundidad de al menos 46 puertas nativas. La causa no es la fidelidad: a $\theta \approx 0,2265$, $m = 7$ cae más allá del cruce geométrico $((2m+1)\theta \approx 3,4 \text{ rad} > \pi)$, fijado por la amplitud codificada y no por el dispositivo. La mitad algorítmica –reescalar la amplitud a a/c [14], o apuntar a un nivel de cola menor– ya es viable hoy; el hito de fidelidad restante lo sitúa el sector a finales de la década de 2020 [33], calendario ajeno a este trabajo.

(b) ZNE eficiente en profundidad, sin *folding* global

Por debajo del umbral ninguna variante de ZNE mejora la estimación cruda: el presupuesto de disparos rinde más repitiendo la medida que amplificando el ruido. Esa conclusión no se ha probado por encima del umbral –hoy inalcanzable–, donde la mitigación podría volver a ser útil; el *folding* global demostró aquí un sobrecoste de $4,5\times$ frente a los esquemas de inserción de identidad [9, 31], que logran mitigación comparable a una fracción de esa sobrecarga.

(c) Verificación empírica de la escala $O(\varepsilon^{-1})$

La ganancia asintótica está demostrada analíticamente y en simulación, pero no medida en hardware. Con el registro de $n = 4$ qubits la precisión queda limitada por discretización, no por ε , por debajo de $\varepsilon \approx 0,002$, y ampliar el registro apenas mueve el CVaR porque el coste añadido es de simulación clásica, no de llamadas a IQAE –fijas en $(N + 1) \cdot O(\varepsilon^{-1})$ por iteración–. Verificarlo exigirá más qubits de registro y promediar las consultas sobre varias semillas, pues el número de rondas es adaptativo.

(d) Validación OOS por ventanas móviles en regímenes heterogéneos

A diferencia de las otras tres, es ejecutable hoy sin avances de hardware. La ventana OOS de este trabajo cubre un único régimen, precisamente el menos favorable para minimizar riesgo de cola, y con los datos actuales tanto $p \approx 0,12$ como $p = 0,329$ son consistentes con una ausencia de ventaja real. Solo una muestra mayor –más regímenes históricos, incluidos los bajistas (COVID-19 en 2020, subidas de tipos en 2022), y más universos independientes– permitiría distinguirla de una ventaja que la potencia estadística actual no alcanza a detectar.

Estas cuatro condiciones no agotan el trabajo futuro del Capítulo 6 –repetir la muestra de Garnet con fidelidad igualada, o extender el umbral de *aliasing* a los otros niveles de Basilea IV–, pero son las que, cumplidas simultáneamente, demostrarían ventaja práctica.

7.3. Las tres preguntas del Capítulo 1, respondidas

Q1: ¿cuántas rondas de Grover aguanta el hardware? Hasta $m = 2$ de forma fiable; en $m = 3$ el umbral de *aliasing* inutiliza ambos procesadores IQM (Contribución C1, con $\text{depth}(m) = D_A + 6m$ confirmada en Emerald y Garnet). El matiz: solo se ha verificado para el nivel de cola del 5 % de PA, no para los demás niveles de Basilea IV (Sección 6.3).

Q2: ¿se mantiene el orden relativo del subgradiente bajo ruido NISQ? Sí, en el punto de partida: margen $\approx 1,83\times$ la precisión objetivo –igual en PA y PB, así que no es atribuible al QUBO– y coseno $> 0,9998$ en los 46 experimentos. El matiz: es una condición suficiente, no necesaria, y solo de inicio de trayectoria; se viola en los pesos convergidos, donde KKT iguala los CVaR marginales.

Q3: ¿mejora ZNE la precisión de IQAE? No por debajo del umbral de *aliasing*, que es el único régimen accesible hoy: ninguna variante mejora la estimación cruda y el *folding* global resulta contraproducente. El matiz: por encima del umbral, inalcanzable y no probado aquí, es el objeto de la condición (b) (Sección 7.2).

En conjunto, las tres respuestas confirman el marco que la Sección 1.7 ya adelantaba: una caracterización experimental de *dónde* funciona IQAE en hardware IQM real, *dónde* deja de hacerlo y por qué –no una demostración de ventaja computacional cuántica ya vigente para la optimización de carteras bajo riesgo CVaR–; las cuatro condiciones de la Sección 7.2 miden la distancia concreta que aún la separa.

Disponibilidad de datos y código

Todos los datos de retornos, los registros de medición en hardware IQM, las formulaciones QUBO para D-Wave, las especificaciones de los modelos de ruido y los registros de optimización de cada uno de los 46 experimentos están públicamente disponibles en el repositorio del proyecto, junto con el código fuente completo que genera cada figura, tabla y cifra de esta memoria:

<https://github.com/NachoLopezLeis/QuantumCVaRPortfolioOptimization-TFMIgnacioLopez>. Todas las rutas de fichero citadas a lo largo del documento –`results/exp_<EXP>/`, `config/*.yaml`, `results/hardware_validation/`, `results/euler_decomposition/` y los guiones de `scripts/`– son relativas a la raíz de ese repositorio. Una web complementaria resume el pipeline, los resultados y la comparación estructurada con el trabajo previo:

`https:`
`//nacholopezleis.github.io/QuantumCVaRPortfolioOptimization-TFMIgnacioLopez/.`

BIBLIOGRAFÍA

- [1] S. Woerner and D. J. Egger, “Quantum risk analysis,” *npj Quantum Inf.*, vol. 5, p. 15, 2019.
- [2] S. Herbert, R. Guichard, and D. Ng, “Noise-aware quantum amplitude estimation,” *IEEE Trans. Quantum Eng.*, vol. 5, pp. 1–12, 2024.
- [3] D.-L. Vu, B. Cheng, and P. Rebentrost, “Low-depth amplitude estimation without really trying,” *ACM Trans. Quantum Comput.*, vol. 6, no. 4, pp. 29:1–29:23, 2025.
- [4] S. Dasgupta and T. S. Humble, “Impact of unreliable devices on stability of quantum computations,” *ACM Trans. Quantum Comput.*, vol. 5, no. 4, pp. 22:1–22:22, 2024.
- [5] E. Pelofske, V. Russo, R. LaRose, A. Mari, D. Strano, A. Bäertschi, S. Eidenbenz, and W. J. Zeng, “Increasing the measured effective quantum volume with zero noise extrapolation,” *ACM Trans. Quantum Comput.*, vol. 5, no. 3, pp. 21:1–21:18, 2024.
- [6] Basel Committee on Banking Supervision, “Minimum capital requirements for market risk,” tech. rep., Bank for International Settlements, 2019.
- [7] V. Skarlatos and N. Konofaos, “Quantum subgradient estimation for conditional value-at-risk optimization,” 2025. Preprint; v2 (15 May 2026).
- [8] C. Acerbi and B. Székely, “Back-testing expected shortfall,” *Risk Mag.*, vol. 27, no. 11, pp. 76–81, 2014.
- [9] A. He, B. Nachman, W. A. de Jong, and C. W. Bauer, “Zero-noise extrapolation for quantum-gate error mitigation with identity insertions,” *Phys. Rev. A*, vol. 102, no. 1, p. 012426, 2020.
- [10] L. Braine, D. J. Egger, J. Glick, and S. Woerner, “Quantum algorithms for mixed binary optimization applied to transaction settlement,” *IEEE Trans. Quantum Eng.*, vol. 2, pp. 1–8, 2021.
- [11] G. Carrascal, P. Hernamperez, G. Botella, and A. del Barrio, “Backtesting quantum computing algorithms for portfolio optimization,” *IEEE Trans. Quantum Eng.*, vol. 5, pp. 1–20, 2024.
- [12] G. Agliardi, D. Alevras, V. Kumar, R. Lo Nardo, G. Compostella, S. Kumar, M. Proissl, and B. Mehta, “Portfolio construction using a sampling-based variational quantum scheme.” arXiv preprint, 2025. arXiv:2508.13557 [quant-ph].
- [13] D. J. Egger, C. Gambella, J. Marecek, S. McFaddin, M. Mevissen, R. Raymond, A. Simonetto, S. Woerner, and E. Yndurain, “Quantum computing for finance: State-of-the-art and future prospects,” *IEEE Trans. Quantum Eng.*, vol. 1, pp. 1–24, 2020.
- [14] N. Stamatopoulos, D. J. Egger, Y. Sun, C. Zoufal, R. Iten, N. Shen, and S. Woerner, “Option pricing using quantum computers,” *Quantum*, vol. 4, p. 291, 2020.
- [15] G. Brassard, P. Høyer, M. Mosca, and A. Tapp, “Quantum amplitude amplification and estimation,” *Contemp. Math.*, vol. 305, pp. 53–74, 2002.
- [16] D. Grinko, J. Gacon, C. Zoufal, and S. Woerner, “Iterative quantum amplitude estimation,” *npj Quantum Inf.*, vol. 7, p. 52, 2021.

- [17] P. Artzner, F. Delbaen, J.-M. Eber, and D. Heath, "Coherent measures of risk," *Math. Financ.*, vol. 9, no. 3, pp. 203–228, 1999.
- [18] R. T. Rockafellar and S. Uryasev, "Optimization of conditional value-at-risk," *J. Risk*, vol. 2, no. 3, pp. 21–41, 2000.
- [19] D. Tasche, "Expected shortfall and beyond," *J. Banking Financ.*, vol. 26, no. 7, pp. 1519–1533, 2002.
- [20] D-Wave Systems, "Hybrid solver service: An overview," white paper, D-Wave Systems Inc., 2020.
- [21] C. C. McGeoch and P. Farré, "Milestones on the quantum utility highway: Quantum annealing case study," *ACM Trans. Quantum Comput.*, vol. 5, no. 1, pp. 1:1–1:28, 2024.
- [22] S. Harwood, C. Gambella, D. Trenev, A. Simonetto, D. Bernal Neira, and D. Greenberg, "Formulating and solving routing problems on quantum computers," *IEEE Trans. Quantum Eng.*, vol. 2, pp. 1–17, 2021.
- [23] R. N. Mantegna, "Hierarchical structure in financial markets," *Eur. Phys. J. B*, vol. 11, no. 1, pp. 193–197, 1999.
- [24] F. Pozzi, T. Di Matteo, and T. Aste, "Spread of risk across financial markets: Better to invest in the peripheries," *Sci. Rep.*, vol. 3, p. 1665, 2013.
- [25] M. E. J. Newman, "Fast algorithm for detecting community structure in networks," *Phys. Rev. E*, vol. 69, no. 6, p. 066133, 2004.
- [26] M. Grande and J. Borondo, "Embedding pairs trading in market networks: a network science approach to portfolio construction," *Humanit. Soc. Sci. Commun.*, vol. 12, no. 1, pp. 1–11, 2025.
- [27] D. P. Kingma and J. Ba, "Adam: a method for stochastic optimization," in *Int. Conf. Learn. Represent. (ICLR)*, 2015.
- [28] D. N. Politis and J. P. Romano, "A circular block-resampling procedure for stationary data," in *Exploring the Limits of Bootstrap* (R. LePage and L. Billard, eds.), pp. 263–270, New York: Wiley, 1992.
- [29] P. H. Kupiec, "Techniques for verifying the accuracy of risk measurement models," *J. Derivatives*, vol. 3, no. 2, pp. 73–84, 1995.
- [30] P. F. Christoffersen, "Evaluating interval forecasts," *Int. Econ. Rev.*, vol. 39, no. 4, pp. 841–862, 1998.
- [31] H. P. Patil, P. Li, J. Liu, and H. Zhou, "Folding-free ZNE: a comprehensive quantum zero-noise extrapolation approach for mitigating depolarizing and decoherence noise," in *Proc. IEEE Int. Conf. Quantum Comput. Eng. (QCE)*, pp. 886–895, 2023.
- [32] T. Giurgica-Tiron, Y. Hindy, R. LaRose, A. Mari, and W. J. Zeng, "Digital zero noise extrapolation for quantum error mitigation," in *Proc. IEEE Int. Conf. Quantum Comput. Eng. (QCE)*, pp. 306–316, 2020.
- [33] IBM Quantum, "IBM quantum development roadmap." <https://www.ibm.com/quantum/roadmap>, 2023. Accessed: April 2026.

ACRÓNIMOS

CVaR Conditional Value-at-Risk.

IQAE Iterative Quantum Amplitude Estimation.

MST Minimum Spanning Tree.

NISQ Noisy Intermediate-Scale Quantum.

QAE Quantum Amplitude Estimation.

QUBO Quadratic Unconstrained Binary Optimization.

TMFG Triangulated Maximally Filtered Graph.

VaR Value-at-Risk.

ZNE Zero-Noise Extrapolation.

APÉNDICES

APÉNDICES

Este capítulo reúne el material de trazabilidad del TFM que no forma parte de su argumento principal: las tablas completas de resultados por experimento, la descomposición de Euler del riesgo de cartera, la configuración de parámetros de los 46 experimentos y los identificadores de trabajos en hardware IQM (Apéndices A, B, F y G del paper, respectivamente). Los Apéndices C, D y E del paper –aplicación financiera, rendimiento fuera de muestra y pseudocódigo– ya están incorporados a los Capítulos 5 y 3 y no se repiten aquí.

A.1. Resultados completos por experimento

La Tabla A.1 recoge el CVaR dentro de muestra (*in-sample*, IS), la mejora de CVaR respecto a la cartera de partida y las métricas fuera de muestra (*out-of-sample*, OOS).

Experimento	Cartera	CVaR IS	Mejora CVaR	CVaR OOS	Sharpe OOS	Experimento	Cartera	CVaR IS	Mejora CVaR	CVaR OOS	Sharpe OOS
A1_PA	PA	0.02768	25.6 %	0.01839	0.681	D2_PA	PA	0.02952	20.7 %	0.02255	0.686
A1_PB	PB	0.03834	19.5 %	0.02357	0.660	D2_PB	PB	0.03844	19.3 %	0.02404	0.531
A1_PC	PC	0.03105	11.6 %	0.01779	0.807	D2_PC	PC	0.03173	9.7 %	0.01668	0.827
A1_n3_PA	PA	0.02768	25.6 %	0.01838	0.686	E1_PA	PA	0.02776	25.4 %	0.01912	0.621
A1_n5_PA	PA	0.02769	25.6 %	0.01839	0.682	E1_PB	PB	0.03841	19.4 %	0.02371	0.664
A1_n6_PA	PA	0.02770	25.6 %	0.01835	0.686	E1_PC	PC	0.03119	11.2 %	0.01696	0.825
A3_PA	PA	0.02768	25.6 %	0.01837	0.688	E2_PA	PA	0.02927	21.3 %	0.02227	0.673
A3_PB	PB	0.03834	19.5 %	0.02358	0.660	E2_PB	PB	0.03876	18.7 %	0.02360	0.662
A3_PC	PC	0.03106	11.6 %	0.01782	0.808	E2_PC	PC	0.03127	11.0 %	0.01702	0.846
B2_PA	PA	0.03723	0.0 %	0.02703	-0.076	E3_PA	PA	0.02955	20.6 %	0.02261	0.672
B2_PB	PB	0.04155	12.9 %	0.02464	0.778	E3_PB	PB	0.03842	19.4 %	0.02402	0.536
B2_PC	PC	0.03255	7.4 %	0.02356	0.245	E3_PC	PC	0.03169	9.8 %	0.01670	0.826
C2_PA	PA	0.03022	18.9 %	0.02761	-0.005	S1_PA_30	PA ₃₀	0.02470	35.5 %	0.01927	0.979
C2_PB	PB	0.04041	15.2 %	0.02373	0.818	S1_PA_100	PA ₁₀₀	0.02395	37.9 %	0.01795	1.099
C2_PC	PC	0.03509	0.1 %	0.01938	1.247	S1_PB_30	PB ₃₀	0.03545	17.2 %	0.01856	0.582
D1_PA	PA	0.02950	20.7 %	0.02251	0.689	S1_PB_100	PB ₁₀₀	0.03313	26.8 %	0.01967	0.973
D1_PB	PB	0.03844	19.3 %	0.02404	0.531	S1_PC_30	PC ₃₀	0.03534	17.8 %	0.01892	0.680
D1_PC	PC	0.03172	9.7 %	0.01668	0.826	S1_PC_100	PC ₁₀₀ *	0.02562	26.8 %	0.01721	0.954
A4_PA [†]	PA	0.02505	24.3 %	0.01859	0.731						

[†]A4_PA se entrenó sobre 2018-01-03–2023-12-29 ($T = 1,508$ días, CVaR inicial de pesos iguales 0,033087), no sobre la ventana 2020-01-03–2023-12-29/ $T = 1,005$ días que usan las demás 36 filas de esta tabla, incluida A3_PA (Sección 6.3 del Capítulo 6).

Tabla A.1: Métricas clave de las ejecuciones de optimización. Los experimentos se agrupan por familia; PA/PB/PC son los tres universos de activos; S1 son los de escalabilidad ($N \in \{10, 30, 100^*\}$, con *S1_PC_100 en realidad $N = 96$ tickers; S1_PA_100 y S1_PB_100 sí usan 100). Por compacidad, esta tabla recoge las 36 configuraciones principales más, como fila 37, la ablación separada A4_PA ($\varepsilon = 0,001$), externa a la batería canónica de 46; las diez ablaciones adicionales de n_q sobre PB/PC (A1) y sobre PA (B2, D1) que completan esa batería se omiten aquí y están disponibles en el repositorio público (Sección 7.3), ejecución a ejecución (`results/exp_<EXP>/`).

A.2. Descomposición de Euler del riesgo de cartera: cartera PA

La descomposición de Euler del riesgo –contribución al riesgo (CR) por activo e índice de concentración HHI, definidos en la Sección 4.2.2 del Capítulo 4– se detalla aquí para el experimento A1_PA bajo el optimizador híbrido cuántico. La Tabla A.2 recoge las diez contribuciones individuales.

Activo	Peso	CVaR marg.	CR (pérdida)	CR %
CLX	0.279	0.0302	0.00845	30.52 %
CBOE	0.290	0.0264	0.00764	27.62 %
CHRW	0.178	0.0260	0.00463	16.73 %
NEM	0.108	0.0234	0.00252	9.09 %
DXCM	0.040	0.0263	0.00105	3.80 %
FMC	0.033	0.0266	0.00088	3.18 %
OXY	0.016	0.0464	0.00073	2.62 %
AAL	0.019	0.0342	0.00065	2.34 %
NFLX	0.019	0.0305	0.00058	2.10 %
AAP	0.018	0.0304	0.00055	1.98 %
Total	1.000	—	0.02768 ^a	100.00 %

^aEl `cvar_loss` archivado (0,02769) difiere del valor canónico recalculado a partir de los pesos reportados (0,02768) por un desfase de contabilidad en el seguimiento del mejor-hasta-ahora del optimizador, no por un efecto de *vintage* de datos (causa raíz completa en la Sección 4.2.2 del Capítulo 4).

Tabla A.2: Descomposición de Euler del CVaR, cartera PA, optimizador híbrido cuántico (A1_PA).
 $CVaR_{0,05} = 0,02768$, $HHI = 0,2102$, N efectivo = 4,76. Error de aditividad: $1,25 \times 10^{-16}$ (precisión de máquina).

A.3. Configuración completa de parámetros

La Tabla A.3 complementa los hiperparámetros presentados en el Capítulo 3 con la configuración completa de los 46 experimentos, agrupada por etapa del *pipeline*: selección del universo (QUBO), circuito de estimación de amplitud (IQAE) y optimización de pesos (Adam). Los parámetros no listados quedan en su valor por defecto, según los paquetes referenciados (`qiskit-aer` 0.17.2, `scipy` 1.16.3).

Datos	
Universo	374 S&P 500 (2020–2025)
Ventana de entrenamiento	2020-01-03 a 2023-12-29 ($T = 1,005$ días) ^b
Ventana OOS	2024-01-02 a 2025-12-30 ($T_{\text{OOS}} = 501$ días)
Rendimientos	$r_{i,t} = \ln(P_{i,t}/P_{i,t-1})$
Confianza CVaR	$\alpha = 0,05$
Semilla aleatoria	42
Selección de universo QUBO (Etapas 1)	
Solver	Annealing simulado (defecto); D-Wave híbrido BQM opcional
Tamaño objetivo	$N_{\text{sel}} \in \{10, 30, 100\}$
Pesos del objetivo	$\lambda_p = 1,0, \lambda_\pi = 0,5, \lambda_b = 0,3$
Penalización de cardinalidad	$P_{\text{card}} = 500$ ($N_{\text{sel}} = 10$); 100,000 ($N_{\text{sel}} = 100$)
Penalización de comunidad	$P_{\text{comm}} = 100$ ($N_{\text{sel}} = 10$), tope $M_c = 3$ por comunidad
	$P_{\text{comm}} = 10,000$ ($N_{\text{sel}} = 100$), tope $M_c = 15$ por comunidad
Límite de tiempo del solver	10 s (solo backend híbrido)
Circuito IQAE (Etapas 2)	
Qubits (circuito real, simulado)	$n = 4$ de distribución + 1 ancilla, 5 físicos (fijo para todo N)
Niveles de pérdida	$2^4 = 16$ bins discretizados
Precisión	$\varepsilon \in \{0,001, 0,002, 0,005, 0,010\}$
Disparos / circuito	4,096 (simulador); hardware 1,024–8,192 por trabajo (Sección 4.4.3; 1,024 solo en el estudio de disparos del bloque D de la oleada 3)
Profundidad (sonda de hardware, 2 qubits)	$D_A + 6m$ con $D_A = 4$ (CZ nativa de IQM; ver Capítulo 4, Sección 4.4.1)
Profundidad (circuito real, <i>offline</i>)	1,476 CZ 782 en $m = 0$; 5,395 en $m = 1$ (nunca ejecutado en hardware)
Prob. de fallo	$\delta = 0,05$
Optimizador Adam (Etapas 3)	
Optimizador	Adam [27], en espacio de parámetros softmax
Parametrización	Softmax, $\tau = 1,0$ (use_softmax=True, todas las ejecuciones)
Tasa de aprendizaje	$\eta = 0,05$
Decaimiento de la tasa	$\times 0,995$ por iteración
β_1, β_2	0,9, 0,999
ε_A	10^{-8}
Cota superior de peso	$u_i = 0,3$ para todo i (recorte-y-renormalización)
Convergencia (gradiente)	$\ \nabla_\theta \text{CVaR}\ _2 < \delta_c = 10^{-5}$
Convergencia (CVaR rel.)	cambio relativo máx. $< 10^{-3}$ en ventana de 5 iteraciones
Paciencia (parada temprana)	$P = 20$ iteraciones sin mejora
Iteraciones máx.	50 (familias con $N_{\text{sel}} = 10$, y también S1_PC_100); 80 (S1, $N_{\text{sel}} = 30$); 100 (S1_PA_100, S1_PB_100) — S1_PB_100 (100/100) y S1_PC_100 (50/50) agotan su tope sin converger
Simulación de ruido (serie E)	
Error de un qubit	$p_{1q} = p_{2q}/10$
Error de dos qubits	$p_{2q} \in \{10^{-3}, 5 \cdot 10^{-3}, 10^{-2}\}$
Modelo de ruido	Depolarizante (Qiskit Aer)
Backtesting bootstrap	
Remuestreos	$B = 5,000^a$
Tamaño de bloque	21 días de negociación
Método	Bloque circular [28]

^aNingún artefacto persistido registra el número de remuestreos de la ejecución publicada (el valor por defecto del script es 10,000; 5,000 coincide con su ejemplo documentado de modo rápido); véase la nota de reproducibilidad del Capítulo 5. ^bA4_PA (fuera de esta batería, Sección 6.3 del Capítulo 6) usa una ventana más larga, $T = 1,508$ días.

Tabla A.3: Configuración completa de parámetros a lo largo de los 46 experimentos.

A.4. Identificadores de trabajos en hardware

Los identificadores de trabajo (UUID) de hardware IQM devueltos por Amazon Braket para las Wave 1/2/3 y por IQM Resonance para la comparación de ZNE sustentan, uno a uno, los resultados de la Sección 4.4. Como muestra representativa: 019d8400-4e8b-74a2-aa8c-d2294ea545e5 (Wave 1, Emerald, $m = 0$), 019d840d-12c6-7740-b803-eacdb98a898e (Wave 2, Emerald, $m = 1$), 019d8411-3723-7080-84a0-455e670792f4 (Wave 3, Emerald, $m = 4$), 019d8411-d65b-7fc0-8a72-bcfa85c9b242 (Wave 3 bloque G, Garnet, $m = 0$), 019d8417-fc21-7c12-95cd-e24e69da79ba (post-Wave 3, Garnet, $m = 0$) y, para la comparativa de ZNE en Emerald a $m = 1$, 019e7d35-8363-70b1-a1f5-c9462f88d38a (bruto), 019e7d35-d9d1-7270-9516-a8aa2a2a0e03 (*folding* global, $\lambda = 3$) y 019e7d36-4f3f-7a23-b41e-1e400e940061 (FIIM, $\lambda = 3$). La lista completa por bloque está en `results/hardware_validation/` del repositorio.



Universidad Autónoma
de Madrid