

Fisher-information-inspired collective feedback in Kuramoto–XY Hamiltonians: exact small-system structure and a hardware falsification

Miroslav Šotek^{1,2,*}

¹*ANULUM, Marbach SG, Switzerland*

²*ORCID: 0009-0009-3560-0851*

(Dated: May 25, 2026)

(Nodo DOI: [10.5281/zenodo.20382125](https://doi.org/10.5281/zenodo.20382125))

Abstract

We introduce and characterise a collective Fisher-information-inspired magnetisation-feedback term for Kuramoto–XY quantum simulations. The model augments the heterogeneous XY Hamiltonian with a diagonal term proportional to minus λ times total magnetisation squared, normalised by the qubit count. The term is diagonal in the computational basis and therefore creates explicit magnetisation-sector energy shifts without breaking the magnetisation conservation of the ideal XY model. Artefact-generated exact diagonalisation on $n = 4, 6, 8$ qubits shows that increasing λ lowers the ground energy, widens the spectrum, and increases the spectral gap in the tested grid. For example, the exact gap changes from 1.1317 to 12.6060 at $n = 4$, from 0.4468 to 13.7801 at $n = 6$, and from 0.2827 to 13.7173 at $n = 8$ when λ is increased from 0 to 4. The full-spectrum adjacent-gap mean shifts downward at $n = 6, 8$ for larger λ , and the $n = 4$ mean low-energy bipartition entropy decreases from about 0.676 bits at $\lambda \leq 1$ to about 0.490 bits at $\lambda = 8$. A small $n = 4$ VQE benchmark further shows that a K_{ij} -informed ansatz gives lower median relative energy error than generic `TwoLocal` and `EfficientSU2` baselines for $\lambda \in \{0, 1, 4\}$. Finally, we executed the corresponding $n = 4$ IBM Heron r2 hardware tests on `ibm_kingston`. A single-sample pilot was followed by a repeated randomized run with 166 circuits and 339968 shots. The repeated run is a negative validation result: at $\lambda = 4$ relative to $\lambda = 0$, magnetisation leakage increased by 0.0896 on average across matched state/depth comparisons (Fisher combined $p = 7.49 \times 10^{-55}$), while exact-state retention decreased by 0.0797 on average. A full 16-state readout-confusion-matrix inversion on the same repeated run preserves this negative conclusion: magnetisation leakage still increases by 0.0917 ($p = 8.14 \times 10^{-55}$) and state retention still decreases by 0.0815 ($p = 1.19 \times 10^{-48}$). We then executed a dedicated replication/ZNE lane on `ibm_marrakesh` and `ibm_fez`. After full readout mitigation and linear ZNE, the mean $\lambda = 4$ minus $\lambda = 0$ magnetisation-leakage deltas remained positive in the $n = 4$ replication: +0.0599 on Marrakesh and +0.1937 on Fez. Larger $n = 6, 8$ Fez/Marrakesh replication/ZNE lanes were then executed. They preserve the falsification boundary rather than promoting FIM protection: all four larger lanes show positive mean raw scale-1 magnetisation-leakage deltas, with linear-ZNE deltas of 0.4021, 0.1480, 0.2541, and 0.4203 for Fez $n = 6$, Fez $n = 8$, Marrakesh $n = 6$, and Marrakesh $n = 8$, respectively. Thus the tested digital Trotter implementation falsifies the simple hypothesis that this FIM term acts as a hardware coherence-protection mechanism. The paper is deliberately bounded: the artefacts support exact small-system Hamiltonian findings and a

backend/circuit-specific negative hardware result on IBM Heron devices. They do not demonstrate hardware coherence improvement, quantum advantage, strict many-body localisation, or universal protection.

I. INTRODUCTION

Kuramoto synchronisation, XY spin Hamiltonians, and NISQ quantum simulation provide a compact setting for studying how oscillator-network structure appears in quantum dynamics [1–3]. In the standard mapping, heterogeneous oscillator frequencies and couplings are encoded as local Z fields and $XX+YY$ exchange terms. This construction preserves total magnetisation because the XY interaction swaps excitations without creating or destroying them.

This paper continues the SCPN Kuramoto–XY hardware programme, which first established a bounded parity-sector and excitation-number correlated leakage observation on IBM Heron hardware [4] and then used reduced-Pauli checks to localise correlator-level structure in the same families [5]. Here we test a distinct Hamiltonian modification: a Fisher-information-inspired magnetisation-feedback term. The contribution is a negative design rule: the direct digital Trotter implementation of this FIM term is not a coherence-protection mechanism on the tested superconducting hardware.

This manuscript studies a minimal extension of that model: a collective magnetisation-feedback term inspired by Fisher-information and self-consistency ideas. The term is

$$H_{\text{FIM}}(\lambda) = -\frac{\lambda}{n}M^2, \quad M = \sum_{i=1}^n Z_i. \quad (1)$$

It is not introduced here as a validated error-correction mechanism. Rather, the question is narrower and falsifiable: what does this term do to the exact small-system energy landscape, sector structure, spectral diagnostics, and low-energy entanglement, and does the most direct digital hardware implementation show any corresponding leakage suppression?

The contribution is threefold. First, we define the FIM-augmented Kuramoto–XY Hamiltonian and record the exact magnetisation-sector consequences of Eq. (1). Second, we generate committed JSON/CSV artefacts for spectra, level-spacing ratios, bipartition entropy,

* protoscience@anulum.li

ideal sector-conservation checks, and VQE ground-state scoring. Third, we report IBM Heron r2 raw-count tests that falsify the simple claim that the tested FIM Trotter construction improves hardware coherence on `ibm_kingston`. This negative result is part of the contribution: it separates a valid Hamiltonian-structure effect from an unsupported NISQ robustness claim.

II. MODEL

The base Hamiltonian is the heterogeneous Kuramoto-XY model

$$H_{XY} = - \sum_i \omega_i Z_i - \sum_{i < j} K_{ij} (X_i X_j + Y_i Y_j), \quad (2)$$

where $K_{ij} = K_{ji}$ is the oscillator coupling matrix and ω_i are natural frequencies. The sign convention follows the implementation in `scpn-quantum-control`. The augmented Hamiltonian is

$$H(\lambda) = H_{XY} + H_{\text{FIM}}(\lambda). \quad (3)$$

Since M is diagonal in the computational basis, M^2 assigns the same energy shift to every state in a fixed magnetisation sector. For a basis state with k spin-down qubits, the magnetisation is $M = n - 2k$, and the FIM shift is

$$\Delta E_M(\lambda) = -\lambda M^2/n. \quad (4)$$

Thus Eq. (1) separates sectors by a known collective diagonal contribution.

This exact-sector effect is the reason the hardware test is scientifically sharp. If the widened spectral gaps and magnetisation-sector shifts translated directly into a NISQ protection mechanism, the $\lambda > 0$ circuits would be expected to suppress leakage or improve state retention relative to the $\lambda = 0$ references. The experiments below test precisely that expected hardware consequence and falsify it for the direct Trotter implementation.

The same identity also explains the digital hardware burden. Because

$$M^2 = nI + 2 \sum_{i < j} Z_i Z_j, \quad (5)$$

the $\lambda > 0$ term is not a single local correction after compilation: it adds an all-to-all pairwise $Z_i Z_j$ layer, up to a global energy shift. On a sparse heavy-hex device this all-pair interaction

pattern must be routed and Trotterised, so any ideal spectral gap effect competes directly with extra two-qubit-gate and depth overhead.

The key scientific boundary follows immediately. Both H_{XY} and H_{FIM} commute with total magnetisation, so the ideal Hamiltonian does not create unitary leakage between magnetisation sectors:

$$[H(\lambda), M] = 0. \quad (6)$$

Therefore any IBM measurement of magnetisation leakage or survival asymmetry is not a direct ideal-Hamiltonian leakage effect. It must be interpreted as hardware, noise, state-preparation, transpilation, layout, or readout behaviour.

III. ARTEFACT-FIRST PROTOCOL

All numerical statements in this manuscript are generated from committed artefacts under `data/scpn_fim_hamiltonian/`. The scripts are:

- `scripts/analyse_fim_spectrum.py`,
- `scripts/analyse_fim_level_spacing.py`,
- `scripts/analyse_fim_entanglement.py`,
- `scripts/analyse_fim_sector_survival.py`,
- `scripts/benchmark_fim_vqe_ground_state.py`, and
- `scripts/analyse_fim_ibm_repeated_followup.py`,
- `scripts/analyse_fim_readout_matrix_mitigation.py`.

The claim boundary is maintained separately in `docs/campaigns/scpn_fim_claim_boundary_2026-05-05.md`. This separation is intentional: paper text is allowed to use only claims that remain valid after checking the generated artefacts.

n	λ	E_0	width	gap
4	0	-6.3030	12.6060	1.1317
4	4	-22.3030	24.6021	12.6060
6	0	-10.7930	21.5860	0.4468
6	4	-34.7930	40.2628	13.7801
8	0	-12.7557	25.2287	0.2827
8	4	-44.4730	51.4015	13.7173

TABLE I. Exact dense-spectrum summaries generated from `fim_spectrum_summary_2026-05-05.json`. The data support a small-system energy-landscape claim, not a hardware robustness claim.

IV. EXACT SPECTRAL EFFECTS

Table I summarises the main exact-spectrum trend. Increasing λ lowers the ground energy and increases the gap and spectral width for the tested $n = 4, 6, 8$ grid. This is expected from Eq. (4): sectors with large $|M|$ are shifted downward.

V. LEVEL-SPACING AND ENTANGLEMENT DIAGNOSTICS

The adjacent-gap ratio is used only as a small-system spectral diagnostic. It is not by itself a proof of strict many-body localisation. For $n = 6$, the full-spectrum mean adjacent-gap ratio changes from 0.4334 at $\lambda = 0$ to 0.4071 at $\lambda = 4$. For $n = 8$, it changes from 0.4368 at $\lambda = 0$ to 0.4048 at $\lambda = 4$ and 0.3852 at $\lambda = 8$. This supports cautious language such as “localisation-like spectral shift” rather than a definitive MBL claim.

The low-energy bipartition-entropy artefact shows a related but also bounded trend. For $n = 4$, the mean entropy over the lowest generated eigenstates is approximately 0.676 bits for $\lambda \leq 1$, then decreases to 0.584 bits at $\lambda = 2$, 0.502 bits at $\lambda = 4$, and 0.490 bits at $\lambda = 8$. The ground state entropy is zero in this grid, consistent with the large- $|M|$ sector structure selected by the FIM shift.

Ansatz	$\lambda = 0$	$\lambda = 1$	$\lambda = 4$
<i>K_{ij}</i> -informed	0.372%	0.413%	0.848%
TwoLocal	7.164%	11.236%	9.155%
EfficientSU2	7.708%	7.228%	9.050%

TABLE II. Median relative ground-state energy error from `fim_vqe_ground_state_summary_2026-05-05.json`. This is an $n = 4$ ansatz-scoring result, not a general optimisation theorem.

VI. IDEAL SECTOR CONSERVATION

The sector-survival artefact checks the Frobenius norm of the commutator with M and the off-sector Hamiltonian blocks. For all generated $n = 4, 6, 8$ cases and all tested λ values, the commutator norm with M and the maximum off-sector block norm are zero in the dense exact representation. The ideal unitary sector leakage column is therefore exactly zero by construction.

This result is scientifically important because it blocks an incorrect interpretation. The FIM term can create sector energy barriers and change which sectors dominate low-energy eigenstates, but it does not create or suppress ideal magnetisation leakage. A hardware experiment can still be meaningful, but only as a test of whether the modified circuit family exhibits λ -correlated survival or leakage under real noise and controls.

VII. VQE GROUND-STATE SCORING

The $n = 4$ VQE artefact compares a *K_{ij}*-informed ansatz against `TwoLocal` and `EfficientSU2` at equal repetition count, three random seeds, and fixed COBYLA iteration budget. The topology-informed ansatz has lower median relative error at all tested FIM strengths:

This result is useful for protocol design because it suggests that the same topology-informed construction used in the methods paper remains appropriate for the FIM-augmented Hamiltonian. It does not imply that the ansatz is optimal at larger n or under hardware noise.

VIII. IBM HARDWARE VALIDATION

The IBM hardware component was designed as a falsification test, not as a demonstration of advantage or error mitigation. Because $[H(\lambda), M] = 0$ in the ideal model, magnetisation leakage on hardware is not an ideal FIM transition. It is a real-device observable that combines circuit depth, transpilation, readout, state preparation, relaxation, and coherent gate errors. The question was therefore narrow: does increasing λ produce a reproducible reduction of measured leakage or increase of state retention under the tested digital Trotter construction?

The first pilot executed 61 circuits and 249856 shots on `ibm_kingston` as job `ibm-run-4c0bd60c3fc2c532`. It produced complete raw counts, but it had one sample per λ /depth/state condition and therefore supported only descriptive conclusions. Its direction was already unfavourable to a simple protection claim: at $\lambda = 4$ relative to $\lambda = 0$, mean magnetisation leakage increased by 0.1156 across 15 matched state/depth comparisons.

We then executed a repeated randomized follow-up on the same backend as job `ibm-run-cf4835290f607387`. The repeated design used:

- $n = 4$ qubits,
- $\lambda \in \{0, 4\}$,
- depths $\{2, 4, 6\}$,
- five representative magnetisation-sector initial states,
- five randomized replicates per matched condition,
- a 16-basis-state readout baseline,
- 2048 shots per circuit,
- 166 total circuits, and
- 339968 total shots.

Table III summarises the repeated hardware result. The sign is clear and reproducible in the opposite direction from the protection hypothesis: $\lambda = 4$ increased leakage and reduced

Observable	mean Δ_{4-0}	Fisher p	positive	negative
State retention	-0.0797	1.11×10^{-48}	0/15	15/15
Magnetisation leakage	$+0.0896$	7.49×10^{-55}	14/15	1/15
Parity leakage	$+0.0615$	6.27×10^{-48}	13/15	2/15
Readout-corrected parity leakage	$+0.0644$	6.27×10^{-48}	13/15	2/15

TABLE III. Repeated IBM Heron r2 follow-up on `ibm_kingston`. Positive leakage deltas mean $\lambda = 4$ leaked more than $\lambda = 0$ within matched initial state and depth. The readout correction is a state-specific parity-flip correction from the available basis-state readout baselines.

exact-state retention relative to $\lambda = 0$ for the tested implementation. The result is therefore a hardware falsification of this specific implementation, not a positive coherence result.

Because the repeated design included all 16 computational-basis readout baseline circuits, we also applied a literal $2^n \times 2^n$ readout confusion-matrix inversion offline. The calibration matrix is well-conditioned for this dataset (condition number 1.049). The mitigated results preserve the same sign and scale: state retention changes by -0.0815 on average ($p = 1.19 \times 10^{-48}$), magnetisation leakage changes by $+0.0917$ ($p = 8.14 \times 10^{-55}$), and parity leakage changes by $+0.0642$ ($p = 6.19 \times 10^{-48}$). This removes the scalable-readout-mitigation objection for the repeated FIM follow-up while leaving the negative hardware interpretation unchanged.

IX. DISCUSSION

The generated artefacts support a coherent but limited story. The Fisher-information-inspired feedback term is not a mysterious source of ideal sector leakage. It is a collective diagonal term that reshapes the magnetisation-sector energy landscape. In exact small systems, this reshaping widens gaps, shifts adjacent-gap diagnostics, and changes low-energy bipartition entropy. Those are scientifically valid offline findings.

The IBM result answers the tested hardware question negatively. The FIM term does reshape the exact Hamiltonian spectrum, but in the implemented digital Trotter circuits the $\lambda = 4$ instances carry enough additional hardware burden that measured leakage increases rather than decreases. This separates mathematical sector engineering from actual NISQ

robustness and prevents an incorrect coherence-protection claim.

The live-transpilation metadata makes this burden explicit. In the repeated follow-up artefact, the $\lambda = 0$ main circuits averaged transpiled depth 92.13 and 24.0 two-qubit gates, with maxima 134 and 36. The $\lambda = 4$ main circuits averaged transpiled depth 354.65 and 103.05 two-qubit gates, with maxima 540 and 158. Thus the negative hardware result is not mysterious: the tested digital implementation asks the processor to realise an all-to-all $Z_i Z_j$ feedback graph, and the added noise burden overwhelms the ideal small-system gap widening.

This distinction is the main scientific value of the hardware section. It shows that a diagonal collective term can be mathematically meaningful while still being operationally counterproductive after compilation to a noisy superconducting device. Reporting the negative result is therefore not a failed experiment; it is the evidence that fixes the claim boundary.

The negative result should not be overgeneralised. It does not prove that all FIM-like Hamiltonian designs are unhelpful, nor does it exclude other encodings, native analog implementations, different layouts, lower-depth constructions, adaptive feedback, or future devices. It does show that the direct all-pair digital implementation tested here is not a hardware-protection mechanism on this backend and calibration window.

The next hardware discriminator was a replication-and-stress protocol on two additional IBM Heron backends. The new lane used the same $n = 4$ FIM/baseline observable family with full 16-state readout calibration and local-folding ZNE scales $\{1, 3, 5\}$ for the dominant leakage and retention observables. The result strengthens the falsification rather than reversing it: on both `ibm_marrakesh` and `ibm_fez`, $\lambda = 4$ continues to increase magnetisation leakage and reduce state retention relative to $\lambda = 0$ after readout mitigation and linear ZNE.

X. RELATION TO OTHER PROTECTION STRATEGIES

It is useful to distinguish the FIM term from established protection strategies. Dynamical decoupling applies control pulses designed to average specific noise components away while preserving the intended logical evolution. Quantum error correction encodes logical information redundantly and uses syndrome information to detect or correct errors. The FIM term in Eq. (1) is neither of these. It is a Hamiltonian design term that reshapes

Backend	raw scale-1	raw linear ZNE	mitigated linear ZNE
<code>ibm_marrakesh</code>	+0.0527	+0.0585	+0.0599
<code>ibm_fez</code>	+0.1897	+0.1862	+0.1937

TABLE IV. Mean $\lambda = 4$ minus $\lambda = 0$ magnetisation-leakage deltas from the dedicated replication/ZNE lane. Positive values mean the FIM circuits leaked more than the matched baseline circuits. The mitigated column applies full 16-state readout-matrix inversion before the linear ZNE fit.

magnetisation-sector energies, while leaving the ideal magnetisation-sector decomposition intact.

This distinction explains why the negative IBM result is not surprising in hindsight. The direct digital implementation adds all-to-all ZZ structure without adding syndrome extraction, active correction, or a decoupling sequence targeted to the backend noise spectrum. The correct comparator is therefore not “FIM versus no protection” in the abstract, but “this compiled FIM circuit family versus other concrete hardware-control mechanisms at the same observable and layout boundary.” Future experiments should compare the same leakage and retention observables against backend-native dynamical-decoupling schedules, lower-depth native interaction decompositions, and explicit QEC or repetition-code scaffolds only after each comparator has its own preregistered circuit family and resource budget.

To make that comparison falsifiable, the replication/ZNE lane used $\lambda \in \{0, 4\}$, depths $\{2, 4\}$, four magnetisation-sector states, three replicates, local-folding noise scales $\{1, 3, 5\}$, and full 16-state readout calibration. The `ibm_marrakesh` job `d86ubgdg7okc73e110qg` and the `ibm_fez` job `d86ubg9789is73906ofg` each executed 160 circuits at 1024 shots. The readout calibration matrices were well-conditioned, with condition numbers 1.042 and 1.077, respectively. Table IV summarises the main result.

The same lane shows retention moving in the opposite direction. After readout mitigation and linear ZNE, the mean state-retention delta is -0.0340 on `ibm_marrakesh` and -0.1690 on `ibm_fez`. Thus the replication/ZNE extension does not rescue the direct digital FIM construction. It makes the claim boundary stronger: the observed failure is no longer a single-backend Kingston-only artefact, although it remains IBM-Heron and circuit-family specific.

Backend	n	Raw scale-1	Raw linear ZNE	Mitigated linear ZNE
ibm_fez	6	0.4091	0.4021	0.4257
ibm_fez	8	0.1479	0.1480	0.1539
ibm_marrakesh	6	0.2474	0.2541	0.2672
ibm_marrakesh	8	0.4231	0.4203	0.4483

TABLE V. Larger-width FIM replication/ZNE lanes submitted on 2026-05-20. Values are mean absolute $\lambda = 4$ minus $\lambda = 0$ magnetisation-leakage deltas from the retrieved raw-count artefacts. The table addresses the small-system limitation, but it remains an IBM-Heron circuit-family result rather than evidence for backend-general protection.

Finally, we executed larger-width replication/ZNE lanes at $n = 6$ and $n = 8$ on the same two Heron backends. These lanes use the same $\lambda = 4$ versus $\lambda = 0$ leakage discriminator, local-folding noise scales $\{1, 3, 5\}$, and full 2^n computational-basis readout calibration. Table V summarises the mean magnetisation-leakage deltas. The larger-width evidence does not reverse the negative hardware conclusion: every scale-1 mean delta is positive, and linear ZNE leaves the deltas positive at both widths and both backends. The values are backend- and width-dependent, so this is not a universal leakage law; it is a stronger falsification of the narrow claim that the tested FIM digital circuit family protects the measured magnetisation sector on IBM Heron hardware.

XI. LIMITATIONS

The study is limited to exact dense diagonalisation at $n = 4, 6, 8$, an $n = 4$ VQE scoring grid, and IBM Heron hardware lanes at $n = 4, 6, 8$. The level-spacing and entropy results are localisation-like diagnostics only. They do not prove strict many-body localisation. The VQE result uses a small optimiser budget and should be treated as an ansatz-design signal, not a theorem. The IBM hardware result is multi-backend within IBM Heron devices but still circuit-family specific and IBM-only. Full readout-matrix mitigation is available for the repeated and replication FIM follow-ups because they include complete computational-basis readout baseline circuits, but this does not correct gate errors, Trotter overhead, layout routing, or calibration drift. ZNE is also not a universal correction: it is a stress test

for monotone noise-scaling behaviour and can fail when folded circuits change coherent-error structure, calibration sensitivity, or dominant leakage channels. Datasets without complete basis calibration remain limited to exact-state or parity-readout corrections. The Fisher combined p-values in the hardware table combine matched-condition Welch tests and should be read as compact evidence of sign stability across the repeated design, not as proof of backend-general behaviour. No hardware coherence improvement, quantum advantage, or universal protection claim is made.

At the time of execution, authenticated live QPU access for the hardware component of this paper was limited to IBM Quantum hardware. The software stack has provider-neutral execution and artefact-packaging routes for other gate-model, trapped-ion, neutral-atom, photonic, annealing, analogue, and simulator targets, but those routes are not evidence of live non-IBM execution. We plan to rerun the same FIM and baseline payloads on additional hardware families when compute allocations and provider access become available. Until those reruns exist, the hardware conclusion is restricted to the IBM Heron backends, layouts, circuit family, and calibration windows reported here.

XII. DATA AND CODE AVAILABILITY

Code, scripts, and artefacts are maintained in the [project GitHub repository](#). The SCPN/FIM artefacts are under `data/scpn_fim_hamiltonian/`. The manuscript claim boundary is recorded in `docs/campaigns/scpn_fim_claim_boundary_2026-05-05.md`. The IBM raw-count and analysis artefacts include:

- `fim_ibm_pilot_raw_counts_2026-05-05_ibm-run-4c0bd60c3fc2c532.json`,
- `fim_ibm_pilot_analysis_2026-05-05_ibm-run-4c0bd60c3fc2c532.json`,
- `fim_ibm_repeated_followup_raw_counts_2026-05-05_ibm-run-cf4835290f607387.json`, and
- `fim_ibm_repeated_followup_analysis_2026-05-05_ibm-run-cf4835290f607387.json`,
- `fim_readout_matrix_mitigation_summary_2026-05-05_ibm-run-cf4835290f607387.json`,

- `fim_readout_matrix_mitigation_rows_2026-05-05_ibm-run-cf4835290f607387.csv`, and
- `fim_readout_matrix_mitigation_comparisons_2026-05-05_ibm-run-cf4835290f607387.csv`,
- `fim_replication_zne_submission_ibm_marrakesh_20260520T164759Z.json`, and
- `fim_replication_zne_submission_ibm_fez_20260520T164759Z.json`,
- `fim_replication_zne_raw_counts_ibm_marrakesh_20260520T165004Z.json`,
- `fim_replication_zne_analysis_ibm_marrakesh_20260520T165004Z.json`,
- `fim_replication_zne_raw_counts_ibm_fez_20260520T165004Z.json`, and
- `fim_replication_zne_analysis_ibm_fez_20260520T165004Z.json`.

AI USAGE DISCLOSURE

AI-assisted tools were used for drafting and editing support. Numerical claims, artefact provenance, scientific framing, authorship, and final responsibility were verified and retained by the author.

-
- [1] Y. Kuramoto, in *International Symposium on Mathematical Problems in Theoretical Physics* (Springer, 1975) pp. 420–422.
 - [2] J. A. Acebrón, L. L. Bonilla, C. J. P. Vicente, F. Ritort, and R. Spigler, [Reviews of Modern Physics](#) **77**, 137 (2005).
 - [3] F. Barahona, [Journal of Physics A: Mathematical and General](#) **15**, 3241 (1982).
 - [4] M. Sotek, Backend-sensitive hardware observation of parity-sector and excitation-number correlated decoherence asymmetry in the xy hamiltonian (2026), preprint manuscript and public repository artefacts.
 - [5] M. Sotek, Reduced-pauli entanglement checks for dla-sector leakage mechanisms on ibm heron hardware (2026), preprint manuscript and public repository artefacts.