

How gently can a protected quantum memory be read?

Rate limits on logical readout from the structure of the code

Dylan Henderson^{1,*}

¹*Cogitan, Coppell, Texas, USA*

(Dated: August 6, 2026)

Quantum error-correcting codes protect logical information by hiding it from the environment; yet the same information must remain accessible to the experimenter. We quantify this tension for weak continuous readout of bosonic codes, to leading order in the measurement rate and in the single-mode Markovian limit. For a logical observable read through a physical meter M at measurement strength γ , we separate the induced logical decoherence into a *conjugate* cost — dephasing of the partner logical operators, which is bounded below by the information-gain rate, saturating it for quantum-limited readouts — and a *self* cost: corruption of the very operator being read. We derive the *readability limit*, a lower bound on the self cost set by the code structure: $\Gamma_{\text{self}} \geq \gamma \max(2\kappa^2, \frac{1}{2}\varepsilon_{\text{KL}}^2) - O(\gamma^2/g)$, where κ is the within-codespace off-diagonal action of M and ε_{KL} is the logical-confusion (off-diagonal) part of the Knill–Laflamme residual of the measurement back-action $E = P_{\perp}MP$. The within-code term is an exact operator identity; the leak term is a recovery-independent converse from the code’s Knill–Laflamme data, established as a rate at $d = 2$ for direct-return stabilizers and supported numerically beyond them, saturating at the cells where the within-code term dominates and loose by up to two orders of magnitude where it does not. The which-path (diagonal) part of the residual, by contrast, feeds only the conjugate cost. A readout corrupts what it reads at least to the extent that it rotates the logical within the code space or its back-action falls outside the code’s correctable error set. We verify the bound cell-by-cell across cat and Gottesman–Kitaev–Preskill (GKP) codes and survey number-type codes, which *within the Hermitian meter menu considered here* are essentially unreadable-while-protected: their back-action is a Knill–Laflamme-uncorrectable error, so the confusion floor is maximal. The canonical phase measurement conventionally used to read number-phase codes is a POVM — the spectrum of the non-Hermitian Susskind–Glogower phase operator — and so lies outside the present framework; whether the bound extends to POVM meters is left open. Among the readable families, the $7.8\times$ readout advantage of cat over GKP qubits we find at matched photon number decomposes into two contributions of distinct origin: a meter-scale gain — the cat’s protected variable is an unbounded quadrature whose signal grows with photon number, against GKP’s bounded modular one — and a $\approx 1.5\times$ residual ratio of self-disturbance rates (at a single photon number) that a meter-normalized readability isolates. We map the resulting readability–protection trade-off across the GKP lattice family. The readability limit turns readout quality into a code-*design* objective and gives a recipe for self-reporting memories — engineer the logical readout so that its back-action is an error the code already corrects. Within the standard bosonic menu this is achieved exactly by symmetry (the cat parity readout is exactly quantum-non-demolition) and, for tuned modular meters, to the controlled leading order.

I. INTRODUCTION

A quantum error-correcting code protects logical information by ensuring that the environment, coupling through its physical noise channels, learns nothing about the encoded state [1–3]. The same concealment applies to the experimenter: whatever channel is used to *read* the logical information is, physically, another environment. The qualitative form of this tension is folklore — gaining information about encoded data disturbs it — but its *rate-level* form, for a specified meter acting on a specified code, has not been stated: given a code and a physical observable M monitored at strength γ , how fast must the logical operator being read degrade? This paper answers that question.

The tension between protecting a logical degree of freedom and coupling to it has recently been sharpened at the structural level by Chen *et al.* [4], who prove a protection–sensitivity trade-off $d \cdot F = O(n^2)$ for principal families of stabilizer probes (non-degenerate, quantum LDPC, and generalized Shor codes) under bounded-degree K -local signal Hamiltonians, and bypass it with “asymmetric” codes whose signal-aligned logical direction is deliberately left at constant effective weight. Their results — like the HNLS (“Hamiltonian not in Lindblad span”) condition of QEC-assisted metrology, introduced in [5] and applied in [6] — are statements about the quantum Fisher information of the encoded probe, and so presuppose an optimal measurement; indeed, Chen *et al.* explicitly leave the measurement stage — extracting the enhanced signal while preserving protection against complementary errors — as an open problem. Our theorem addresses precisely this stage: the operator-selective asymmetry we quantify is the dynamical, rate-level coun-

* dylan@cogitan.ai

terpart of their combinatorial weight asymmetry. Where they show a code must leave the signal direction unprotected for information to accumulate, we quantify the price of reading it out.

A closely related obstacle, in a different regime, was recently demonstrated by Halaski and Koch [7]: continuous parity readout in circuit QED corrupts logical information because the three-body meter coupling $Z_1 Z_2 \hat{M}$ that faithful parity transfer requires is replaced, in practice, by a sum of two-body couplings $(Z_1 + Z_2) \hat{M}$. Their setting is multi-qubit and the monitored operator is a stabilizer rather than a logical, so it lies outside the single-mode, logical-readout scope in which our theorem is proven (cf. the open directions of the Discussion); we therefore draw the connection as a qualitative analogy rather than a special case. In our language their substitution acts as a hardware-forced within-code rotation $\kappa \neq 0$: what they call the *information-induced* phase (their Eq. 11), which carries a Wiener increment and is therefore stochastic, is the direct analogue of the $2\gamma\kappa^2$ term of Theorem 1, and their *interaction-induced* phase (their Eq. 10) the deterministic face of the same rotation. Making this precise — reducing their result to ours — would require extending the theorem to the multimode, stabilizer-monitored setting, which we leave open. The same analogy indicates why their two remedies — native multi-body couplings ($\kappa = 0$) and erasure encodings (relaxed back-action correctness) — are natural escape routes: they are the two ways to send the within-code rotation and the leak-confusion terms separately to zero.

Contributions: (i) the perceive/disturb figure of merit with the self/conjugate split (Sec. II); (ii) the conjugate identity as anchor (Sec. III); (iii) the self-disturbance bound (Sec. IV, the main result); (iv) the blindness corollary, explicit in QEC metrology and restated here at rate level (Sec. V); (v) numerical verification and the readability atlas of bosonic codes (Sec. VI); (vi) design consequences (Sec. VII).

II. SETUP AND FIGURES OF MERIT

We consider a d -dimensional logical code with code-space projector P and codewords $\{|\mu\rangle\}_{\mu=0}^{d-1}$. The code is stabilized dissipatively: a Lindbladian $\mathcal{L}_{\text{stab}}$ has the code space $\text{ran } P$ as its dark space and confines the complementary subspace with a gap g . Throughout, $\mathcal{D}[L]\rho = L\rho L^\dagger - \frac{1}{2}\{L^\dagger L, \rho\}$ denotes the Lindblad dissipator, $P_\perp = \mathbb{1} - P$ the projector onto the complement of the code space, and $\bar{n}$ the mean photon number; logical operators are written in qubit notation, with Z_P and X_P the logical Pauli operators on the code (at $d = 2$). A logical observable is read by weakly and continuously monitoring a Hermitian meter M , modeled by the additional dissipator $\mathcal{L}_M = \gamma \mathcal{D}[M]$ at measurement strength γ .

Signal (the read-basis distinguishability, equal to the leading Fisher/SNR content of the homodyne record for

quadrature meters) for resolving the logical O eigenbasis:

$$\mathcal{S}(M, O) = \sum_{i < j} |\langle i|M|i\rangle - \langle j|M|j\rangle|^2, \quad (1)$$

over O -eigenstate codewords. Two limitations of $\mathcal{S}$ belong here rather than in a footnote, since the atlas compares codes through it. First, it retains only the diagonal, pointer-shift content of M ; the off-diagonal elements $\langle i|M|j\rangle$ broaden the pointer distributions and degrade discrimination, so $\mathcal{S}$ is least faithful precisely when κ is large — the self-damaging regime the theorem is about. Second, it carries no detection efficiency, added noise, or integration time. Those cancel in a ratio taken within one meter type but not across types: readability values are directly comparable between two quadrature readouts, and only indicatively so between a quadrature and a photon-number or parity meter, which require different detection chains altogether. Disturbance: the induced decay rates of logical coherences, extracted from the slow spectrum of $\mathcal{L}_{\text{stab}} + \gamma \mathcal{D}[M]$, resolved per logical operator: Γ_{self} (on the operator being read) and Γ_{conj} (on its conjugates). Readability: $\mathcal{P} = \gamma \mathcal{S} / \Gamma_{\text{self}}$.

Decomposition of the meter with respect to the code: $M = PMP + (P_\perp MP + \text{h.c.}) + P_\perp MP_\perp$, with $PMP = \bar{m}P + \frac{\Delta m}{2}Z_P + \kappa X_P$ (qubit notation; a phase choice on the read basis removes the Y_P component without loss of generality, so $\kappa = |\langle 0|PMP|1\rangle|$ throughout). We write $\vec{n} = (\kappa, 0, \Delta m/2)$ for the Bloch vector of the traceless part of PMP , and set $n_z \equiv \Delta m/2$ and $n^2 \equiv |\vec{n}|^2 = \kappa^2 + n_z^2$, so that $\kappa^2 = n^2 - n_z^2$; these are used throughout the appendices. The back-action error is $E = P_\perp MP$, with Knill–Laflamme data $\Lambda = PMP_\perp MP$. The residual of Λ splits into two parts with sharply different roles: its *off-diagonal* part in the read basis,

$$\varepsilon_{\text{KL}}^2 \equiv \frac{|\Lambda_{01}|^2}{\text{Tr } \Lambda}, \quad (2)$$

which measures the indistinguishability of the leaked states and feeds the *self* channel, and its diagonal traceless part (codeword-dependent leak rates), which is which-path information and feeds only the *conjugate* channel. This normalization renders $\varepsilon_{\text{KL}}^2$ homogeneous with κ^2 .

III. THE CONJUGATE IDENTITY

For a code-diagonal meter ($\kappa = 0$, $E = 0$, $PMP = \sum_\mu m_\mu |\mu\rangle\langle\mu|$), the dissipator $\gamma \mathcal{D}[M]$ acts on the logical coherence $|\mu\rangle\langle\nu|$ as pure dephasing at rate $(\gamma/2)|m_\mu - m_\nu|^2$. For a qubit this is precisely $(\gamma/2)\mathcal{S}(M, Z)$: the conjugate coherences decay at the rate at which the meter record distinguishes the Z eigenstates. This is the measurement-induced-dephasing/information-gain correspondence of quantum-limited continuous measurement [8–10], which we use as a validation anchor rather

than claim: in general, *per codeword pair*,

$$\Gamma_{\text{conj}}^{(\mu\nu)} \geq \frac{\gamma}{2} |m_\mu - m_\nu|^2, \quad (3)$$

with equality for quantum-limited readouts, and excess conjugate dephasing whenever the meter carries back-action without information. For $d = 2$ this reads $\Gamma_{\text{conj}} \geq (\gamma/2)\mathcal{S}(M, Z)$. For $d > 2$ only the pairwise statement holds, and the summed functional $\mathcal{S}$ is not a proxy for it: the ratio of $\mathcal{S}$ to any individual pair term is unbounded when that pair is nearly degenerate, so no fixed combinatorial factor relates the two. Numerically, at $d = 2$ we find saturation to 6–8% for the natural readouts (cat via quadrature, ratio 1.08; GKP via modular quadrature, 1.06), over-dephasing only where κ or the back-action is large, and no violations of the pairwise inequality (Table I; the $d > 2$ rows report the worst single-pair coherence rate). The per-pair $|m_\mu - m_\nu|^2$ are not tabulated, so the $d > 2$ pairwise check cannot be reproduced from the table alone; it is reported from the underlying run.

IV. THE SELF-DISTURBANCE BOUND

Theorem 1 (Readability limit). *Let Γ_{self} denote the damage rate of the logical operator being read: the initial slope $-\frac{d}{dt} \ln C_Z|_{t \rightarrow 0^+}$ of the logical Z autocorrelation under the adiabatically eliminated code-space generator (Appendix A) or, equivalently by Lemma 1, the initial decay rate of the Helstrom distinguishability of the read-basis codewords — a quantity involving no choice of observable. To leading order in γ/g and for $\kappa^2, \varepsilon_{\text{KL}}^2 \lesssim \|M\|^2$,*

$$\Gamma_{\text{self}}(M) \geq \gamma \max(2\kappa^2, c_2 \varepsilon_{\text{KL}}^2) - O(\gamma^2/g), \quad (4)$$

where the constant of the first term is exact, and $c_2 = 1/2$ follows from combining the per-leak-event Helstrom floor ($1/4$ at small overlap) with the leak rate (Appendix C): each leak event forces any recovery — at any speed — to distinguish the leaked states, whose overlap is precisely Λ_{01} , and the resulting confusion is a logical error on the operator being read; cf. the approximate-correction converse [11, 12]. The additive form $2\kappa^2 + c_2 \varepsilon_{\text{KL}}^2$ holds at leading order under the confinement hypotheses of Appendix D (dark space, gap g , weak measurement $\gamma\|M\|^2/g \leq 1/32$), there proven with explicit constants for direct-return stabilizers: the within-code and leak contributions occupy orthogonal operator blocks of the slope functional $-\text{Tr}[Z \mathcal{L}(\rho)]$, and the only cross term pairs the off-diagonal block with the code-perp coherences, which are $O(\gamma\|E\|\|M\|/g)$ by Lemma 2 and so fall into the $O(\gamma^2/g)$ remainder; the two terms therefore add at leading order, as confirmed numerically at every cell tested. The max form follows a fortiori. The bound is stated for a logical qubit ($d = 2$); for $d > 2$ it applies to each codeword pair, with κ and Λ_{01} the corresponding pairwise quantities.

Status of the two terms. We state plainly what is established and what is not, because the two terms of the

max have different standing and are proven under different idealizations.

The within-code term is an exact operator identity (Appendix B), nonperturbative in κ , and holds for the initial slope with no hypotheses beyond the decomposition of *PMP*. The leak term is a converse: the per-leak-event Helstrom floor is proven, but its promotion to a *rate* runs through the flux identity of Appendix C, which is established at $d = 2$ under the confinement hypotheses (G1)–(G3) of Appendix D, and only for direct-return stabilizers. Cascaded recovery basins — including the two-photon dissipation that stabilizes physical cat qubits — violate the confinement hypothesis as stated, for a reason made explicit in that appendix, and are supported numerically rather than derived; the cat case additionally has unbounded K , which makes the quantitative remainder vacuous without a truncation hypothesis. The g -uniformity of the per-event floor is argued physically (measurement jumps are incoherent point events, so damage accumulates linearly rather than being Zeno-suppressed) rather than bounded.

The two terms also privilege different times. At the defining $t \rightarrow 0^+$ the within-code constant is exact, but the leak contribution to the slope is *bare* leakage $\gamma\epsilon^2$ rather than the post-recovery damage the converse describes; since bare leakage exceeds post-recovery damage, the leak term remains a valid lower bound there, but a loose one. In the extraction window the ordering reverses: the leak term is the physically meaningful quantity and the within-code term is reduced by the factor of Eq. (A2). The max-form is therefore a valid bound under both conventions and tight under neither uniformly. The *additive* form is claimed only at leading order under (G1)–(G3), and is not needed for anything below; the max-form follows from either term a fortiori.

Proof sketch. (i) *Within-code term, exact.* Restricted to the code space, $\mathcal{D}[PMP] = \mathcal{D}[\vec{n} \cdot \vec{\sigma}]$ with $\vec{n}$ the Bloch vector of Sec. II, and the qubit identity $(\vec{n} \cdot \vec{\sigma}) Z (\vec{n} \cdot \vec{\sigma}) = 2n_z(\vec{n} \cdot \vec{\sigma}) - n_z^2 Z$ gives $d\langle Z \rangle/dt = -2\gamma\kappa^2\langle Z \rangle + (\text{transverse})$, using $n^2 - n_z^2 = \kappa^2$: the constant 2 is an operator identity, nonperturbative in κ . (ii) *Leak-return term, converse.* The back-action $E = P_\perp MP$ leaks codewords at rate $\gamma\Lambda_{\mu\mu}$, $\Lambda = PE^\dagger EP$; the stabilizer returns them through its recovery basin. Whatever recovery the stabilizer implements, its logical damage is bounded below by that of the *optimal* recovery for the error set $\{E\}$, which by the approximate-QEC converse is $\Theta(\varepsilon_{\text{KL}}^2)$. The diagonal (rate-asymmetry) part of the KL violation damages only the conjugate operators — it is a which-path measurement through the leak channel — which is why Γ_{self} is second order in the residual while Γ_{conj} is first order. (iii) The two contributions act through orthogonal operator sectors and add at leading order. $\square$

Numerically, the exact within-code constant is confirmed to three digits at the cat-parity cell (predicted $2\gamma\kappa^2 = 0.0999$, measured 0.100) and to 4% at GKP-mod- p (0.0709 vs 0.0737, remainder consistent with the leak term); the corrected bound is respected at every

(code, meter) cell with a nonzero bound — eleven of the twelve, the exception being the blind GKP-parity cell whose bound is exactly zero (Fig. 1c). Saturation, however, is confined to the cells where the within-code term dominates: cat-parity ($1.0\times$) and GKP-mod- p ($1.04\times$) sit essentially on the bound, while the leak-dominated and mixed cells run from $4\times$ to $316\times$ above it (GKP- q is the extreme, 3.1×10^{-3} measured against a bound of 9.8×10^{-6}). The bound is therefore predictive of the magnitude of Γ_{self} only in the κ -dominated regime; elsewhere it is a valid but weak floor, and what sets the measured rate at those cells is not established here. It is loosest at the unbounded-meter cells (photon number), which fall outside the weak-measurement regime $\gamma\|M\|^2/g \leq 1/32$ under which the leak term is proven, and where uncontrolled higher-order terms in γ/g inflate Γ_{self} above the leading bound (the full $d = 2$ grid is Table II; the constant $c_2 = 1/2$ — what the converse bounding chain yields, distinct from the operator-identity constant of the within-code term, and not claimed tight: two steps in the chain are lossy, the Helstrom relaxation $(1 - \sqrt{1-x})/2 \geq x/4$ except as $x \rightarrow 0$ and $\max(\epsilon_0^2, \epsilon_1^2) \leq \text{Tr}\Lambda$ by up to a factor of two — and its normalization convention are derived in Appendix C).

V. BLINDNESS COROLLARY

Corollary 1. *If the meter’s within-code action is trivial — the diagonal Knill–Laflamme condition $PMP \propto P$, which holds whenever the code corrects M as an error (the converse need not hold: a code can be blind to M without correcting it) — then $S(M, O) = 0$ for every logical O .*

Protection against a channel implies blindness to it. This is not new, and the prior statement is more explicit than an acknowledgement of “folklore” would suggest: it appears in QEC-assisted metrology as the pair of code-design conditions $PL_kP \propto P$ (Knill–Laflamme [2]) and $PHP \not\propto P$ (a signal that survives encoding) [5, 6], in the same notation used here, and Faist *et al.* [13] give the same algebra for transversally acting symmetry generators, summarising it as the statement that one cannot measure what one can correct. We restate it at the rate level for completeness — it is the boundary case of Theorem 1. It should not be conflated with the HNLS condition, which asks the different, existential question of whether *some* code corrects the whole Lindblad span while leaving the signal visible; failure of HNLS degrades the achievable scaling to the standard quantum limit rather than driving the signal to zero.

VI. THE READABILITY ATLAS OF BOSONIC CODES

The atlas surveys the principal single-mode bosonic code families — cat codes [14, 15], GKP codes [16],

TABLE I. Best logical- Z readout per code at matched $\bar{n} \approx 1.9$: signal $\mathcal{S}$, self- and conjugate-disturbance rates, and readability $\mathcal{P} = \gamma\mathcal{S}/\Gamma_{\text{self}}$. Every code is read through the canonical observable conjugate to its protected variable; none of these canonical readouts achieves strictly zero self-disturbance (tuned meters can; Sec. VII).

code	readout	$\mathcal{S}$	Γ_{self}	Γ_{conj}	$\mathcal{P}$
cat $d=2$	q	15.2	1.17×10^{-3}	0.41	650
cat $d=3$	p	17.1	7.90×10^{-3}	0.108	108
cat $d=4$	q	30.4	1.94×10^{-2}	0.127	78
GKP $d=2$	mod- q	2.84	1.7×10^{-3}	0.075	83
GKP $d=3$	mod- q	3.51	3.15×10^{-3}	0.047	56

and number- and rotation-symmetric codes [17, 18] — each read through the menu of physically available meters (quadrature, photon number, parity, and modular quadrature). All rates are extracted from the slow spectrum of the full Liouvillian $\mathcal{L}_{\text{stab}} + \gamma\mathcal{D}[M]$ (measurement strength $\gamma = 0.05$ in stabilizer units, photon number matched at $\bar{n} \approx 1.9$ across codes). For every $d = 2$ cell the rates were additionally re-extracted directly from the windowed autocorrelation of Appendix A on the full dense Liouvillian — no mode identification of any kind — agreeing with the spectral extraction at the 3–12% level; the two definitions coincide in the window, as the theory requires, and the quoted $d = 2$ self-rates are the autocorrelation values (e.g. GKP mod- q : 1.71×10^{-3} , stable to $< 3\%$ across the entire window). The $d > 2$ rows initially carried triple-run spectral values (spread $< 1\%$); completing their autocorrelation cross-check — which requires adjoint normalization, $C(0) = \text{Tr}[Z^\dagger Z P]/d$, because the qudit Weyl logical is non-Hermitian and $\text{Tr}[Z^2 P]$ vanishes identically — revealed that the spectral extraction *understates* the $d > 2$ self-rates by 20–50% (mode misidentification among the near-degenerate logical modes), so the quoted $d > 2$ self-rates are now the windowed autocorrelation values, consistent with the $d = 2$ convention. At $d > 2$ several logical modes contribute and the windowed slope drifts across the window by an amount that varies strongly by row — 23% for cat $d = 4$, 13% for cat $d = 3$, but only 5% for GKP $d = 3$, much the most stable of the three (values quoted at the lower edge); these rows carry atlas (not theorem-verification) weight.

Three structural facts emerge. *First*, readable-while-protected is universal, and for the canonical readouts never perfect: every family admits a near-QND readout through the observable conjugate to its protected variable — plain quadrature for cat codes, modular quadrature for GKP (whose logical is invisible to plain homodyne, $\mathcal{S}_q \sim 10^{-22}$) — with back-action landing dominantly on the conjugate, but Γ_{self} is nonzero for every *canonical* readout tested — blind meters such as GKP-parity correctly floor at numerical zero, exactly as the blindness corollary requires — while tuned me-

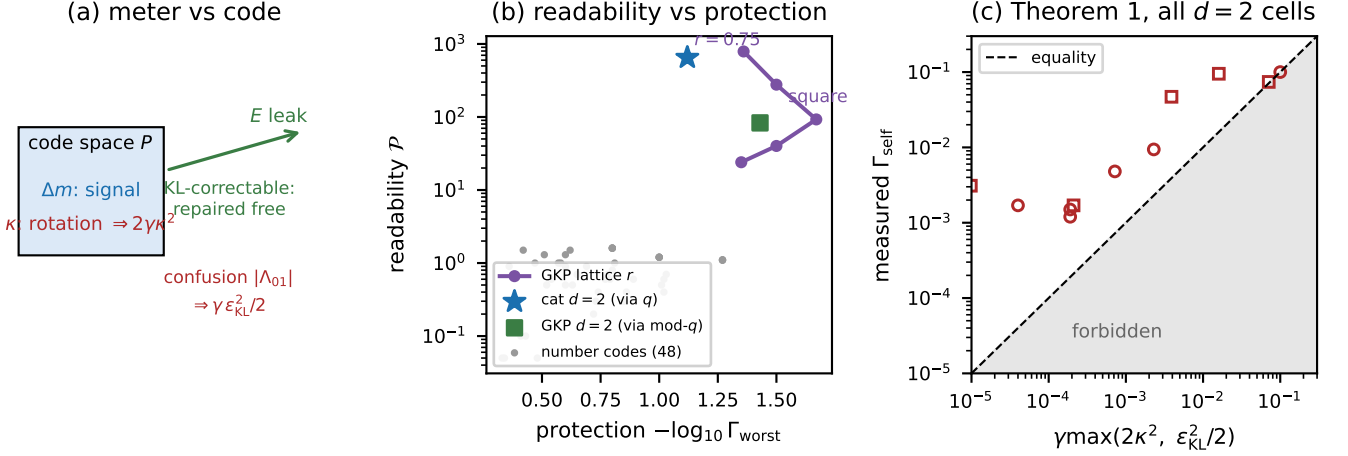


FIG. 1. (a) A meter M acting on a codespace decomposes into signal (Δm), within-code rotation (κ , self-damage with exact constant 2), and leakage E whose damage is free if Knill–Laflamme-correctable and otherwise floors at the per-event confusion rate. (b) The readability–protection plane at $\bar{n} \approx 1.9$: the GKP lattice family traces a pure trade-off through the Pareto-optimal square point; the cat qubit is the most readable undeformed code, number codes are confined to the lower left (protection scores conservative; see text). Number-code and lattice-curve readabilities use the spectral extraction. (c) Theorem 1 verified on the $d = 2$ (code, meter) cells with nonzero bound: every measured self-rate lies on or above the bound with derived constants; the cat–parity cell saturates, and the blind GKP–parity cell floors below the axis.

ters within the same physical menu can achieve exactly zero within-code rotation and zero logical confusion (Sec. VII). *Second*, the raw readability hierarchy at $d = 2$ — cat/quadrature ($\mathcal{P} \sim 650$) $\gg$ GKP/modular ($\mathcal{P} \sim 83$) $\gg$ number-type codes ($\mathcal{P} \sim 1$) — has two distinct origins. It is also, as a comparison, cross-modal: plain homodyne, modular-quadrature and number/parity readouts need different detection chains, and $\mathcal{S}$ carries no efficiency or added-noise budget to calibrate between them (Sec. II). The hierarchy should be read as a statement about idealized meters of each kind, not as a prediction of measured signal-to-noise across a common apparatus. The number-code floor is the bound at work: a number-code’s back-action E is relative dephasing of its codeword components — a KL-violating error set for loss-designed codes, so the $\varepsilon_{\text{KL}}^2$ term of the bound is maximal: the readout channel *is* the noise channel. In a coarse grid of 48 generalized number codes (three Fock-sector spacings $\times$ four term-counts $\times$ four amplitude envelopes), read through the number and parity channels, none escaped this floor, and the family is dominated by cat and GKP on both axes (the readability deficit, $\sim 10^2$ – 10^3 , is stabilizer-independent; the protection scores are conservative, as the search employed a biased recovery pump).

This carries a scope restriction, and it is the one place where the atlas judges a code family through a meter that family was not designed for. Number-phase codes [18] are conventionally read in the dual basis by *canonical phase measurement* — a POVM, equivalent to measuring the spectrum of the non-Hermitian Susskind–Glogower phase operator, and one whose modular phase uncertainty vanishes for ideal number-phase codes. No such meter appears in the menu above, and none can: Sec. II models

readout as $\gamma \mathcal{D}[M]$ for Hermitian M . The floor reported here is therefore a statement about what the Hermitian, weakly-monitored menu can do with these codes, *not* a claim that number-phase codes cannot be read while protected. Whether an analogue of Theorem 1 holds for POVM meters — and in particular whether the vanishing phase uncertainty of ideal number-phase codes evades the $\varepsilon_{\text{KL}}^2$ floor or merely relocates it — is the sharpest open question this atlas raises.

The cat-over-GKP gap, by contrast, decomposes into a meter-scale gain times a residual ratio of self-disturbance rates. Because $\mathcal{P} = \gamma \mathcal{S}/\Gamma_{\text{self}}$ identically, *any* ratio of readabilities splits this way; the physical content is that the two factors have distinct origins. The meter-scale factor is set by the code geometry: the cat’s protected variable is an unbounded quadrature, so its signal grows with photon number ($\mathcal{S} = 8\bar{n}$ at $d = 2$ for $q = (a + a^\dagger)/\sqrt{2}$, exact for coherent states and accurate to $O(e^{-4\alpha^2})$, some 2% at $\bar{n} \approx 1.9$, for the orthogonalized cat codewords), whereas GKP’s is modular and bounded, $\mathcal{S} = O(1)$. The residual factor is the part the bound constrains. Indeed, for any $d = 2$ meter $PMP = \bar{m}P + \frac{\Delta m}{2}Z_P + \kappa X_P$ the within-code meter variance is $V = (\Delta m/2)^2 + \kappa^2$, so the signal-to-variance ratio obeys the meter-independent identity $\mathcal{S}/V = 4/(1 + 4\kappa^2/\Delta m^2) \approx 4$: a scale-invariant readability $\gamma(\mathcal{S}/V)/\Gamma_{\text{self}} \approx 4\gamma/\Gamma_{\text{self}}$ isolates the residual factor, and the $7.8\times$ raw advantage we find resolves into a meter-scale gain times a $\approx 1.5\times$ ratio of self-disturbance rates (171 vs. 117). This residual ratio is a code-*and*-meter quantity read off at a single photon number ($\bar{n} \approx 1.9$); with both Γ_{self} near the extraction floor and carrying 3–12% uncertainty, the value 1.5 is

resolved only to $\approx \pm 0.2$, and we do not claim it as a structural constant. The quoted raw hierarchy refers to the $d = 2$ representatives of each family; at $d = 3$ the raw cat advantage narrows substantially (to a factor of order 2, from a single $d = 3$ comparison on atlas-weight rows; no $d=4$ GKP comparator was run), and the scale-invariant ranking is not settled at $d > 2$. *Third*, within the Gaussian (lattice) family the bound manifests as a readability-protection Pareto trade-off with no interior optimum (Fig. 1b): deforming the square GKP lattice by aspect ratio r buys readout at protection's expense ($r = 0.75$: $\mathcal{P} = 789$ at worst-case logical rate 4.3×10^{-2} ; $r = 1$: $\mathcal{P} = 93$ at 2.2×10^{-2} ; the lattice sweep uses the spectral extractor throughout, so its $r = 1$ value is the spectral 93 rather than the autocorrelation 83 quoted for the same cell in Table I), the square lattice is Pareto-optimal, and no Gaussian deformation dominates it — consistent with the fact that Gaussian relatives merely re-select which quadrature is protected versus read.

VII. DESIGN CONSEQUENCES

The readability limit (Theorem 1) converts readout quality from an afterthought into a design objective with an explicit recipe: a *self-reporting* memory is a code-meter pair with (i) *PMP* diagonal in the read basis ($\kappa = 0$) and (ii) back-action inside the correctable set ($\varepsilon_{\text{KL}} = 0$). Both conditions are properties the code designer controls — and, contrary to the impression given by the canonical cells of Table I, the leading-order bound can be driven to zero within the physically available menu (quadratures, photon number, parity, modular quadratures, and their few-tone combinations). For the symmetry route below this is exact ($E = 0$, true QND); for tuned meters it holds to the controlled leading order in γ/g , with residual higher-order self-disturbance not excluded. Symmetry provides one route: if the read-basis codewords lie in distinct eigenspaces of a unitary R with $[M, R] = 0$, then $\kappa = \Lambda_{01} = 0$ identically; the parity readout of a two-component cat in its parity ($\pm$) basis is of this type, with $E = 0$ — exactly QND. The physical primitive is the dispersive QND parity measurement demonstrated in [19] and used as the error syndrome in [20]; note that in those experiments parity carries *no* logical information by construction — both codewords of the cat code of [20] are even-parity states — so the same operator is being repurposed here as a logical readout rather than a syndrome. Note also that this cell is quoted in the $\pm$ parity read basis, in which $\kappa = 0$; Table II reports the cat/parity cell in the atlas's coherent-state basis, where the same meter is a maximal within-code rotation. Tuning provides another: for real codewords the two perfection conditions are two real equations, so a two-parameter meter family generically contains exact zeros. For square-lattice GKP the phase-shifted modular meter $\cos(\sqrt{\pi}q + \phi)$ obeys, exactly at finite energy Δ , $\kappa(\phi) = A \cos \phi$ and $\Lambda_{01}(\phi) = -C \cos 2\phi + O(A^2)$,

with $C = o e^{-\pi\Delta^2}/(1 - o^2)$ (o the raw codeword overlap) and A doubly exponentially small in $1/\Delta^2$ (zero for ideal combs). At $\phi = 3\pi/4$ the logical confusion vanishes ($\cos 2\phi = 0$) with $O(1)$ signal, while $\kappa = -A/\sqrt{2}$ does not vanish: at fixed wavevector these two forms have no common zero in ϕ , since $\kappa = 0$ forces $\cos 2\phi = -1$ and hence $\Lambda_{01} = C \neq 0$. What makes the point useful is that A is doubly exponentially small, so the residual rotation at $\phi = 3\pi/4$ already sits far below C . Allowing the wavevector to vary alongside the phase, a genuine common zero $\kappa = \Lambda_{01} = 0$ exists within $O(A/C)$ of $(\sqrt{\pi}, 3\pi/4)$ in the resulting two-parameter family — certified numerically by a winding number ± 1 of (κ, Λ_{01}) about that point, not by an analytic transversality argument — one quarter-period phase shift away from the canonical mod- q readout of Table I. At such a tuned meter the leading-order bound vanishes, but because $E \neq 0$ (unlike the exact-QND cat $\pm$ /parity case) a residual $O(\gamma/g)$ self-disturbance is not excluded; the remaining leak is which-path, damaging only the conjugates at leading order. A two-tone modular meter does the same for the two-component cat's protected basis at every α^2 scanned. In restricted sub-menus the obstruction is real but soft: quadratic meters on the two-component cat floor at $\min \max(\kappa^2, \varepsilon_{\text{KL}}^2) \approx 0.5 e^{-4\alpha^2}$ — the squared codeword overlap — so such pairs are *asymptotically* self-reporting. This reframes several practical facts. The cat qubit's reputation for accessible readout amounts, in the convention where the computational states are the coherent states $|\pm\alpha\rangle$, to the statement that logical Z is discriminated by a quadrature measurement whose back-action is dephasing — an error in the already-dominant phase-flip channel, so the measurement is *bias-preserving* rather than back-action-free. (The opposite axis convention is standard in the Kerr-cat literature, where $|\pm\alpha\rangle$ are the X eigenstates; the statement must be read with the convention fixed.) GKP requires the harder modular readout for the same role; and codes whose logical variable coincides with a noise generator (number codes under dephasing) are unreadable-while-protected *by any meter in the Hermitian menu considered here*: the bound floors their self-disturbance whenever the back-action is a Knill-Laflamme-uncorrectable error, with the magnitude of the resulting deficit ($\mathcal{P} \sim 1$) confirmed numerically. The canonical phase readout these codes are designed around is a POVM and is not covered by that statement. It also identifies lattice choice as a *readability* knob for GKP codes, distinct from the distance-optimal (hexagonal) choice and complementary to noise-bias engineering by lattice deformation [21], which tunes the decoding threshold rather than the readout back-action — a deliberately squeezed lattice buys nearly an order of magnitude in readout at a factor-two protection cost, which may be the right trade at readout-limited operating points.

VIII. DISCUSSION

The readability limit is leading order in γ/g , single-mode, Markovian, and — because readout enters as $\gamma\mathcal{D}[M]$ — restricted to Hermitian meters; the strong-measurement regime, multimode codes, POVM meters such as the canonical phase measurement, and the fault-tolerant syndrome setting (where M is a stabilizer rather than a logical) are open. The sharpest internal gap is the one named above: a g -uniform derivation of the leak term as a rate, valid for cascaded recovery basins rather than only direct-return stabilizers, which would put the two halves of the max on the same footing. The functional is in the tradition of information–disturbance trade-offs [22, 23], though it is not one of them specialised: Shitara *et al.* measure information by the classical Fisher information and disturbance by the average loss of quantum Fisher information, whereas here information is read-basis distinguishability and disturbance is a logical decoherence rate. What the code structure adds, and what those formulations have no room for, is the resolution of the disturbance into self and conjugate channels; the theorem is the statement that the code’s Knill–Laflamme data controls the self channel. Together with the results of QEC metrology [4, 5], it suggests a picture in which what a code can sense, what it can hide, and what it can tell you are governed by the same correctability structure. Of those three, only the last is established here; the first is cited rather than derived and the second is already explicit in the metrology literature.

ACKNOWLEDGMENTS

This work received no specific funding. Numerical results were obtained with QuTiP [24]. Generative AI tools (Anthropic Claude, 2026) were used throughout — for literature search, code, manuscript preparation, and the analytic derivations — under the author’s direction. The author verified all results and takes full responsibility for the content and claims of this work. The figure is a schematic-and-data plot of the paper’s computed values, produced with author-directed plotting code; no generative image model was used to create or alter it. The author declares no competing interests.

DATA AVAILABILITY

The code and data that reproduce the numerical results and the figure of this article are openly available at github.com/CogitanAI/gentle-readout, archived at Zenodo (doi:10.5281/zenodo.21464317, which resolves to the current version).

Appendix A: Operational definition of the self-disturbance rate

Let $C_Z(t) = \text{Tr}[Z_L^\dagger e^{\mathcal{L}t}(Z_L P/d)]$ be the infinite-temperature autocorrelation of the read logical operator under the full generator $\mathcal{L} = \mathcal{L}_{\text{stab}} + \gamma\mathcal{D}[M]$; the adjoint normalisation $C_Z(0) = \text{Tr}[Z^\dagger Z P]/d$ is required at $d > 2$, where the Weyl logical is non-Hermitian and $\text{Tr}[Z^2 P]$ vanishes identically. We define the self-disturbance rate as the *initial slope*

$$\Gamma_{\text{self}} \equiv -\frac{d}{dt} \ln C_Z(t) \Big|_{t \rightarrow 0^+}, \quad (\text{A1})$$

which is the convention under which the within-code constant of Theorem 1 is exact. This is not an artefact of projecting onto the read axis: Lemma 1 shows that the Helstrom distinguishability of the evolved read-basis codewords — an operational quantity involving no choice of observable — decays at the same initial rate $2\gamma\kappa^2$.

Extraction at finite time, and the resulting deficit. A numerical extraction cannot be taken at $t \rightarrow 0^+$: there the slope still reflects the bare leakage $\gamma\epsilon^2$ rather than the corrected rate. Rates are therefore read at the lower edge t_- of the window $1/g \ll t_- \ll 1/(\gamma\|PMP\|^2)$, after stabilizer transients have closed but before the measurement has appreciably depolarized the code. The two conventions do not coincide, and we state the difference explicitly rather than absorb it. For the within-code channel the slope is monotone non-increasing in t , with the exact form

$$\begin{aligned} -\frac{d}{dt} \ln C_Z(t) &= \frac{2\gamma\kappa^2}{1 + \frac{n_z^2}{n^2} (e^{2\gamma n^2 t} - 1)} \\ &= 2\gamma\kappa^2 [1 - 2\gamma n_z^2 t + O((\gamma n^2 t)^2)], \end{aligned} \quad (\text{A2})$$

so the windowed value lies *below* the defining initial slope by the one-sided relative factor $2\gamma n_z^2 t_-$, and the lower edge of the window is the closest approach to it. We flag two consequences rather than bury them. First, this deficit is $O(\gamma^2 n_z^2 \kappa^2 t_-)$ in absolute terms; since the window requires $t_- \gg 1/g$, it is *not* covered by the $O(\gamma^2/g)$ remainder of Theorem 1, and the two must be tracked separately. Second, it is proportional to $n_z^2 = (\Delta m/2)^2$ and therefore vanishes identically for a pure within-code rotation ($\Delta m = 0$), which is exactly the regime in which the $2\kappa^2$ term is the binding one. The cells at which the bound is reported to saturate — cat/parity and GKP/mod- p — are of this type, so their agreement is not degraded by the deficit; at the remaining cells the measured rate exceeds the bound by factors of 8 or more, far above the tens of percent the deficit can contribute at an admissible t_- . No tabulated verification is threatened by it, but the initial-slope convention is the one the theorem constrains, and the quoted numbers are systematically the smaller of the two.

A spectral extraction instead reports the decaying-mode eigenvalue $2\gamma n^2 = 2\gamma\kappa^2 + \frac{\gamma}{2}\Delta m^2 \geq 2\gamma\kappa^2$. Its excess over the slope value is $\frac{\gamma}{2}\Delta m^2 = \frac{\gamma}{2}\mathcal{S}(M, Z)$,

the information-gain rate, which lower-bounds Γ_{conj} but equals it only when $\kappa = 0$ — every operator orthogonal to $\vec{n}$ in fact decays at $2\gamma n^2$. The bound of Theorem 1 holds under either convention, and is tight under the slope convention used throughout. Two remarks. (i) C_Z may retain a nonzero plateau: the component of Z_L along the conserved direction of the measurement (the $\vec{n}$ axis of Appendix B) does not decay — this component is precisely the QND-readable part of the logical, and the theorem bounds the decay of its complement. (ii) At times $t \lesssim 1/g$ the slope reflects bare leakage $\gamma\epsilon^2$ rather than the corrected rate; the window excludes this regime.

Appendix B: The within-code term: exact constant

Restricted to the code space, $PMP = \bar{m}P + \vec{n} \cdot \vec{\sigma}$ with $\vec{n} = (\kappa, 0, \Delta m/2)$ in the phase convention fixed in Sec. II; the identity part drops from the dissipator. (Nothing below uses the convention: for a general $\vec{n} = (\kappa_x, \kappa_y, \Delta m/2)$ every step holds with $\kappa^2 = \kappa_x^2 + \kappa_y^2$.) Using $(\vec{n} \cdot \vec{\sigma})\sigma_z(\vec{n} \cdot \vec{\sigma}) = 2n_z(\vec{n} \cdot \vec{\sigma}) - n^2\sigma_z$,

$$\mathcal{D}^\dagger[\vec{n} \cdot \vec{\sigma}](Z) = 2n_z(\vec{n} \cdot \vec{\sigma}) - 2n^2Z. \quad (\text{B1})$$

Decompose $Z = (n_z/n^2)(\vec{n} \cdot \vec{\sigma}) + Z_\perp$: the first component is conserved, the transverse component decays at $2\gamma n^2$, so

$$C_Z(t) = \frac{n_z^2}{n^2} + \frac{\kappa^2}{n^2}e^{-2\gamma n^2 t}, \quad -\dot{C}_Z(0^+) = 2\gamma\kappa^2, \quad (\text{B2})$$

using $n^2 - n_z^2 = \kappa^2$. The constant 2 is an operator identity, nonperturbative in κ , and is the coefficient of the initial slope at $t \rightarrow 0^+$. The windowed slope obeys the exact Eq. (A2) of Appendix A: at the lower edge t_- the within-code self-rate is $2\gamma\kappa^2[1 - 2\gamma n_z^2 t_-]$ to leading order, so $2\gamma\kappa^2$ is the supremum of this contribution over the window and is attained only in the defining $t \rightarrow 0^+$ limit. We emphasise that this one-sided deficit is *not* contained in the $-O(\gamma^2/g)$ remainder of Theorem 1: the window imposes $t_- \gg 1/g$, so $\gamma^2 n_z^2 \kappa^2 t_-$ exceeds γ^2/g by the factor $gt_- \gg 1$. It is a separate, explicitly computable gap between the defining convention and the extraction convention, and it vanishes identically at $\Delta m = 0$.

The constant is moreover basis-free, which forestalls the objection that the read-axis projection is chosen to make it small:

Lemma 1 (Basis-free self-rate). *Let $\Phi_t = e^{\gamma\mathcal{D}[\vec{n} \cdot \vec{\sigma}]t}$ and let $T(t) = \frac{1}{2}\|\Phi_t(|0\rangle\langle 0|) - \Phi_t(|1\rangle\langle 1|)\|_1$ be the Helstrom distinguishability of the evolved read-basis codewords — an operational quantity involving no choice of observable. Then, exactly,*

$$T(t) = \sqrt{\frac{n_z^2}{n^2} + \frac{\kappa^2}{n^2}e^{-4\gamma n^2 t}}, \quad -\dot{T}(0^+) = 2\gamma\kappa^2. \quad (\text{B3})$$

Proof. $\Phi_t(Z)$ has Bloch vector $\vec{v}(t) = (n_z/n^2)\vec{n} + [\hat{z} - (n_z/n^2)\vec{n}]e^{-2\gamma n^2 t}$ with the two components orthogonal,

of squared norms n_z^2/n^2 and κ^2/n^2 ; since $\frac{1}{2}\|\vec{v} \cdot \vec{\sigma}\|_1 = |\vec{v}|$, the display follows, and differentiation at 0^+ gives $2\gamma\kappa^2$. $\square$

The optimal discrimination of the physically evolved codewords thus degrades at exactly the rate the theorem bounds.

Appendix C: The leak term: per-event converse with explicit constant

Setup. Unravel $\gamma\mathcal{D}[M]$ into a no-jump evolution and jump events. Because the code space is the *dark space* of $\mathcal{L}_{\text{stab}}$, a trajectory prepared in the code undergoes no stabilizer dynamics until the first measurement jump: the process on each trajectory factorizes exactly as (jump) followed by (everything after), and “everything after” — stabilizer at any rate g , arbitrarily interleaved — composes into a single CPTP map $\mathcal{R}$ that returns population to the code space. This is the factorization that a channel-level treatment lacks. It is not, however, exact: the post-jump state is the coherent superposition $M|\mu\rangle = M_c|\mu\rangle + E|\mu\rangle$, not a mixture of a code branch and a leaked branch, so conditioning on $P_\perp$ below discards the code-perp coherence. That coherence is bounded by $6\gamma\|E\|m/g$ in trace norm (Lemma 2(i)), so the factorization holds to leading order in $\gamma\|E\|m/g$ and the discarded piece falls inside the remainder of Theorem 1.

Per-event floor. A jump applies M ; conditioning on the leaked branch (projection $P_\perp$) maps the logical eigenstates $|\mu\rangle$ to $|w_\mu\rangle = E|\mu\rangle/\epsilon_\mu$, with overlap $s = \langle w_0|w_1\rangle$ and $\Lambda_{01} = \epsilon_0\epsilon_1 s$. For initial $|0\rangle$ vs $|1\rangle$, the recovered states $\rho'_\mu = \mathcal{R}(|w_\mu\rangle\langle w_\mu|)$ satisfy, by monotonicity of the trace distance under any $\mathcal{R}$, $T(\rho'_0, \rho'_1) \leq T(|w_0\rangle, |w_1\rangle) = \sqrt{1 - |s|^2}$. The average Z -population error per leak event is therefore bounded by the Helstrom limit,

$$\begin{aligned} \delta &\equiv \frac{1}{2}[(1 - \langle 0|\rho'_0|0\rangle) + (1 - \langle 1|\rho'_1|1\rangle)] \\ &\geq \frac{1 - \sqrt{1 - |s|^2}}{2} \geq \frac{|s|^2}{4}. \end{aligned} \quad (\text{C1})$$

No recovery, at any speed, can beat state distinguishability.

Flux identity and rate. For ρ supported on the code space the code block of the full generator is exactly

$$\begin{aligned} \dot{\rho}_c &= \gamma[M_c\rho M_c - \tfrac{1}{2}\{M_c^2, \rho\}] - \tfrac{\gamma}{2}\{\Lambda, \rho\} \\ &\quad + \gamma\mathcal{R}(E\rho E^\dagger) + O(\gamma^2/g), \end{aligned} \quad (\text{C2})$$

with $M_c = PMP$ and $\mathcal{R}$ the return channel of the quasi-static-leak lemma (Appendix D, Lemma 2). Because $\{Z, \Lambda\} = 2\text{diag}(\epsilon_0^2, -\epsilon_1^2)$ *exactly* — the off-diagonal Λ_{01} cancels identically in the anticommutator — the no-jump term moves the Z -populations only through the leak rates, and combining loss and return for initial codeword $|\mu\rangle$ gives the exact per-codeword balance $-\frac{d}{dt}[(-1)^\mu\langle Z\rangle_\mu]_{\text{leak}} = 2\gamma\epsilon_\mu^2\delta_\mu$, hence

$$\Gamma_{\text{self}}^{(\text{leak})} = \gamma(\epsilon_0^2\delta_0 + \epsilon_1^2\delta_1)[1 + o(1)], \quad (\text{C3})$$

with $\delta_\mu = 1 - \langle \mu | \mathcal{R}(|w_\mu\rangle\langle w_\mu|) | \mu \rangle$. (The naïve weighting event-rate $\times$ mean error, $\gamma \frac{\text{Tr}\Lambda}{2} \bar{\delta}$, is *not* a valid accounting of the slope — it overstates the damage whenever the frequently leaking codeword is well recovered — and is not used.) With the Helstrom floor $\delta_0 + \delta_1 \geq |s|^2/2$ and $|\Lambda_{01}|^2 = \epsilon_0^2 \epsilon_1^2 |s|^2$,

$$\begin{aligned} \gamma(\epsilon_0^2 \delta_0 + \epsilon_1^2 \delta_1) &\geq \gamma \min(\epsilon_0^2, \epsilon_1^2) \frac{|s|^2}{2} = \frac{\gamma |\Lambda_{01}|^2}{2 \max(\epsilon_0^2, \epsilon_1^2)} \\ &\geq \frac{\gamma}{2} \frac{|\Lambda_{01}|^2}{\text{Tr}\Lambda}, \end{aligned} \quad (\text{C4})$$

i.e., $c_2 = 1/2$ in the conventions of the main text, uniformly over return channels, via a strictly sharper intermediate. An independent four-level model with explicit stabilizer at rate g confirms the structure of Eq. (C3): the within-code windowed slope equals $2\gamma\kappa^2$ exactly (not the spectral eigenvalue $2\gamma n^2$), the flux identity reproduces the measured slope to within $\sim 10\%$, and the which-path case ($s = 0$, $\epsilon_0 \neq \epsilon_1$) shows no leading self-damage — its residual self-rate is $O(\gamma/g)$, consistent with the diagonal residual feeding only the conjugate. All measured (code, meter) cells with a nonzero bound satisfy it with this constant. Saturation is reached only where the within-code term dominates ($1.0\times$ at cat-parity, $1.04\times$ at GKP-mod- p); the best confusion-dominated cell is GKP-mod- q at $8.1\times$, and the worst, GKP- q , is loose by a factor of 316.

Quasi-static leak sector. The step above uses that, for t in the window, the leaked population and code-perp coherences are $O(\gamma\|\Lambda\|/g)$ in trace norm, and the return flux into the code block equals $\gamma \mathcal{R}(E\rho_c E^\dagger)$ for a CP trace-nonincreasing $\mathcal{R}$ with trace deficit $O(\gamma/g)$, uniformly over the window. This is proven, with explicit constants, in Appendix D (which in fact yields more: the return channel $\mathcal{R}$ is exactly trace preserving). Numerically, the windowed slope at $g = 5\text{--}2000$ stays bounded with no Zeno suppression and converges to within a few percent of the flux value, consistent with an $O(1/g)$ remainder well within the proven bound.

Uniformity in g and linear accumulation. The per-event floor is channel-independent (monotonicity holds for every CPTP $\mathcal{R}$, including any interleaving of recovery with subsequent dynamics), and the event rate is set by γ alone; damage therefore accumulates linearly in t with multi-event corrections $O((\gamma t)^2)$. The quadratic (Zeno-type) suppression familiar from coherent leakage does not apply: measurement jumps are incoherent point events. The diagonal residual ($\epsilon_0 \neq \epsilon_1$, $s = 0$) permits perfect population restoration and damages only the conjugate coherences — it is which-path information about Z , dephasing X, Y at first order in the rate asymmetry. By the exact cancellation $\{Z, \Lambda\} = 2 \text{diag}(\epsilon_0^2, -\epsilon_1^2)$, the no-jump back-action $-(\gamma/2)\{E^\dagger E, \cdot\}$ cannot act on the Z -populations through Λ_{01} at leading order in γt from a state diagonal in the read basis. (It does so at second order: the same term generates a coherence $\dot{\rho}_{01} = -\frac{\gamma}{2}\Lambda_{01}(\rho_{00} + \rho_{11})$, which feeds back into the populations at $O(\gamma^2 t^2 |\Lambda_{01}|^2)$.) Its off-diagonal part dephases

TABLE II. Verification of Theorem 1 across the $d = 2$ (code, meter) grid: measured Γ_{self} against the two bound terms with the derived constants (within-code coefficient 2 exact as an operator identity; $c_2 = 1/2$ what the converse bounding chain yields, not claimed tight). The bound $\Gamma_{\text{self}} \geq \gamma \max(2\kappa^2, \epsilon_{\text{KL}}^2/2)$ is respected at every cell; saturation occurs at the κ -dominated cells. The two photon-number cells (bottom) use an *unbounded* meter, outside the weak-measurement regime $\eta = \gamma\|M\|^2/g \leq 1/32$ under which the leak term is proven; they are shown for context, where the leading bound is expectedly loose.

code	meter	Γ_{self}	$2\gamma\kappa^2$	$\gamma\epsilon_{\text{KL}}^2/2$
cat	q	1.2×10^{-3}	1.9×10^{-4}	1.8×10^{-4}
cat	p	1.5×10^{-3}	1.9×10^{-4}	1.8×10^{-4}
cat	parity	1.0×10^{-1}	1.0×10^{-1}	0
cat	mod- q	1.7×10^{-3}	4.0×10^{-5}	1.2×10^{-5}
cat	mod- p	9.4×10^{-3}	2.3×10^{-3}	4.8×10^{-4}
GKP	q	3.1×10^{-3}	~ 0	9.8×10^{-6}
GKP	p	4.7×10^{-2}	~ 0	3.9×10^{-3}
GKP	parity	1.4×10^{-8}	~ 0	0
GKP	mod- q	1.7×10^{-3}	2.0×10^{-4}	2.1×10^{-4}
GKP	mod- p	7.4×10^{-2}	7.1×10^{-2}	3.4×10^{-7}
<i>Unbounded meter, outside the proven regime (context only):</i>				
cat	n	4.8×10^{-3}	7.2×10^{-4}	3.8×10^{-4}
GKP	n	9.5×10^{-2}	1.6×10^{-2}	5.4×10^{-3}

only the conjugates. Combining Appendices B and C: the within-code and leak contributions are separate summands of the same slope functional, so under Lemma 2 they add at leading order, giving Theorem 1 in additive form and a fortiori in max form.

Appendix D: Quantitative adiabatic elimination

Setting: $\mathcal{H} = \mathcal{H}_c \oplus \mathcal{H}_\perp$; $\mathcal{L} = \mathcal{L}_{\text{stab}} + \gamma \mathcal{D}[M]$ with stabilizer jumps $\{L_k\}$; $m = \|M\|$, $E = P_\perp M P$, $\Lambda = P E^\dagger E P$ (so $\|\Lambda\| = \|E\|^2$), $K = \sum_k L_k^\dagger L_k$, $\beta = \|K\|/g$. Hypotheses: (G1) *dark space*, $L_k P = 0$; (G3) *weak measurement*, $\eta \equiv \gamma m^2/g \leq 1/32$. The remaining hypothesis (G2) is the confinement gap, stated below once the operator it constrains has been identified — which is not the one a reader might expect, and is exactly what confines the lemma to direct-return stabilizers. Split each stabilizer jump out of the leaked block by where it lands, $L_k P_\perp = P L_k P_\perp + P_\perp L_k P_\perp$, and collect the code-landing flux and the perp-internal recycling into

$$\mathcal{J}_c(\sigma) = \sum_k (P L_k P_\perp) \sigma (P L_k P_\perp)^\dagger, \quad \mathcal{J}_\perp(\sigma) = \sum_k (P_\perp L_k P_\perp) \sigma (P_\perp L_k P_\perp)^\dagger \quad (\text{D1})$$

Write $\mathcal{L}'_\perp = -\frac{1}{2}\{K_\perp, \cdot\} + \mathcal{J}_\perp$ for the *code-landing-free* generator, with $K_\perp = P_\perp K P_\perp$: this propagates the leaked block including its internal cascade, and omits only the flux that reaches the code. Inserting $\mathbb{1} = P + P_\perp$

between $L_k^\dagger$ and L_k splits $K_\perp$ exactly into two positive pieces, and the trace balance of $\mathcal{L}'_\perp$ is

$$\frac{d}{d\tau} \text{Tr}[e^{\mathcal{L}'_\perp \tau} \sigma] = -\text{Tr}[G e^{\mathcal{L}'_\perp \tau} \sigma], \quad G = K_\perp - \sum_k A_k^\dagger A_k = \sum_k B_k^\dagger B_k, \quad (\text{D2})$$

with $A_k = P_\perp L_k P_\perp$ and $B_k = P L_k P_\perp$, so that $G \geq 0$ and $\text{Tr}[G\sigma] = \text{Tr}[\mathcal{J}_c(\sigma)]$ identically: the trace lost by $\mathcal{L}'_\perp$ is exactly the code-landing flux. The confinement hypothesis is a lower bound on this drain operator, *not* on $K_\perp$:

(G2) *confinement gap*, $G \geq g P_\perp$.

This is the hypothesis the population step consumes, and the distinction is not cosmetic. Since $\langle \psi | G | \psi \rangle = \sum_k \|P L_k |\psi\rangle\|^2$, (G2) demands that *every* leaked state decay into the code space directly at rate at least g . Direct-return stabilizers ($A_k = 0$, the engineered dissipation of the numerics) satisfy it with $G = K_\perp$. A cascaded basin generically does not: G annihilates any leaked state that cannot reach the code in one hop, so a deep cascade violates (G2) however fast its individual rungs are. A three-level example makes this sharp — code $|0\rangle$, leaked $\{|1\rangle, |2\rangle\}$, $L_1 = \sqrt{g}|1\rangle\langle 2|$, $L_2 = \sqrt{g}|0\rangle\langle 1|$ gives $K_\perp = g P_\perp$ but $G = g|1\rangle\langle 1|$, and the leaked population starting at $|2\rangle$ has zero initial drain.

Extending to cascaded basins therefore requires more than reinterpreting g . Replacing (G2) by an exponential envelope $\|e^{\mathcal{L}'_\perp \tau} \sigma\|_1 \leq C_{\text{ret}} e^{-g\tau} \|\sigma\|_1$ recovers the population bound only through variation of constants, at the cost of a factor $C_{\text{ret}} > 1$ (strictly, whenever $A_k \neq 0$, since a cascade's survival curve has zero initial slope) and a correspondingly tightened weak-measurement threshold. The cross-coherence step also acquires a genuinely new source $\sum_k B_k \rho_{\perp\perp} A_k^\dagger$, absent by construction when $A_k = 0$, which carries the factor $(1 + \beta')$ with $\beta' = \sum_k \|B_k\| \|A_k\|/g$. Both routes are straightforward but neither is carried out here, so the quantitative lemma below is stated for $A_k = 0$.

A second obstruction is independent of all of this and applies to the cat case in particular. For two-photon dissipation, $L = \sqrt{\kappa_2}(a^2 - \alpha^2)$ is unbounded, so $\beta = \|K\|/g$ and hence C_β diverge and remainder (iii) is vacuous. Applying the lemma to a physical cat qubit requires a truncation hypothesis bounding K on the relevant subspace, which we do not supply.

Lemma 2 (Quasi-static leak sector). *For code-supported $\rho(0)$ and all $t \geq 0$: (i) $\text{Tr}[P_\perp \rho(t)] \leq \frac{3}{2} \gamma \|\Lambda\|/g$ and $\|P\rho(t)P_\perp\|_1 \leq 6 \gamma \|E\|m/g$. (ii) The return channel*

$$\mathcal{R}(\sigma) = \int_0^\infty \mathcal{J}_c(e^{\mathcal{L}'_\perp \tau} \sigma) d\tau \quad (\text{D3})$$

is completely positive and exactly trace preserving, and for $t \geq t_ = g^{-1} \ln(1/\eta)$ the slope identity of Appendix C [Eq. (C3)] holds with remainder (iii) $|\text{err}| \leq C_\beta \gamma^2 \|E\|^2 m^2/g$, with the explicit (non-optimized) constant $C_\beta = 26 + 61\beta + 2\beta^2$.*

Proof sketch (full bookkeeping in the companion note).

(i) By the dark-space property the stabilizer sees only the leaked block, so $\dot{p} \leq -gp + \gamma \|\Lambda\| + 2\gamma \|E\|m/g$ for the leaked population p , while Duhamel against the contraction $\|X e^{-K\tau/2}\|_1 \leq e^{-g\tau/2} \|X\|_1$ bounds the cross coherence q ; the drain step uses (G2) in the form $\text{Tr}[G\rho_{\perp\perp}] \geq g \text{Tr}[\rho_{\perp\perp}]$. The resulting linear system closes under (G3), giving the stated suprema (a sup-norm bootstrap). (ii) $\mathcal{R}$ is an integral of CP maps. Its trace on positive σ equals $\text{Tr}\sigma$ by Eq. (D2): the instantaneous trace loss of $\mathcal{L}'_\perp$ is identically the code-landing flux $\text{Tr}[\mathcal{J}_c(\cdot)]$, so integrating from 0 to ∞ and using $e^{\mathcal{L}'_\perp \tau} \sigma \rightarrow 0$, which (G2) supplies, accounts for the whole of $\text{Tr}\sigma$. Trace preservation is stronger than the monotonicity needed by the Helstrom step of Appendix C. (iii) Four error sources — cross-block couplings of $\mathcal{D}^\dagger[M](Z_L)$, non-code-driven sources of the leaked block, freezing $\rho_{cc}(s) \rightarrow \rho_{cc}(t)$ inside the memory integral, and the e^{-gt} tail beyond t_* — are bounded by 26, 58β , $2\beta(1+\beta)$, and β times $\gamma^2 \|E\|^2 m^2/g$ respectively. $\square$

On an independent four-level model ($\beta = 1$, $C_\beta = 89$) the windowed slope stays bounded as g ranges over 5–2000 (no Zeno suppression) and converges to within a few percent of the additive prediction, consistent with the $O(1/g)$ remainder; the proven bound $C_\beta \gamma^2 \|E\|^2 m^2/g$ holds comfortably. With Lemma 2, the additive form of Theorem 1 is proven under (G1)–(G3) — for direct-return stabilizers and $d = 2$ — with an explicit remainder; the constant C_β is an upper bound, not claimed tight. Closing the cascaded case, and supplying the truncation hypothesis that would bring two-photon-dissipation cat qubits inside the quantitative statement, is the sharpest piece of unfinished business in this appendix.

-
- [1] B. Schumacher and M. A. Nielsen, Quantum data processing and error correction, Phys. Rev. A **54**, 2629 (1996), quant-ph/9604022.
[2] E. Knill and R. Laflamme, Theory of quantum error-correcting codes, Phys. Rev. A **55**, 900 (1997), quant-ph/9604034.
[3] D. Kretschmann, D. W. Kribs, and R. W. Spekkens, Complementarity of private and correctable subsystems in quantum cryptography and error correction, Phys.

- Rev. A **78**, 032330 (2008), 0711.3438.
[4] J. Chen, R. Luo, Z. Du, Y. Yan, Y. Zhou, and X. Ma, Bypassing the protection-sensitivity incompatibility in quantum-error-corrected metrology via asymmetric codes (2025), arXiv preprint (v2, June 2026), 2512.20426.
[5] S. Zhou, M. Zhang, J. Preskill, and L. Jiang, Achieving the heisenberg limit in quantum metrology using quantum error correction, Nat. Commun. **9**, 78 (2018),

- 1706.02445.
- [6] D. Layden, S. Zhou, P. Cappellaro, and L. Jiang, Ancilla-free quantum error correction codes for quantum metrology, *Phys. Rev. Lett.* **122**, 040502 (2019), 1811.01450.
 - [7] A. Halaski and C. P. Koch, Obstacles to continuous quantum error correction via parity measurements (2026), arXiv preprint, 2603.02106.
 - [8] A. A. Clerk, M. H. Devoret, S. M. Girvin, F. Marquardt, and R. J. Schoelkopf, Introduction to quantum noise, measurement, and amplification, *Rev. Mod. Phys.* **82**, 1155 (2010), 0810.4729.
 - [9] H. M. Wiseman and G. J. Milburn, *Quantum Measurement and Control* (Cambridge University Press, 2009).
 - [10] V. B. Braginsky and F. Y. Khalili, *Quantum Measurement*, edited by K. S. Thorne (Cambridge University Press, 1992).
 - [11] C. Bény and O. Oreshkov, General conditions for approximate quantum error correction and near-optimal recovery channels, *Phys. Rev. Lett.* **104**, 120501 (2010), 0907.5391.
 - [12] C. Bény, Conditions for the approximate correction of algebras, in *Theory of Quantum Computation, Communication, and Cryptography (TQC 2009)*, LNCS, Vol. 5906 (Springer, 2009) pp. 66–75, 0907.4207.
 - [13] P. Faist, S. Nezami, V. V. Albert, G. Salton, F. Pastawski, P. Hayden, and J. Preskill, Continuous symmetries and approximate quantum error correction, *Phys. Rev. X* **10**, 041018 (2020), 1902.07714.
 - [14] P. T. Cochrane, G. J. Milburn, and W. J. Munro, Macroscopically distinct quantum-superposition states as a bosonic code for amplitude damping, *Phys. Rev. A* **59**, 2631 (1999), quant-ph/9809037.
 - [15] M. Mirrahimi, Z. Leghtas, V. V. Albert, S. Touzard, R. J. Schoelkopf, L. Jiang, and M. H. Devoret, Dynamically protected cat-qubits: a new paradigm for universal quantum computation, *New J. Phys.* **16**, 045014 (2014), 1312.2017.
 - [16] D. Gottesman, A. Kitaev, and J. Preskill, Encoding a qubit in an oscillator, *Phys. Rev. A* **64**, 012310 (2001), quant-ph/0008040.
 - [17] M. H. Michael, M. Silveri, R. T. Brierley, V. V. Albert, J. Salmilehto, L. Jiang, and S. M. Girvin, New class of quantum error-correcting codes for a bosonic mode, *Phys. Rev. X* **6**, 031006 (2016), 1602.00008.
 - [18] A. L. Grimsmo, J. Combes, and B. Q. Baragiola, Quantum computing with rotation-symmetric bosonic codes, *Phys. Rev. X* **10**, 011058 (2020), 1901.08071.
 - [19] L. Sun, A. Petrenko, Z. Leghtas, B. Vlastakis, G. Kirchmair, K. M. Sliwa, A. Narla, M. Hatridge, S. Shankar, J. Blumoff, L. Frunzio, M. Mirrahimi, M. H. Devoret, and R. J. Schoelkopf, Tracking photon jumps with repeated quantum non-demolition parity measurements, *Nature* **511**, 444 (2014), 1311.2534.
 - [20] N. Ofek, A. Petrenko, R. Heeres, P. Reinhold, Z. Leghtas, B. Vlastakis, Y. Liu, L. Frunzio, S. M. Girvin, L. Jiang, M. Mirrahimi, M. H. Devoret, and R. J. Schoelkopf, Extending the lifetime of a quantum bit with error correction in superconducting circuits, *Nature* **536**, 441 (2016), 1602.04768.
 - [21] L. Hänggeli, M. Heinze, and R. König, Enhanced noise resilience of the surface-gottesman-kitaev-preskill code via designed bias, *Phys. Rev. A* **102**, 052408 (2020), 2004.00541.
 - [22] C. A. Fuchs and A. Peres, Quantum-state disturbance versus information gain: Uncertainty relations for quantum information, *Phys. Rev. A* **53**, 2038 (1996), quant-ph/9512023.
 - [23] T. Shitara, Y. Kuramochi, and M. Ueda, Trade-off relation between information and disturbance in quantum measurement, *Phys. Rev. A* **93**, 032134 (2016), 1505.01320.
 - [24] J. R. Johansson, P. D. Nation, and F. Nori, QuTiP 2: A Python framework for the dynamics of open quantum systems, *Comput. Phys. Commun.* **184**, 1234 (2013).