

Dynamic Quantum Resource Partitioning Architecture for Noise-Adaptive Quantum Sensing

Deepinder Sidhu*

Department of Computer Science and Electrical Engineering
University of Maryland, Baltimore County, MD 21250, USA
sidhu@umbc.edu

July 18, 2026

Abstract

Dynamic Quantum Resource Partitioning Architecture (DQRPA) provides a system-level framework for operating quantum-enhanced systems under realistic noise and coherence constraints. Rather than relying on monolithic entangled states whose performance collapses outside laboratory conditions, DQRPA partitions a fixed resource budget into coherence-matched logical blocks, selected by maximizing the retained Heisenberg contribution encoded in a platform-specific retention function. This converts empirical coherence characterization into explicit architectural design rules governing block size, routing feasibility, and operating regimes.

We show that DQRPA generically yields a robust enhanced-SQL scaling plateau, avoids catastrophic performance collapse, and enables noise-aware routing across heterogeneous segments. All platform dependence enters exclusively through the retention function, rendering the framework agnostic to the underlying physical implementation. Logical-block abstraction provides a clean interface for incorporating local mitigation strategies without altering architectural scaling behavior.

DQRPA therefore establishes a unifying architectural layer for scalable, noise-aware quantum enhancement across sensing, networking, and hybrid quantum-classical systems, within the constraints of all realistic hardware platforms.

1 The Heisenberg Dream Meets Operational Reality

Heisenberg scaling,

$$\Delta\phi \propto \frac{1}{N}, \quad (1)$$

represents the ideal benchmark for quantum metrology [1], promising order-of-magnitude gains over standard quantum limit (SQL) strategies,

$$\Delta\phi_{\text{SQL}} \propto \frac{1}{\sqrt{N}}. \quad (2)$$

*This work is partially supported by Cyberspace Analytics Corporation. This work is related to U.S. Provisional Patent Application No. 63/958634, filed January 12, 2026, entitled “Dynamic Quantum Resource Partitioning Architecture for Noise-Adaptive Quantum Sensing.”

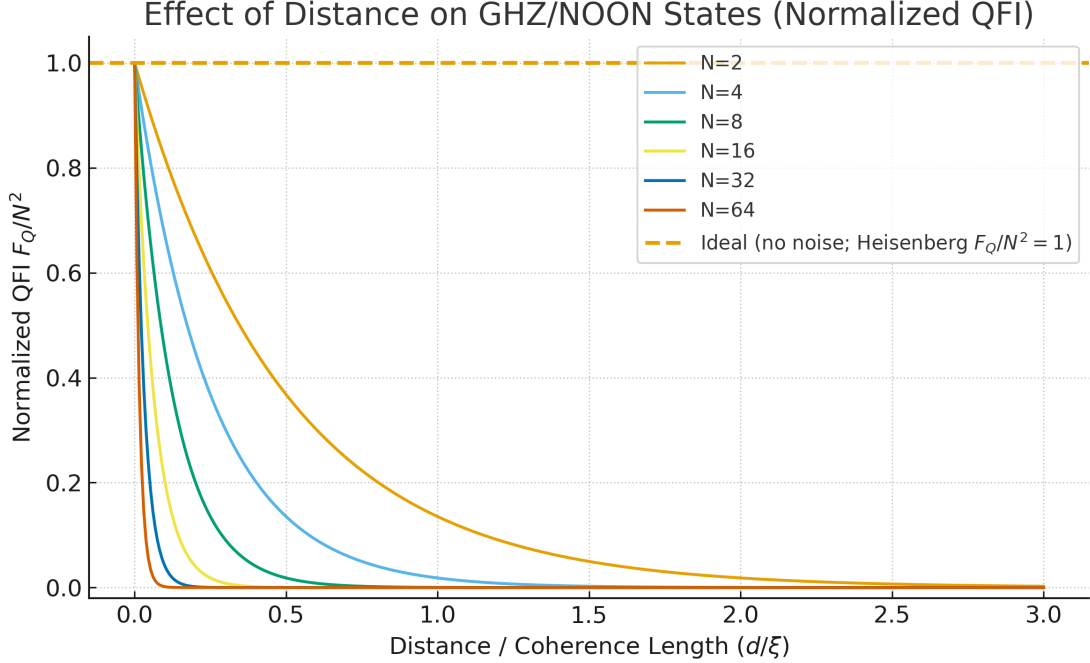


Figure 1: **Fragility of large GHZ/NOON states under noise.** Illustrative normalized QFI F_Q/N^2 versus an effective decoherence parameter (e.g., reduced distance d/ξ) for several particle numbers N . Larger- N GHZ/NOON probes lose their Heisenberg advantage extremely quickly, providing the architectural motivation for noise-aware resource partitioning.

In idealized, noiseless models, maximally entangled GHZ and NOON states can achieve the Heisenberg scaling of Eq. (1): their quantum Fisher information (QFI) scales as

$$F_Q(N) \propto N^2, \quad (3)$$

leading, via the quantum Cramér–Rao bound (QCRB), to

$$\Delta\phi_{\text{HL}}(N) = \frac{1}{\sqrt{F_Q(N)}} \propto \frac{1}{N}. \quad (4)$$

Realistic platforms, however, operate in noisy regimes analyzed by Escher *et al.* and Demkowicz-Dobrzański *et al.*: dephasing, loss, motional decoherence, and control imperfections cause the QFI of large- N GHZ/NOON probes to collapse rapidly as N increases, often eliminating any usable advantage over optimized classical or spin-squeezed strategies [2, 3, 4]. Earlier work in this series introduced the *Heisenberg-saturation* picture and associated precision-reduction laws [5, 6], demonstrating that every noisy platform exhibits a coherence-limited saturation knee beyond which increasing resources ceases to improve—and can even degrade—precision.

Figure 1 illustrates a universal phenomenon: independently of platform details, large- N GHZ/NOON states lose usable QFI after only modest exposure to decoherence. This observation motivates the central architectural question of this paper:

Given a fixed total resource budget N , how should quantum resources be distributed across entangled blocks so as to retain the maximum achievable fraction of Heisenberg scaling under realistic noise?

The conventional answer—concentrating all resources into a single entangled state of size N —is suboptimal and often impractical. The analysis developed here shows that the conventional strategy—concentrating all resources into a single entangled block of size N —is fundamentally suboptimal under realistic noise. Instead, the breakdown of Heisenberg scaling reveals a coherence-limited boundary that constrains how quantum resources may be distributed across entangled subsystems. Identifying how to operate near this boundary motivates architecture-aware resource partitioning strategies, which we do not construct explicitly in the present work.

Throughout this work, achievable phase precision $\Delta\phi$ is expressed in terms of the quantum Fisher information associated with a single architecture-level interrogation. Platform-dependent constants and numerical prefactors are suppressed or treated as $\mathcal{O}(1)$ when discussing scaling with total resource N .

One important exception is the number of statistically independent repetitions, denoted ν_w . With ν_w repetitions, the QCRB becomes

$$\Delta\phi \geq \frac{1}{\sqrt{\nu_w F_Q}}, \quad (5)$$

introducing a global factor of $\nu_w^{-1/2}$.

This factor does not alter any architectural comparisons presented in this paper (SQL versus Heisenberg behavior, saturation knees, or retention laws), since it rescales all strategies uniformly. Accordingly, all results should be interpreted as *single-shot or per-architecture baselines*; repetition improves absolute precision without changing optimal resource-partitioning conclusions.

2 Architectural Resource Partitioning Under Noise

We consider general sensing architectures in which entanglement depth is treated as an explicit architectural resource subject to a global budget. Rather than committing all resources to a single fragile entangled probe, one may distribute the total budget across multiple entangled blocks of varying size, whose performance is constrained by the platform’s coherence envelope.

2.1 Resource accounting and ideal architectural scaling

The total sensing resource is written as

$$N = \sum_j n_j m_j, \quad (6)$$

where m_j is the size (entanglement depth) of the j th block and n_j is the number of such blocks. This representation encompasses all relevant sensing architectures, from fully separable probes to maximally entangled GHZ/NOON states.

In the ideal (noiseless) limit, a block of size m achieves quantum Fisher information

$$F_Q^{\text{ideal}}(m) \propto m^2, \quad (7)$$

and the ideal QFI of the full architecture is additive across blocks:

$$F_Q^{\text{ideal}}(N) \propto \sum_j n_j m_j^2. \quad (8)$$

The quadratic scaling in Eq. (7) follows directly from the pure-state quantum Fisher information identity in Eq. (3). For maximally entangled m -particle blocks, generator covariances scale as $\mathcal{O}(m^2)$, yielding $F_Q^{\text{ideal}}(m) \propto m^2$.

Two familiar strategies appear as limiting cases of Eqs. (6)–(8):

- **GHZ/NOON strategy:** one block of size $m = N$ yields $F_Q^{\text{ideal}} \propto N^2$ and $\Delta\phi \sim 1/N$.
- **SQL strategy:** N blocks of size $m = 1$ yield $F_Q^{\text{ideal}} = N$ and $\Delta\phi_{\text{SQL}} \sim 1/\sqrt{N}$.

DQRPA interpolates continuously between these extremes by allowing arbitrary block partitions consistent with Eq. (6).

2.2 Noise and effective QFI

In realistic platforms, entangled blocks do not retain their ideal QFI. Noise, loss, dephasing, and control imperfections reduce the usable information content of each block. We capture this reduction through an *architecture-level retention function* $R_{\text{arch}}(m)$, defined so that

$$F_Q^{\text{block}}(m) = m^2 R_{\text{arch}}(m), \quad (9)$$

where $0 \leq R_{\text{arch}}(m) \leq 1$ under a fixed normalization.

The effective noisy QFI of a DQRPA architecture is therefore

$$F_Q^{\text{eff}}(N) = \sum_j n_j m_j^2 R_{\text{arch}}(m_j), \quad (10)$$

and the corresponding achievable phase precision is

$$\Delta\phi_{\text{ach}}(N) = \frac{1}{\sqrt{F_Q^{\text{eff}}(N)}}. \quad (11)$$

Equation (10) is the central performance functional of the DQRPA architecture. All architectural optimization questions reduce to choosing a partition $\{m_j, n_j\}$ that maximizes $F_Q^{\text{eff}}(N)$ subject to the resource constraint (6).

This is a constrained architectural optimization under the fixed total resource budget $\sum_j n_j m_j = N$ encoded by (6). No additional additive cost model (e.g., $m + n = C_{\text{max}}$ or $L + M = C_{\text{max}}$) is assumed in this manuscript; such constraints may be introduced for implementation studies without changing the DQRPA performance functional.

2.3 Relation to sequential depth

Some platforms permit additional coherent sequential interactions of depth L within each block. In that case the ideal block scaling generalizes to $(mL)^2$, and the retention becomes $R_{\text{arch}}(m, L)$. In this paper, we focus on *spatial entanglement depth* m to isolate the architectural role of block partitioning; sequential depth effects are treated explicitly in Ref. [6] and can be incorporated into DQRPA without modifying the resource accounting framework.

2.4 DQRPA as a Viable Architectural Solution

The architectural consequences of Eqs. (6)–(10) are immediate and decisive. By replacing monolithic entanglement with coherence-matched block partitioning, DQRPA provides a concrete and viable architectural solution to the fragility problem illustrated in Fig. 1.

Several essential properties distinguish DQRPA from both single-block GHZ/NOON strategies and fully separable SQL approaches:

- **Elimination of catastrophic fragility.** Because entanglement depth is limited to blocks whose size respects the platform’s coherence envelope, DQRPA avoids the rapid collapse of usable quantum Fisher information that afflicts large GHZ/NOON probes under realistic noise.
- **Existence of an optimal entanglement scale.** The retention function $R_{\text{arch}}(m)$ generically implies an optimal block size m^* at which the retained quantum advantage is maximized. Pushing entanglement beyond this scale is counterproductive, while operating below it forfeits available quantum enhancement.
- **Retention of scalable quantum advantage.** Within the coherence-limited operating regime, DQRPA preserves a substantial fraction of Heisenberg-like scaling, outperforming both monolithic entanglement and SQL strategies over a wide and experimentally relevant parameter range.
- **Multidimensional operational viability.** DQRPA simultaneously respects coherence, control, and resource constraints, making it implementable on present and near-term quantum hardware without assuming fault-tolerant error correction or idealized noise models.
- **Platform-agnostic applicability.** The same architectural principles apply across superconducting qubits, photonic sensors, trapped ions, NV centers, and hybrid systems; platform dependence enters only through the form of $R_{\text{arch}}(m)$.

While DQRPA does not eliminate decoherence or remove the ultimate saturation knee, it provides the only known architectural strategy that converts unavoidable noise from a catastrophic failure mode into a quantifiable and optimizable design constraint.

% =====

3 Coherence-Matched Partitioning and Retention Laws

3.1 Resource partitioning as an optimization problem

A DQRPA architecture provides a concrete and implementable solution to the fragility of monolithic entanglement by partitioning the total resource N into entangled blocks of sizes $\{m_j\}$ with multiplicities $\{n_j\}$ satisfying

$$N = \sum_j n_j m_j. \quad (12)$$

In the absence of noise, a block of size m behaves ideally as an m -particle GHZ/NOON probe with QFI scaling as m^2 . Under realistic noise, this ideal scaling is reduced by an architecture-level retention factor $R_{\text{arch}}(m)$.

The DQRPA design problem is therefore explicit and operational: for fixed N , choose $\{(m_j, n_j)\}$ to maximize the effective QFI $F_Q^{\text{eff}}(N)$ subject to Eq. (12).

Noiseless–noisy optimality flip. When $R_{\text{arch}}(m) = 1$, the optimal architecture is a single block $m = N$, yielding Heisenberg scaling. Under noise, $R_{\text{arch}}(m)$ decays with m , reversing this ordering: large blocks decohere fastest, and optimal architectures consist of many moderate-sized blocks constrained by the platform’s coherence envelope.

This reversal is the architectural origin of Heisenberg saturation.

3.2 Retained QFI, retention, and normalization

Under noise, a block of size m contributes

$$F_Q^{\text{block}}(m) = m^2 R_{\text{arch}}(m), \quad (13)$$

with representative retention profiles including

$$\text{Markovian dephasing: } R_{\text{arch}}(m) = e^{-m\gamma t}, \quad (14)$$

$$\text{Photonic loss: } R_{\text{arch}}(m) = \eta^m, \quad (15)$$

$$\text{Correlated dephasing: } R_{\text{arch}}(m) \sim e^{-cm^2}. \quad (16)$$

Summing over blocks yields the total retained QFI $F_Q^{\text{eff}}(N)$ and the achievable precision $\Delta\phi_{\text{ach}}(N) = 1/\sqrt{F_Q^{\text{eff}}(N)}$. These retention profiles encode the unavoidable coherence constraints imposed by realistic hardware, and therefore determine the feasible operating regime of any scalable quantum sensing architecture.

To quantify survival of ideal scaling at the architecture level, we define the Heisenberg-retention function

$$R(N) = \frac{F_Q^{\text{eff}}(N)}{N^2}, \quad (17)$$

so that

$$\Delta\phi_{\text{ach}}(N) = \frac{1}{N\sqrt{R(N)}}. \quad (18)$$

By construction,

$$0 \leq R(N) \leq 1, \quad (19)$$

with $R(N) = 1$ corresponding to ideal Heisenberg scaling.

A complementary SQL-referenced advantage is

$$G(N) = \frac{F_Q^{\text{eff}}(N)}{N}, \quad (20)$$

which distinguishes retained Heisenberg scaling from practical SQL beating.

3.3 Coherence-matched block size

For a broad class of noise models, the retained contribution $m^2 R_{\text{arch}}(m)$ admits a unique interior maximum. For Markovian dephasing,

$$m^* = \frac{2}{\gamma t}. \quad (21)$$

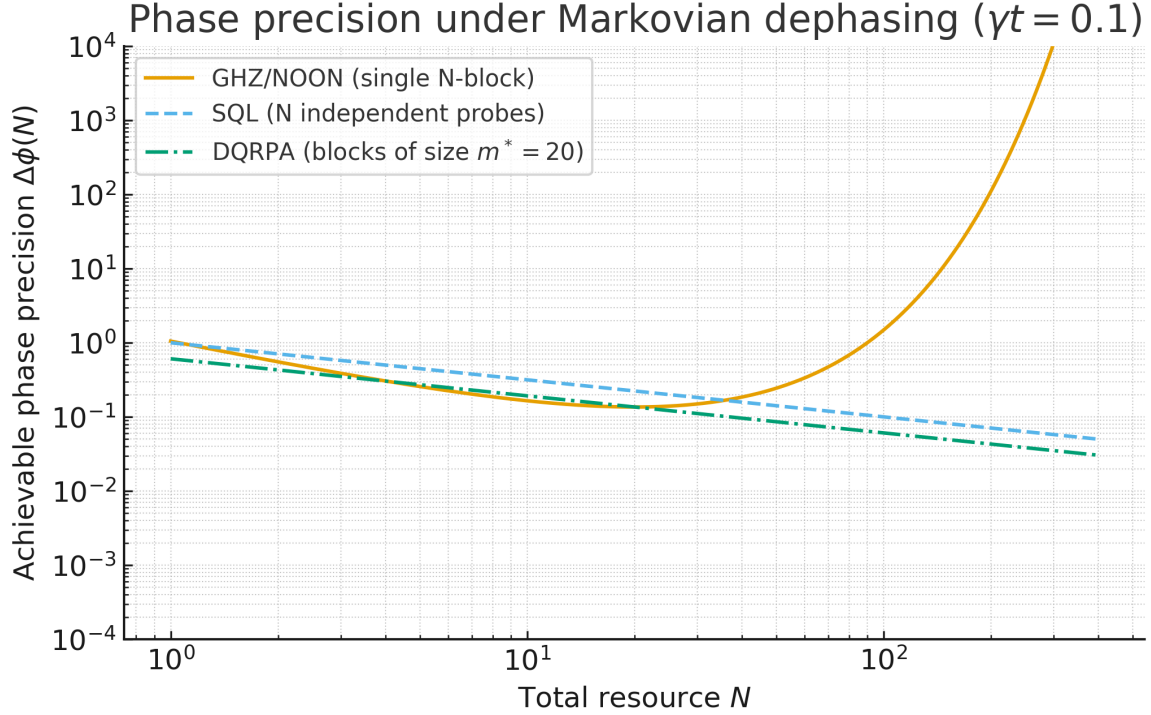


Figure 2: **Phase precision under Markovian dephasing.** Achievable phase precision $\Delta\phi(N)$ versus total resource N for a single N -GHZ/NOON block, the SQL strategy (N independent probes), and a DQRPA architecture built from blocks of size $m^* = 20$. Parameters: dephasing rate–time product $\gamma t = 0.1$. GHZ/NOON briefly beats SQL at small N but then degrades catastrophically; DQRPA stays below SQL over a wide range of N by keeping block size fixed at the noisy optimum m^* .

Importantly, the emergence of the coherence-matched scale m^* is not a consequence of fine-tuning, but a robust architectural feature: broad plateaus of near-optimal performance exist around this scale.

Architectures constructed from $n \approx N/m^*$ such blocks achieve

$$F_Q^{\text{DQRPA}}(N) = N m^* e^{-m^* \gamma t}, \quad (22)$$

with precision

$$\Delta\phi_{\text{DQRPA}}(N) = \frac{e^{\frac{1}{2} m^* \gamma t}}{\sqrt{N m^*}}. \quad (23)$$

When $R_{\text{arch}}(m)$ is non-monotone (e.g., revival structure), m^* is chosen from the maximizing set $\arg \max_m m^2 R_{\text{arch}}(m)$ without altering architectural scaling.

Because the optimization is QFI-based, the location of m^* is independent of the measurement strategy; measurements affect only saturation of the bound, not the architectural optimum itself.

3.4 Architectural interpretation and numerical illustration

Figures 2 and 3 demonstrate that DQRPA provides a viable and scalable architectural solution to the fragility problem illustrated in Fig. 1: precision is maximized not by increasing entanglement depth in-

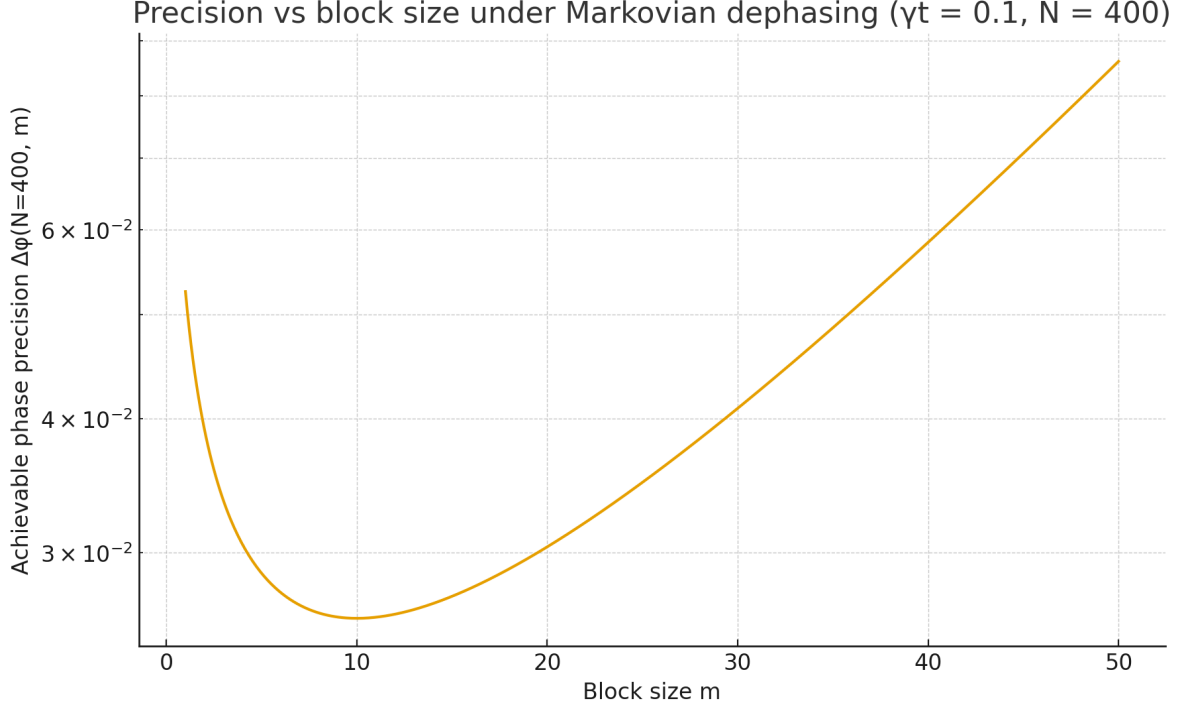


Figure 3: **DQRPA precision as a function of block size for fixed N .** Achievable phase precision $\Delta\phi$ for a DQRPA architecture with total resource $N = 400$ under Markovian dephasing ($\gamma t = 0.1$), plotted as a function of block size m . There is a broad plateau of “good” choices around the coherence-matched optimum $m^* = 20$; for $m \gtrsim 60$ the precision deteriorates rapidly as large blocks become GHZ-like and lose their QFI advantage.

definitely, but by matching entanglement depth to the platform’s coherence envelope. DQRPA therefore operates on an extended enhanced-SQL plateau, preserving quantum advantage while avoiding catastrophic GHZ collapse.

The coherence-matched block size m^* can be understood as the architectural projection of the full (L, M) optimization landscape analyzed in Ref. [5]. Whereas that work exposes the global resource geometry, DQRPA collapses it into an implementation-level design rule: the Heisenberg-saturation knee corresponds to a constrained saddle point of the objective $m^2 R_{\text{arch}}(m)$ under fixed N , marking the boundary of feasible coherence allocation and defining the regime in which scalable, noise-robust quantum enhancement remains operationally viable.

4 DQRPA Implications for Future System Design

Dynamic Quantum Resource Partitioning Architecture (DQRPA) elevates coherence management from a physical-layer concern to a system-level architectural principle. Rather than pursuing monolithic entangled states whose performance collapses outside narrow laboratory regimes, DQRPA provides a viable system-level solution by converting platform-specific coherence behavior into explicit, actionable design rules governing resource allocation, routing, and operational feasibility.

The central insight is that *all platform dependence enters exclusively through the retention profile* $R_{\text{arch}}(m)$, or its path-dependent generalization $R_{\text{arch}}(m; \pi)$. Once this function is characterized (mea-

sured or modeled), the resulting architecture follows deterministically. The concepts below therefore define both the structure of DQRPA and its cross-platform implications for future quantum-enabled systems.

4.1 Core Architectural Rules

- **C1: Coherence-matched blocks (local optimality criterion).** DQRPA replaces monolithic entanglement with coherence-matched logical blocks. The preferred block size m^* is chosen by maximizing the retained Heisenberg contribution of a single block,

$$m^* \in \arg \max_{m \in \mathbb{N}} (m^2 R_{\text{arch}}(m)). \quad (24)$$

This transforms coherence characterization into a directly actionable architectural parameter rather than an abstract performance metric.

- **C2: Enhanced-SQL operating regime (robust scaling plateau).** Partitioning a total resource budget N into $n \approx N/m^*$ coherence-matched blocks yields an effective quantum Fisher information

$$F_Q^{\text{eff}}(N) \approx N m^* R_{\text{arch}}(m^*), \quad \Delta\phi_{\text{ach}}(N) \approx \frac{1}{\sqrt{N m^* R_{\text{arch}}(m^*)}}. \quad (25)$$

The resulting precision scaling is SQL-like but with a coherence-dependent enhancement factor,

$$\frac{\Delta\phi_{\text{SQL}}(N)}{\Delta\phi_{\text{ach}}(N)} \approx \sqrt{m^* R_{\text{arch}}(m^*)}. \quad (26)$$

Crucially, this regime avoids the catastrophic collapse associated with single large GHZ or NOON states in noisy environments, replacing it with a broad, operationally robust performance plateau.

- **C3: Noise-aware routing (path feasibility constraint).** Complex sensing, transport, or processing workflows may be modeled as sequences of spatial, temporal, or network segments s (e.g., time intervals, channel hops, processing stages, or storage operations). Each segment is characterized by a segment-level retention function $r_s(m) \in [0, 1]$, defined as the fraction of ideal m^2 quantum Fisher information preserved by an m -particle entangled block while traversing segment s ,

$$r_s(m) \equiv \frac{F_{Q,s}^{\text{block}}(m)}{m^2}.$$

For a path π composed of multiple segments, the effective architectural retention compounds multiplicatively,

$$R_{\text{arch}}(m; \pi) = \prod_{s \in \pi} r_s(m). \quad (27)$$

DQRPA therefore converts routing into a constrained architectural feasibility problem: only paths for which $R_{\text{arch}}(m; \pi)$ exceeds a minimum coherence threshold admit scalable quantum enhancement.

- **C4: Revival locking (non-Markovian selection rule).** In non-Markovian environments, retention profiles may be non-monotonic and exhibit coherence revivals. DQRPA exploits this structure by

restricting block selection to revival windows $\mathcal{M}_{\text{rev}}$,

$$m^* \in \arg \max_{m \in \mathcal{M}_{\text{rev}}} (m^2 R_{\text{arch}}(m)), \quad (28)$$

effectively locking operation to coherence-favorable regimes rather than treating revivals as incidental physical curiosities.

- **C5: Logical-block abstraction (optional).** Each coherence-matched block may be treated as an effective logical unit, in which case the architectural retention function is modified as

$$R_{\text{arch}}(m) \mapsto R_{\text{arch}}^{(\text{eff})}(m),$$

where $R_{\text{arch}}^{(\text{eff})}(m)$ subsumes any local implementation overhead. Importantly, this abstraction does not alter the routing, partitioning, or scaling structure of DQRPA: the architectural optimizer (Eq. (24)) applies unchanged.¹

4.2 System-Level Design Implications

Taken together, these rules establish DQRPA as a universal system-design framework rather than a platform-specific technique. Large single-block strategies fail because they implicitly assume uniform coherence across scales; DQRPA replaces this assumption with explicit coherence accounting. Retention profiles convert empirical characterization into deterministic architectural decisions, eliminating ad hoc partitioning.

Most importantly, cross-platform generality is structural rather than heuristic. Superconducting circuits, photonic interferometers, trapped ions, NV centers, and hybrid architectures differ only in the form of $R_{\text{arch}}(m)$ and its path-dependent variants. DQRPA supplies a unifying architectural layer that maps these diverse physical behaviors onto a common, feasible optimization problem, enabling scalable, noise-aware, and deployable quantum-enabled system design.

Operational viability. An architecture is operationally viable if it (i) avoids catastrophic performance collapse under realistic noise, (ii) admits scalable resource growth without requiring global error correction, (iii) remains compatible with platform-specific coherence constraints, and (iv) yields deterministic design rules rather than heuristic tuning.

DQRPA satisfies all four criteria by construction. Coherence-matched block partitioning prevents GHZ-like fragility, retention-based optimization replaces idealized scaling assumptions, and platform dependence enters exclusively through the measurable retention profile $R_{\text{arch}}(m)$. As a result, DQRPA provides a viable architectural framework for quantum-enhanced sensing systems operating outside idealized laboratory conditions.

5 Summary and Conclusions

This paper resolves a central tension in quantum metrology: while ideal Heisenberg scaling favors increasingly large entangled probes, realistic hardware environments decohere such probes before their theoretical advantage can be realized. Dynamic Quantum Resource Partitioning Architecture (DQRPA) provides a viable architectural solution to this tension by treating entanglement depth as a tunable system-level resource subject to a fixed global budget.

¹A detailed analysis of how architectural partitioning modifies protection requirements is deferred to future work.

For a general partition $\{(m_j, n_j)\}$ satisfying $\sum_j n_j m_j = N$, the effective noisy quantum Fisher information is

$$F_Q^{\text{eff}}(N) = \sum_j n_j m_j^2 R_{\text{arch}}(m_j), \quad (29)$$

yielding an achievable precision

$$\Delta\phi_{\text{ach}}(N) = \frac{1}{\sqrt{F_Q^{\text{eff}}(N)}}. \quad (30)$$

To quantify the survival of ideal scaling under noise, we introduced the retention function

$$R(N) = \frac{F_Q^{\text{eff}}(N)}{N^2}, \quad (31)$$

so that

$$\Delta\phi_{\text{ach}}(N) = \frac{1}{N\sqrt{R(N)}}. \quad (32)$$

The central architectural outcome is the emergence of a coherence-matched block size m^* , determined directly from the retention profile. In realistic noise regimes, architectures composed of many moderate-sized blocks clustered near m^* maximize retained information and avoid the catastrophic collapse associated with single large GHZ or NOON states. As demonstrated in Figs. 2 and 3, DQRPA operates on a broad enhanced-SQL plateau, delivering robust quantum advantage while remaining compatible with realistic coherence constraints.

DQRPA therefore establishes a scalable, noise-adaptive entanglement architecture for quantum-enhanced sensing and timing systems. Platform dependence enters exclusively through the retention profile $R_{\text{arch}}(m)$, enabling a single architectural optimizer to apply across superconducting, photonic, trapped-ion, NV, and hybrid technologies. This abstraction also provides clean interfaces for routing constraints, non-Markovian revival locking, and optional mitigation layers.

Viewed together with the (L, M) resource-geometry analysis of Ref. [5], DQRPA admits a unifying interpretation: the Heisenberg knee is not a failure of quantum mechanics, but the boundary of feasible coherence allocation under a fixed resource budget. Beyond this boundary, performance loss is architectural rather than physical.

The central lesson is therefore architectural: quantum advantage is not lost because noise exists, but because entanglement is misallocated. By engineering entanglement around coherence constraints rather than idealized scaling laws, DQRPA enables scalable, tunable, and deployable quantum enhancement across realistic hardware platforms.

References

- [1] V. Giovannetti, S. Lloyd, and L. Maccone. Advances in quantum metrology. *Nature Photonics*, 5:222–229, 2011.
- [2] Barbara M. Escher, Roberto L. de Matos Filho, and Luiz Davidovich. General framework for estimating the ultimate precision limit in noisy quantum-enhanced metrology. *Nature Physics*, 7(5):406–411, 2011.
- [3] Rafał Demkowicz-Dobrzański, Marcin Jarzyna, and Jan Kołodyński. Quantum limits in optical interferometry. In Emil Wolf, editor, *Progress in Optics*, volume 60, pages 345–435. Elsevier, 2015.

- [4] L. Pezzè, A. Smerzi, M. K. Oberthaler, R. Schmied, and P. Treutlein. Quantum metrology with nonclassical states of atomic ensembles. *Reviews of Modern Physics*, 90:035005, 2018.
- [5] Deepinder Sidhu. Heisenberg-saturation across quantum sensing platforms: A universal law of noise-limited precision, 2025. Condensed conference manuscript.
- [6] Deepinder Sidhu. Precision reduction in quantum channels: How far from heisenberg must we fall?, 2025. Condensed conference manuscript.