

Quantum AI 2026 Conference

Indiana University Indianapolis
Contact: [Damir Cavar](#) <dcavar@iu.edu>

Table of Contents

Quantum AI 2026 — Conference Program..... 1

Schedule.....2

 Friday, 14 August 2026.....2

 Saturday, 15 August 2026.....3

 Banquet location: Madam Walker Legacy Center4

 Sunday, 16 August 2026.....6

Talks and Abstracts.....8

 Saturday, 15 August 2026.....8

 Sunday, 16 August 2026..... 14

Posters and Abstracts 18

 Poster Overview 18

 Poster Abstracts..... 19

Quantum AI 2026 — Conference Program

**Indiana University Indianapolis — Campus Center, 420 University Blvd., Indianapolis,
IN 46202**

14–16 August 2026

Schedule

Friday, 14 August 2026

Start	End	Session / Talk
3:00 PM	5:00 PM	Registration
5:00 PM	9:00 PM	Reception – Snacks and Drinks

Location: Hine Hall, Tower Ballroom, IU Indianapolis

<https://eventservices.indianapolis.iu.edu/meeting-spaces.html>

Address:

568 Blake St

Indianapolis, IN 46202

Saturday, 15 August 2026

Location: IU Indianapolis, [Campus Center](#)

Address:

420 University Blvd

Indianapolis, IN 46202

Start	End	Session / Talk	Speaker(s)
7:45 AM	8:30 AM	Registration Desk Open	
8:30 AM	9:00 AM	Opening Remarks	Organizers Senator Greg Goode Caroline Chick Jarrold (IU, College of Arts and Sciences)
9:00 AM	9:45 AM	Quantum Technologies in Medicine: Vision and Works in Progress (<i>invited talk</i>)	Charles Bruce and Rickey Carter (Mayo Clinic in Florida)
10:00 AM	10:30 AM	Multi-Objective Quantum Architecture Search for Noise-Resilient Circuits with Barren Plateau Mitigation	Aahan Shah, Armaan Sharma, and Hana Sultana
10:30 AM	11:00 AM	Solving Systems of Linear Equations Using Sample-Based Quantum Diagonalization	Taha Selim
11:00 AM	11:45 AM	Computational Chemistry Enabled by Quantum Computers – An Overview of Cleveland Clinic Contributions in Quantum Computing (<i>invited talk</i>)	Milana Bazayeva (Cleveland Clinic Research)
12:00 PM	1:30 PM	Poster Session & Lunch	
1:30 PM	2:00 PM	Learning Relationship between Quantum Walks and Underdamped Langevin Dynamics	Yazhen Wang
2:00 PM	2:30 PM	Byzantine-Resilient Federated Learning via QUBO-Based Client Selection on Quantum Annealers	Andras Ferenczi, Dagen Wang, and Sutapa Samanta
2:30 PM	3:00 PM	Quantum-Classical Reservoir	Mehdi Khalaj, Lingling Jin,

Start	End	Session / Talk	Speaker(s)
		Computing to Predict Influenza H3N2 Antigenic Distance	and Steven Rayan
3:00 PM	4:00 PM	Poster Session & Coffee Break	
4:00 PM	4:30 PM	Dedicated Neuronal Circuitry as a Hopf Bialgebraic Mechanism for Syntax Supports the Strong Minimalist Thesis: Implications for Large Language Models	Laurent Dekydtspotter, A. Kate Miller, Rodica Frimu, Justina Eshun, and Micheal Adeoye
4:30 PM	5:00 PM	Lean-Typed DisCoCat Reduction Certificates with Externally-Attributed Semantics	Rob Sneiderman
5:00 PM	5:45 PM	From Relational Intelligence to Quantum Entertainment (<i>invited talk</i>)	Bob Coecke (Relational Intelligence)

Banquet location: Madam Walker Legacy Center

See the instructions below for parking and entrance.

MADAM WALKER

LEGACY CENTER

As you make your way to Madam Walker Legacy Center there are a few things to note!

PARKING: Attendees can park in the free surface lots A, B, C and D, as well as another free surface lot across the street (Purdue Lot 92).



ENTRANCE: Attendees need to enter through the 617 Indiana Ave entrance (please note this is not the main theater entrance). Once inside, please take the elevator up to the ballroom floor or Foyer. There will be signage to direct you to the event area.

Sunday, 16 August 2026

Location: IU Indianapolis, [Campus Center](#)

Address:

420 University Blvd

Indianapolis, IN 46202

Start	End	Session / Talk	Speaker(s)
7:45 AM	8:30 AM	Registration Desk Open	
8:30 AM	9:15 AM	Enabling Scientific Discovery with Generative Quantum AI (<i>invited talk</i>)	Jasmine Brewer (Quantinuum)
9:30 AM	10:00 AM	Variational Quantum Transformer Architecture for Synthetic Language Generation	Julian Hager, Michael Koelle, Gerhard Stenzel, Tobias Rohe, Jonas Stein, and Claudia Linnhoff-Popien
10:00 AM	10:30 AM	Quantum Personas: A Q-NLP Formalism for Context-Dependent Character Dynamics	Amir Reza Asadi
10:30 AM	11:00 AM	CLAQS: Compact Learnable All-Quantum Token Mixer with Shared-ansatz for Text Classification	Junhao Chen, Yifan Zhou, Hanqi Jiang, Yi Pan, Wei Zhang, Yingfeng Wang, and Tianming Liu
11:00 AM	11:45 AM	Federal Quantum Algorithm Development Pathways (<i>invited talk</i>)	Ethan Krimins (Quantum Research Sciences)
12:00 PM	1:15 PM	Poster Session & Lunch	
1:15 PM	1:45 PM	The production of meaning in the processing of natural language	Christopher J. Agostino, Quan Le Thien, Nayan D'Souza, and Louis van der Elst
1:45 PM	2:15 PM	Evaluating Hybrid Quantum and Classical Proxies for Reinforcement Learning-Driven Text Instance Selection	Santiago Galiano-Segura, Juan Pablo Consuegra-Ayala, Olga Frances, and Elena Lloret
2:15 PM	2:45 PM	Training Data Set Effects in LLM-Assisted Quantum Portfolio	Pradyot Bathuri

Start	End	Session / Talk	Speaker(s)
		Optimization	
2:45 PM	3:00 PM	Coffee Break	
3:00 PM	3:30 PM	Quantum-Structured World Models (QSWMs) for Predictive Latent Dynamics	Hailong Jiang, Emran Hossain, Feng Yu, Jianfeng Zhu, Guilin Zhang, and Wulan Guo
3:30 PM	4:00 PM	Learning to Program Multi-Chip Quantum Neural Networks via Fast Weights	Samuel Yen-Chi Chen, Chun-Hua Lin, Jiun-Cheng Jiang, Kuo-Chung Peng, Yifeng Peng, Junghoon Justin Park, Huan-Hsin Tseng, Hsin-Yi Lin, Kuan-Cheng Chen, Chen-Yu Liu, and Shinjae Yoo
4:00 PM	4:45 PM	The Next Decade of Quantum AI: A Roadmap from Computation to Discovery (<i>invited talk</i>)	Keeper L. Sharkey (ODE)
5:00 PM		Closing Remarks	Organizers
5:30 PM	8:00 PM	Individual and Small Group Dinners	

All sessions take place at the Campus Center, IU Indianapolis. Saturday's banquet is at 6:30 PM in the Madam Walker Theater, see instructions above; Sunday concludes with individual and small group dinners in Indianapolis.

Talks and Abstracts

Saturday, 15 August 2026

1. Quantum Technologies in Medicine: Vision and Works in Progress

Charles Bruce and Rickey Carter (Mayo Clinic in Florida)

Invited talk · Saturday, 15 August 2026, 8:45 AM–9:30 AM

Abstract to be announced.

2. Multi-Objective Quantum Architecture Search for Noise-Resilient Circuits with Barren Plateau Mitigation

Aahan Shah, Armaan Sharma, and Hana Sultana

Contributed talk · Saturday, 15 August 2026, 9:45 AM–10:15 AM

Variational quantum algorithms on near-term devices face two compounding obstacles: barren plateaus that suppress training gradients in deep or highly entangled circuits, and gate noise that degrades fidelity as circuit depth grows. Existing quantum architecture search (QAS) methods optimize for task accuracy alone, ignoring the trainability–depth–noise trade-off that determines whether a discovered circuit can be trained and executed reliably. We introduce Multi-Objective Quantum Architecture Search (MO-QAS), which jointly optimizes accuracy, circuit depth, and a surrogate-predicted trainability score using NSGA-II. Candidate architectures are encoded as typed bipartite graphs and evaluated by a Graph Transformer surrogate, eliminating expensive quantum simulation from the evolutionary inner loop. The search space extends beyond gate sequences to include cost-function locality, initialization strategy, and classical interface configuration, each of which has demonstrated an impact on gradient variance. On a digitsclassification task, the final Pareto front is dominated by shallow decoded architectures (logical depths 2–3) with hardware-efficient parametrization; direct simulation of Pareto representatives confirms task accuracies of 0.86–0.98. Noise simulation under depolarizing, amplitude-damping, and phase-damping channels confirms 1.3–15× higher state fidelity than depth-matched baselines at 10% error probability.

3. Quantum Circular Convolutional Neural Networks (QCirCNN): A parameter and energy efficient model for computer vision

Daphne Wang, Anthony Walsh, Hugo Fruchet, Ariane Soret, Pierre-Emmanuel Emeriau, and Shane Mansfield

Contributed talk · Saturday, 15 August 2026, 10:15 AM–10:45 AM

Convolutional neural networks (CNN) are foundational architectures in deep learning, especially in computer vision, due to their inductive bias of translation invariance. On the other hand, while research in quantum machine learning has highlighted the potential of quantum models to reduce the number of training parameters, fundamental issues hinder the realisation of convolutional operations in quantum systems. In this work, we propose a quantum circular convolutional neural network (QCirCNN) based on circulant matrices. This model is designed to be implemented on photonic quantum computers using single photons and existing photonic technology. We trained this model on the MNIST and CIFAR10 classification datasets, and observed accuracies on par with classical CNNs, while requiring only 6% of the number of parameters. We then further analysed the resource requirements of the model. We found that the QCirCNN required significantly fewer parameters than classical models to reach the same level of accuracy. In addition, we examine the energy consumption of the quantum model on realistic hardware and identify a regime of optimal energy efficiency where quantum models are more efficient than classical CNNs. These results suggest that QCirCNN offers an efficient and hardware-aligned approach to vision tasks.

4. Solving Systems of Linear Equations Using Sample-Based Quantum Diagonalization

Taha Selim

Contributed talk · Saturday, 15 August 2026, 10:45 AM–11:15 AM

Solving systems of linear equations is essential across a wide range of scientific and engineering applications. Classical algorithms for solving large-scale linear systems can be computationally expensive, particularly as problem size grows. Quantum algorithms, such as the Harrow–Hassidim–Lloyd (HHL) algorithm, have shown promise in providing exponential speedups for solving linear systems under specific conditions — namely, a sparse or efficiently simulable matrix, a well-conditioned system, and efficient preparation of the input state. However, practical implementation of HHL faces significant challenges, including the need for fault-tolerant quantum computers, a large number of ancilla qubits for precise phase estimation, and the fact that extracting the full solution vector via measurement can offset much of the theoretical speedup.

In this work, we propose an alternative approach based on sample-based quantum diagonalization (SQD). SQD has proven efficient and practical for ground-state estimation problems in quantum chemistry and many-body physics, where a shallow quantum circuit samples configurations that are then classically assembled into a reduced subspace matrix and diagonalized. We extend this approach to linear systems $Ax=b$ by using time-evolution of the input vector b under A to bias quantum sampling toward configurations relevant to the specific solution being sought, then solving the resulting reduced linear system classically within the sampled subspace. We benchmark this approach on a range of test linear systems, comparing its accuracy and resource requirements against both classical solvers and the HHL algorithm. We further analyze the computational complexity

of the method, characterize the conditions under which the sampled subspace remains small relative to the full problem size, and evaluate scalability to larger systems.

5. Computational Chemistry Enabled by Quantum Computers – An Overview of Cleveland Clinic Contributions in Quantum Computing

Milana Bazayeva (Cleveland Clinic Research)

Invited talk · Saturday, 15 August 2026, 11:15 AM–12:00 PM

Abstract to be announced.

6. Learning Relationship between Quantum Walks and Underdamped Langevin Dynamics

Yazhen Wang

Contributed talk · Saturday, 15 August 2026, 1:15 PM–1:45 PM

Fast computational algorithms are in constant demand, and their development has been driven by advances such as quantum speedup and classical acceleration. This paper intends to study search algorithms based on quantum walks in quantum computation and sampling algorithms based on Langevin dynamics in classical computation. On the quantum side, quantum walk-based search algorithms can achieve quadratic speedups over their classical counterparts. In classical computation, a substantial body of work has focused on gradient acceleration, with gradient-adjusted algorithms derived from underdamped Langevin dynamics providing quadratic acceleration over conventional Langevin algorithms.

Since both search and sampling algorithms are designed to address learning tasks, we study learning relationship between coined quantum walks and underdamped Langevin dynamics. Specifically, we show that, in terms of the Le Cam deficiency distance, a quantum walk with randomization is asymptotically equivalent to underdamped Langevin dynamics, whereas the quantum walk without randomization is not asymptotically equivalent due to its high-frequency oscillatory behavior. We further discuss the implications of these equivalence and nonequivalence results for the computational and inferential properties of the associated algorithms in machine learning tasks. Our findings offer new insight into the relationship between quantum walks and underdamped Langevin dynamics, as well as the intrinsic mechanisms potentially underlying quantum speedup and classical gradient acceleration.

7. Byzantine-Resilient Federated Learning via QUBO-Based Client Selection on Quantum Annealers

Andras Ferenczi, Dagen Wang, and Sutapa Samanta

Contributed talk · Saturday, 15 August 2026, 1:45 PM–2:15 PM

Federated Learning (FL) is a well-established method for training a global model in a decentralized fashion, thus preserving the privacy of local data and eliminating the need for large centralized infrastructure. FL at scale runs the risk of malicious updates that can impact the global model. To mitigate such risks, a number of Byzantine-resilient aggregation methods have been introduced, one of the most popular ones being MultiKrum. Whereas this method scores the individual gradients against their nearest neighbors, it can easily miss malicious updates that preserve the statistical properties of honest updates. To mitigate such concerns, we propose a quantum annealing approach that reformulates client selection as a Quadratic Unconstrained Binary Optimization (QUBO) problem, encoding pairwise distances into a cost function solved by quantum annealers (QA). Unlike MultiKrum’s greedy per-client scoring, the QUBO formulation jointly optimizes over all possible client subsets to find the mutually closest group of m clients. Our experiments at small scale consisting of 15 clients provided promising results for the most challenging Byzantine attacks. For example, the Advanced LIE attack was detected with 95.11% accuracy with QUBO versus 81.33% for classical MultiKrum on an MNIST-trained model and 97.78% vs. 75.56% on CIFAR-10. Whereas our QUBO method performs well against sophisticated attacks, it fares poorly against simpler attacks, which are handled well by classical MultiKrum. Our experiments show that the two methods complement each other well. Additionally, the QUBO method’s quality degrades as we increase the number of clients. To address these shortcomings, we introduce a MultiSignal ensemble that uses a dual-feature routing gate based on Euclidean and cosine Krum score gaps to classify attacks into four regimes and routes evasion attacks to a suspicion-penalized QUBO with agreement voting. The MultiSignal ensemble performs better than classical MultiKrum in all our experiments, with 95.3% average detection accuracy versus 91.8% for classical. The largest gains are achieved on Sparse Lie, where detection accuracy improves from 72.0% to 95.2%, a 23.2 percentage-point gain, and on Advanced Lie, where it increases from 80.4% to 85.2%, a 4.8 percentage-point gain. These results demonstrate that QUBO-based quantum annealing with the MultiSignal ensemble is a principled and scalable defense against the most challenging Byzantine strategies in federated learning.

8. Quantum-Classical Reservoir Computing to Predict Influenza H3N2 Antigenic Distance

Mehdi Khalaj, Lingling Jin, and Steven Rayan

Contributed talk · Saturday, 15 August 2026, 2:15 PM–2:45 PM

Accurate prediction of antigenic distance between influenza A/H3N2 strains is essential for timely vaccine strain selection, yet traditional hemagglutination inhibition (HI) assays are labour-intensive and limited in throughput. We present FluQRC, a hybrid Quantum-Classical Reservoir Computing framework for sequence-based antigenic distance prediction. On two datasets spanning 1963–2002 (271 strains, 73,441 pairs) and 2003–2025 (888 strains, 788,544 pairs), FluQRC is competitive with four established baselines

on the historical set and outperforms all of them on the larger, more challenging 2003–2025 set, where it attains MAE = 0.369, RMSE = 0.635, and $R^2 = 0.900$ —a 20.3% reduction in MAE and 14.3% reduction in RMSE over the strongest baseline, raising R^2 from 0.862 to 0.900. These results demonstrate the scalability and effectiveness of FluQRC for large-scale antigenic distance prediction.

9. Dedicated Neuronal Circuitry as a Hopf Bialgebraic Mechanism for Syntax Supports the Strong Minimalist Thesis: Implications for Large Language Models

Laurent Dekydtspotter, A. Kate Miller, Rodica Frimu, Justina Eshun, and Micheal Adeoye

Contributed talk · Saturday, 15 August 2026, 3:30 PM–4:00 PM

This paper proposes that syntactic computations in the human language faculty are implemented by a dedicated neuronal mechanism implementing a Hopf bialgebraic structure. Building on the Strong Minimalist Thesis (SMT), we argue that the core operation Merge can be modeled as a product–coproduct interaction over tree-based representations, corresponding to complementary neural processes across bilaterally organized cortical and subcortical circuits. Drawing on evidence from EEG, fMRI, and neurolinguistic studies, we show that synchronous homotopic coactivations across hemispheres support an entangled representational system in which combined structures and accessible substructures are simultaneously maintained. We further argue that externalization processes, including linearization, follow from elementary set-theoretic operations compatible with this architecture. Finally, we consider implications for large language models, suggesting that integrating an innate, SMT-consistent computational mechanism is necessary for capturing the efficiency and structure of human language processing.

10. Polarity Cones: Natural Logic Meets Quantum Linguistics

Peter Ludlow

Contributed talk · Saturday, 15 August 2026, 4:00 PM–4:30 PM

Two research programs have independently converged on a similar picture of semantic space. The natural logic tradition (Ladusaw 1979; van Benthem 2008; Moss 2010, 2012; Icard and Moss 2014) holds that linguistic inference is governed by polarity-sensitive monotonicity: $\uparrow$ and $\downarrow$ markers encode the direction of admissible inference at each position in a derivation. The geometric, or order-embedding, tradition in distributional semantics (Vendrov et al. 2016; Ganea et al. 2018) holds that lexical meaning requires order-theoretic rather than metric geometry, representing concepts as entailment cones. This paper argues that the two programs describe the same structure, related by an order-reversing isomorphism between the entailment order on meanings and the inclusion order

on cones: a polarity marker selects a cone direction, polarity controls which cones may intersect, and the intersection in turn dictates how polarity propagates.

We isolate one closure condition on derivations: Cone Preservation under COPY, the requirement that a COPY operation preserve what a structure denotes — its cone. This does the work of the p-scope junction gates of Ludlow and Živanović (2022) ($\vee$ admits only $\downarrow$, $\wedge$ only $\uparrow$). The restrictedness, and thus the conservativity, of natural-language determiners are likewise instances of this single algebraic invariance condition rather than independent stipulations; non-conservative determiners are excluded by the very condition that licenses the conservative ones. The reduction holds in the distributive fragment in which a determiner’s arguments are interpreted, and rests on a polarity-identification lemma which we prove for the p-scope fragment. We close by conjecturing that the lattice this calculus generates over derivations is orthomodular, the non-classical algebra whose Hilbert-space realization quantum-inspired models of meaning exploit.

11. Lean-Typed DisCoCat Reduction Certificates with Externally-Attributed Semantics

Rob Sneiderman

Contributed talk · Saturday, 15 August 2026, 4:30 PM–5:00 PM

We present a proof-of-concept Lean 4 typing certificate for one fixed DisCoCat pregroup reduction: the canonical “Alice loves Bob” cup contraction from $n \otimes (nr \otimes s \otimes nl) \otimes n$ to s . The artifact uses a dependently typed intermediate representation in which morphisms are indexed by source and target object lists. We are precise about what is machine-checked. What Lean checks for this example is exactly two things: (i) typing—the typed target term `aliceLovesBob : Hom [] [S]` elaborates and two malformed cup routings fail to elaborate (shown inline in §3); and (ii) one toy inequality—the hand-assigned additive annotation Γ_{toy} evaluates to 4 before the contraction and 1 after, with the lemma certifying the toy score inequality (after $\leq$ before) for this single instance. The worked reduction has object-list length at most five; no width theorem and no graph-theoretic treewidth computation is mechanized in this paper, and the tree-width field is a hand-assigned component label, not a verified width. Two further things are reported but are not part of the Lean-verified claim. Semantic preservation $[[D]] = [[D']]$ of DisCoCat denotations is not machine-checked here; it is attributed to known rigid and compact-closed category identities.

12. TBA

Bob Coecke (Relational Intelligence)

Invited talk · Saturday, 15 August 2026, 5:00 PM–5:45 PM

Abstract to be announced.

Sunday, 16 August 2026

13. Enabling Scientific Discovery with Generative Quantum AI

Jasmine Brewer (Quantinuum)

Invited talk · Sunday, 16 August 2026, 8:30 AM–9:15 AM

Abstract to be announced.

14. Variational Quantum Transformer Architecture for Synthetic Language Generation

Julian Hager, Michael Koelle, Gerhard Stenzel, Tobias Rohe, Jonas Stein, and Claudia Linnhoff-Popien

Contributed talk · Sunday, 16 August 2026, 9:30 AM–10:00 AM

We propose a compact NISQ-compatible quantum transformer architecture for synthetic QNLP sequence modelling. The model preserves the autoregressive next-token interface of a classical transformer, but replaces attention and feed-forward sublayers with variational quantum encoder blocks, connector circuits, decoder blocks and a direct two-qubit measurement readout. Token contexts are angle-encoded into small quantum registers, processed by parallel variational heads and encoder integration circuits and conditioned through decoder ancillae to produce a distribution over a four-token vocabulary. We evaluate several architecture variants on deterministic and lexicographic grammar-generation tasks against a compact classical transformer baseline. The quantum models are trainable end-to-end and learn nontrivial grammar structure, including perfect deterministic generation in individual runs and high lexicographic validity in the strongest variant. The classical baseline remains more accurate and stable and the quantum models are sensitive to initialization. The contribution is therefore not a claim of quantum advantage, but a concrete architecture and evaluation of transformer-inspired QNLP sequence modelling under near-term quantum constraints.

15. Quantum Personas: A Q-NLP Formalism for Context-Dependent Character Dynamics

Amir Reza Asadi

Contributed talk · Sunday, 16 August 2026, 10:00 AM–10:30 AM

When prompting language models to adopt a persona, character is typically declared once in a static system prompt, ignoring a basic fact of personality psychology: the same identity expresses itself differently across situations, and how it expresses itself depends on what came before. We present a proof of concept demonstrating that this context-dependence can be modelled as grammar-driven quantum dynamics. A persona is formalized as a contextual quantum state on a behavioural Hilbert space; conditional

signatures (how traits co-vary across situations) are encoded as entanglement. Incoming utterances compile to unitaries via pregroup grammar, and generation is a measurement whose basis is set by the question’s grammar. We instantiate this on a conversational corpus. Our goal is persona fidelity: grounding generation in how a person actually behaves in a given situation significantly improves fidelity to their held-out conduct over a static persona prompt. As a proof of concept that the underlying formalism is well-founded, we further show that the construction realises three exact, non-classical signatures, each with a classical control that fails, confirming that the persona state behaves as a genuine quantum system rather than a classical model in disguise.

16. CLAQS: Compact Learnable All-Quantum Token Mixer with Shared-ansatz for Text Classification

Junhao Chen, Yifan Zhou, Hanqi Jiang, Yi Pan, Wei Zhang, Yingfeng Wang, and Tianming Liu

Contributed talk · Sunday, 16 August 2026, 10:30 AM–11:00 AM

Quantum compute is scaling fast, from cloud QPUs to high throughput GPU simulators, making it timely to prototype quantum NLP beyond toy tasks. However, devices remain qubit limited and depth limited, training can be unstable, and classical attention is compute and memory heavy. This motivates compact, phase aware quantum token mixers that stabilize amplitudes and scale to long sequences. We present CLAQS, a compact, fully quantum token mixer for text classification that jointly learns complex-valued mixing and nonlinear transformations within a unified quantum circuit. To enable stable end-to-end optimization, we apply ℓ^1 normalization to regulate amplitude scaling and introduce a two-stage parameterized quantum architecture that decouples shared token embeddings from a window-level quantum feed-forward module. Operating under a sliding-window regime with document-level aggregation, CLAQS requires only eight data qubits and shallow circuits, yet achieves 91.64% accuracy on SST-2 and 87.08% on IMDB, outperforming both classical Transformer baselines and strong hybrid quantum–classical counterparts.

17. Federal Quantum Algorithm Development Pathways

Ethan Krimins (Quantum Research Sciences)

Invited talk · Sunday, 16 August 2026, 11:00 AM–11:45 AM

Abstract to be announced.

18. The production of meaning in the processing of natural language

Christopher J. Agostino, Quan Le Thien, Nayan D’Souza, and Louis van der Elst

Contributed talk · Sunday, 16 August 2026, 1:15 PM–1:45 PM

Understanding the fundamental mechanisms governing the production of meaning in the processing of natural language is critical for designing safe, thoughtful, engaging, and empowering human-agent interactions. If meaning is constituted rather than retrieved, then the search for context-independent features or circuits in the pursuit of mechanistic interpretability may be fundamentally limited. Experiments in cognitive science and social psychology have demonstrated that human semantic processing exhibits contextuality more consistent with quantum logical mechanisms than classical Boolean theories, and recent works have found similar results in large language models—in particular, clear violations of the Bell inequality in experiments of contextuality during interpretation of ambiguous expressions. In this work, we explore the CHSH $|S|$ parameter—the metric associated with the inequality—across the inference parameter space of models spanning four orders of magnitude in scale and cross-reference our findings with MMLU, hallucination rate, and nonsense detection benchmarks. We find that the interquartile range of the $|S|$ distribution is completely orthogonal to all external benchmarks, while overall violation rate shows weak anticorrelation with all three benchmarks. We investigate how $|S|$ varies with sampling parameters and word order, and discuss the information-theoretic constraints that genuine contextuality imposes on prompt injection defenses and its human analogue, whereby careful construction and maintenance of social contextuality can be carried out at scale, shaping the space of possible interpretations before any particular one is reached. We consider the implications for mechanistic interpretability and how genuine contextuality sets an information-theoretic bound on the decomposability of semantic processing, such that no context-independent assignment of meanings to internal representations can fully account for interpretive behavior.

19. Evaluating Hybrid Quantum and Classical Proxies for Reinforcement Learning-Driven Text Instance Selection

Santiago Galiano-Segura, Juan Pablo Consuegra-Ayala, Olga Frances, and Elena Lloret

Contributed talk · Sunday, 16 August 2026, 1:45 PM–2:15 PM

Fine-tuning pre-trained language models requires significant computational resources, making instance selection and curation essential. While reinforcement learning has automated dynamic subset selection using classical proxy models, the stability and expressivity limits of quantum architectures specifically gate-based Variational quantum circuits remain largely unexplored. In this paper, we propose a proxy-guided reinforcement learning framework for dataset reduction that evaluates and compares hybrid variational quantum circuit against conventional multi-layer perceptron proxies. By testing across distinct geometric clustering topologies on both low-entropy and high-entropy sentiment analysis datasets, we empirically define a clear architectural trade-off. Our findings demonstrate that classical multi-layer perceptron proxies exhibit remarkable topological resilience against high-entropy text noise, whereas hybrid variational quantum circuit exhibits better performance under favorable clustering topologies and well-structured semantic distributions.

20. Training Data Set Effects in LLM-Assisted Quantum Portfolio Optimization

Pradyot Bathuri

Contributed talk · Sunday, 16 August 2026, 2:15 PM–2:45 PM

We observe the possible oversight of using large language models in portfolio optimization pipelines and include in the pipeline a language-conditioned Quantum Approximate Optimization Algorithm (QAOA) with the LLM's role to read financial news from RavenPack dataset. A per-asset conviction tilts the cost Hamiltonian of a cardinality-constrained QAOA portfolio optimizer. With strong bias controls, we ask whether the language adds out-of-sample value to two LLMs, Qwen2.5-14B, Sept. 2024 released during the evaluation window and an out-of-sample Llama 2, cut off Sept. 2022 on the same headlines, the Qwen model shows a mean difference relative to matched random tilts by +0.76%/quarter and Llama 2 -0.23%/quarter ($p = 0.94$) which is indistinguishable from random. The benefit or alpha is inferred only for the model released into the test period consistent with data the model is trained on or the training-horizon hindsight rather than prediction, otherwise a misleading alpha. We contribute a quantified training-horizon look-ahead insight on LLM-for-finance evaluation.

21. Quantum-Structured World Models (QSWMs) for Predictive Latent Dynamics

Hailong Jiang, Emran Hossain, Feng Yu, Jianfeng Zhu, Guilin Zhang, and Wulan Guo

Contributed talk · Sunday, 16 August 2026, 3:00 PM–3:30 PM

World models learn latent states that summarize interaction histories, evolve over time, and support prediction, simulation, or planning. Most existing world models represent these states using classical vectors, probability distributions, recurrent hidden states, or transformer activations. In this paper, we introduce Quantum-Structured World Models (QSWMs), a quantum-inspired framework for predictive world modeling with structured latent states, latent transition operators, and measurement-inspired decoding maps. We study whether mathematical structures inspired by quantum theory, such as complex-valued representations and density-matrix-like latents, provide useful inductive biases for world modeling. We establish three foundational properties: classical inclusion, predictive sufficiency, and structured compactness. We then instantiate complex-valued and density-matrix-like QSWM variants and evaluate them on elementary cellular automata against strong classical baselines. Results show promising local predictive potential for complex-valued QSWMs, while also revealing limitations in long-horizon rollout, density-matrix variants, and latent interpretability.

22. Learning to Program Multi-Chip Quantum Neural Networks via Fast Weights

Samuel Yen-Chi Chen, Chun-Hua Lin, Jiun-Cheng Jiang, Kuo-Chung Peng, Yifeng Peng, Junghoon Justin Park, Huan-Hsin Tseng, Hsin-Yi Lin, Kuan-Cheng Chen, Chen-Yu Liu, and Shinjae Yoo

Contributed talk · Sunday, 16 August 2026, 3:30 PM–4:00 PM

Quantum machine learning (QML) based on variational quantum circuits (VQCs) has shown promise, but practical performance often depends on manually selected and fixed circuit designs. This paper proposes a dynamic multi-chip QML framework for time-series prediction, where an external neural controller updates both VQC rotation parameters and the weighting coefficients of multiple quantum processors at each time step. This enables input-dependent activation of different VQC components during inference. Numerical simulations on several sequential benchmarks show that the proposed framework achieves comparable or superior performance to conventional single-chip VQC models. These results suggest a flexible route toward task-adaptive and programmable QML architectures.

23. The Next Decade of Quantum AI: A Roadmap from Computation to Discovery

Keeper L. Sharkey (ODE)

Invited talk · Sunday, 16 August 2026, 4:00 PM–4:45 PM

Abstract to be announced.

Posters and Abstracts

Posters are presented during the poster sessions on Saturday, 15 August (12:00–1:15 PM and 2:45–3:30 PM) and Sunday, 16 August (12:00–1:15 PM).

Poster Overview

#	Poster Title
1	Quantum Compositional NLP for Arabic: Grammar, Morphology, and Word Sense in Circuit Topology
2	Adapting Quantum-Inspired Word-Sense Disambiguation to Real Quantum Hardware: A Feasibility Study
3	Security Analytics for Twin-Field QKD Telemetry Using QLSTM, Transformers, and Anomaly Learning
4	What do AI practitioners need to know about “Quantum AI”?

#	Poster Title
5	EmotiQ: A Quantum Residual Specialist Framework for Text Emotion Classification
6	Benchmarking the Bures-Metric Quantum Natural Gradient for Variational Quantum Natural Language Inference
7	QC-Stark: A Multi-Task Benchmark Revealing Capability Dissociations in LLMs for Quantum Computing
8	Automatic Pregroup Supertagging for Scalable Hindi Quantum Natural Language Processing
9	QGENT: Semantic Parameter Generation for OOV Robustness in Quantum NLP
10	Beyond SAT Traces: Representation as Inductive Bias for Neural Formal Reasoning in Regular Logic
11	Conformal Risk Control for Quantum Shot Selection
12	Learning When to Go Deeper: A Cost-Aware Adaptive Controller for Quantum Architecture Search and Noise Mitigation on NISQ Devices

Poster Abstracts

P1. Quantum Compositional NLP for Arabic: Grammar, Morphology, and Word Sense in Circuit Topology

We present the first pregroup grammar-based quantum compositional NLP (QNLP) system for Arabic — a morphologically rich, freewordorder language with a formal correspondence to quantum tensor products identified independently in the computational linguistics literature. Our system converts Arabic sentences into quantum circuits whose topology mirrors grammatical structure: Subject-Verb-Object (SVO) and Verb-Subject-Object (VSO) word orders produce circuits with structurally distinct wiring, determined entirely by the pregroup grammar derivation. We conduct three controlled experiments on word order, morphological tense, and verb sense disambiguation, comparing quantum circuit methods against AraVec and AraBERT baselines. The central finding is a clean causal ablation: quantum circuits encoding only grammar topology achieve exactly 50.0% accuracy on a matched-pair word order task (zero variance across all 150 fold-seed evaluations); adding a single parameterised entangling layer raises accuracy to 64.9% (95% CI [62.8, 66.3]). This 15-percentage-point gain is entirely attributable to parameterised entanglement — the zero-variance baseline establishes a true zero point, not an estimate. We additionally introduce the first vocabulary-controlled Arabic word sense disambiguation dataset and characterise SPSA label inversion in symmetric binary quantum classification tasks, showing that ancilla qubit encoding measurably reduces inversion rates.

P2. Adapting Quantum-Inspired Word-Sense Disambiguation to Real Quantum Hardware: A Feasibility Study

Word-sense disambiguation (WSD) is one of the most challenging lexical-semantic tasks in Natural Language Processing, since a system must select the intended meaning of a word among competing candidate senses in context. This makes WSD a suitable test bed for studying whether quantum-inspired representations, in particular superposition-based semantic states, can model alternative senses and their contextual compatibility. Starting from the existing QiWSD quantum-inspired representation mechanism, this paper studies how its scoring strategy can be adapted to executable quantum circuits. We frame the work as a feasibility study focused on circuit realization and quantum-hardware execution, rather than as an attempt to advance the state of the art in WSD. The method builds sense representations from WordNet glosses and examples, encodes target contexts with BERT-based representations, and compares candidate senses through classical, quantum-only, and hybrid scoring variants. We evaluate the approach on CoarseWSD-20 dataset using two target words subsets with different ambiguity levels: apple, treated as a simple two-sense case, and seal, treated as a more challenging four-sense few-shot case. Experiments are conducted in simulation and on quantum hardware using IBM Quantum and the BSC/QIBO infrastructure. The results show that quantum and quantum-inspired scores can produce meaningful WSD decisions, while the hybrid QiWSD+ variant is the most consistent configuration in terms of accuracy.

P3. Security Analytics for Twin-Field QKD Telemetry Using QLSTM, Transformers, and Anomaly Learning

Twin-Field quantum key distribution (TF-QKD) overcomes the repeaterless rate–distance bound, yet its field security rests on phase stabilization, reference-light handling, and implementation details that surface directly in device telemetry. We cast TF-QKD security monitoring as a sequence-learning problem and compare a Transformer, an LSTM, a hybrid quantum LSTM (QLSTM), and a one-class autoencoder across clean, drifted, asymmetric-loss, and unknown-attack regimes. We find no quantum advantage: the Transformer is the strongest supervised classifier both in-domain and under distribution shift, while the autoencoder is more useful as a ranking-based anomaly scorer than as a hard-thresholded detector. The QLSTM tracks the classical recurrent baseline closely—most tightly under asymmetric loss—but never overtakes the Transformer here. Attention maps, drift traces, and t-SNE projections agree that the telemetry carries real physical structure, with phase-lock error, the phase-basis QBER, and reference power the most discriminative signatures. The operational message is simple: reliable TF-QKD monitoring should pair a closed-set sequence classifier with an independent open-set anomaly detector rather than collapse both roles into one score.

P4. What do AI practitioners need to know about “Quantum AI”?

Despite agreement that quantum computing will dramatically change the field of artificial intelligence, AI practitioners lack an actionable framework for evaluating quantum AI claims or understanding if their workthroughs would benefit from quantum computing. To address this gap, this work proposes Quantum AI Technology Readiness Levels (QAITRLs), a spin on the aerospace industry-accepted technology readiness levels (TRLs). While literature primarily focuses on hypothetical speedups, these QAITRLs enable practitioner decision-making by addressing the end-to-end AI workflows that are relevant to practitioners.

This work’s first contribution, the QAITRLs are a nine level framework for understanding the readiness of quantum techniques for AI practitioners, adapted from NASA’s TRLs. Quantum technologies are evaluated on three axes: their complexity-theoretic foundation, their physical realization and hardware fidelity, and their statistical validity and benchmarking. The first axis, complexity-theoretic foundations, focuses on whether the potential quantum advantage is based on a rigorous theoretical justification, the second, physical realization and hardware fidelity, focuses on how feasible the technology is to implement in hardware, and the final statistical validity and benchmarking, asks whether the technology has been empirically validated with reproducible, industry-standard metrics. Prior work, like the QMetric framework, provides a set of objective and reproducible metrics that go beyond basic metrics like accuracy. Approaches are assigned to QAITRLs based on their weakest axis. The work’s second and third contributions extend from the QAITRLs, providing AI practitioners with a glossary mapping quantum techniques to classical ML concepts, and a decision guide for whether to pursue or invest in quantum approaches.

At the lowest levels (1-2) of the QAITRLs are works like the Harrow, Hassidim, and Lloyd (HHL) algorithm, which propose asymptotic speedups, but abstract practical hardware and software implementation and cannot be empirically benchmarked. Then, between levels 3 and 6, there are examples of QAI research that runs on simulators or noise intermediate-scale quantum (NISQ) devices. These technologies deal with real-world noise and connectivity constraints, and are comparable to classical baselines on real-world data. The final tiers of the QAITRLs, 7 to 9, are, so far, unreachable, requiring systems that are deployed in production environments with repeatable advantages over classical baselines and clearly characterized risk bounds and maintenance requirements.

As mentioned, most QAI technologies sit around QAITRL 3-5, while no systems reach levels 7-9 in our framework. Theoretically mature algorithms, like HHL, are promising but have the weakest hardware and benchmarking foundations, while less sophisticated algorithms have been applied to NISQ hardware. When considering barriers to higher QAITRL levels, we find barren plateaus, which are similar to vanishing gradients in that the quantum algorithm’s loss landscape appears flat, is the main barrier to training many QAI algorithms. Additionally, we find that dequantization, where tasks thought to require quantum computers can be efficiently solved by quantum-inspired classical algorithms,

may limit the QAITRL that algorithms are on. Currently, the most favorable approach harness quantum-native data from domains like chemistry or materials science.

The practitioner decision guide operationalizes these findings into a structured five-step process. Practitioners begin by assessing the if the workload is classically hard, since, if the classical baselines are underexplored quantum approaches are unlikely to offer near-term benefit. Then, the data provenance is evaluated, since quantum-native data sources are most favorable for hybrid quantum-classical methods, and NISQ devices are strictly limited by the size and shape of the dataset. Next, problem structure is examined to distinguish sparse, QUBO-amenable problems from kernel-based tasks. Finally, the practitioners timeline and risk tolerance are key metrics for whether QAI technologies will be practical or even possible.

This work is directly along the lines of the “Quantum Computing for AI” track because it gives conference-goers insight into the current state of quantum approaches to artificial intelligence, helping them understand whether, and when, they will be ready for deployment and includes a domain-agnostic framework that can be applied to medical, chemistry, and cyber applications.

P5. EmotiQ: A Quantum Residual Specialist Framework for Text Emotion Classification

Emotion classification from text is challenging because many affective states are distinguished by subtle linguistic signals and semantically overlapping boundaries. We investigate whether compact quantum modules can serve as targeted residual specialists that assist traditional classifiers on difficult emotion decisions. We introduce EmotiQ, a residual specialist framework for text emotion classification using a frozen transformer classifier as the primary decision maker, identifies challenging emotion boundaries from validation behavior, and trains boundary specialists, including quantum encodings, data re-uploading specialists, prototype and fidelity specialists, and kernel-based specialists. Specialists are permitted to intervene only when they satisfy validation-based admission criteria. We evaluate the framework across six text emotion classification datasets and establish traditional model baselines. Experimental results show that the admitted EmotiQ policy improves mean macro-F1 metrics over the baseline, and achieves positive gains on five of six datasets. Among the quantum components, ZZ interaction feature maps are the most consistently effective, yielding a mean macro-F1 improvement of +0.0030 and positive gains on four datasets. These findings suggest that current quantum NLP methods can be deployed as residual specialists that provide localized improvements under strong classical supervision.

P6. Benchmarking the Bures-Metric Quantum Natural Gradient for Variational Quantum Natural Language Inference

Training variational quantum models requires choosing between parameter-shift gradients, which are exact but cost $O(P)$ forward evaluations, and simultaneous

perturbation stochastic approximation (SPSA), which uses only two samples but produces high-variance estimates that can degrade optimisation on small supervised tasks. Whether the cheap gradient is usable depends on the variance that results from different choices of the SPSA perturbation scale, learning rate, and gain-decay schedule. We varied those quantities across a broad grid on a 6-qubit, 60-parameter QNLI classifier and compared the best configurations to parameter-shift AdamW and BuresQNG. AdamW-style SPSA with $c_0=0.01$, $\eta=0.10$, $\gamma=0.10$ reached $55\% \pm 11\%$ test accuracy, improving over the default configuration ($49\% \pm 6\%$) but remaining 16–19 percentage points below the parameter-shift baselines because the two-sample SPSA gradient estimate has too much variance for reliable optimisation of 60 parameters in 40 epochs. Classical-gain SPSA and Bures-preconditioned SPSA performed worse, at 51% and 46% respectively. Bures-preconditioning a noisy two-sample SPSA gradient amplifies perturbation noise.

P7. QC-Stark: A Multi-Task Benchmark Revealing Capability Dissociations in LLMs for Quantum Computing

We introduce QC-Stark, a benchmark evaluating large language models (LLMs) on 11 quantum computing (QC) tasks spanning circuit construction, debugging, compilation, error correction, and simulation. Across 2,200 evaluations (8 models $\times$ 11 tasks $\times$ 5 difficulty levels $\times$ 5 seeds), we find that overall rankings mask substantial per-task variation: Spearman ρ between overall and per-task rankings drops below 0.5 for 3 of 11 tasks. We identify a debugging paradox: reasoning-optimized models score significantly lower than the weakest overall model on circuit debugging (o4-mini: 0.10 vs. LLaMA-70B: 0.27). A 2-parameter Item Response Theory (IRT) model validates measurement quality (reliability = 0.98, 81% items with discrimination $a > 1.0$), and prompt sensitivity analysis confirms ranking robustness ($\rho = 0.96$) across prompt conditions. All tasks are auto-verified via execution; code and data are publicly available.

P8. Automatic Pregroup Supertagging for Scalable Hindi Quantum Natural Language Processing

Quantum Natural Language Processing (QNLP) employs pregroup grammars to translate grammatical structure into diagrammatic representations and quantum circuits. Recent Hindi QNLP work has shown that Hindi-specific pregroup grammars can support grammar-sensitive compositional models, but grammatical type assignment is still largely manual, limiting scalability. This paper formulates automatic Hindi pregroup supertagging as a token-level classification task. Using a manually annotated corpus of 380 Hindi sentences, we evaluate lexical, contextual, prompting-based, lexical-repair, and suffix/morphology-aware methods. Our experimental evaluation demonstrates that simple lexical and contextual models are strong in this low-resource setting: contextual backoff achieves the best completed accuracy of 64.56%, while raw Qwen2.5 prompting reaches only 11.65%. Lexical repair raises LLM-assisted prediction to 64.08%, demonstrating the value of constraining generative outputs with symbolic grammar knowledge. Diagnostic analysis further reveals that seen and unambiguous tokens are much easier than unseen tokens,

and suffix/morphology features improve karaka-token accuracy but not overall performance. These results establish automatic Hindi pregroup assignment as a practical step toward scalable multilingual QNLP pipelines.

P9. QGENT: Semantic Parameter Generation for OOV Robustness in Quantum NLP

Quantum Natural Language Processing (QNLP) represents sentence meaning through grammar-derived quantum circuits, but many current pipelines remain tied to the vocabulary observed during training. When a test sentence contains an out-of-vocabulary (OOV) word, the circuit may still be constructible, but the word-specific gate parameters are unavailable. We present QGENT, a semantic parameter-generation framework for improving OOV robustness in QNLP. QGENT separates circuit topology from parameter assignment: a lambeq-based QNLP model supplies the grammar-derived circuit structure and learned parameters for observed words, while GPT-based semantic embeddings and a neural mapping layer estimate gate parameters for unseen words. We evaluate QGENT in low-data sentiment-classification settings on SST-2 and Food-IT, comparing against grammar-only QNLP and FastText-based OOV handling. The results are preliminary, but they suggest that semantic embeddings can serve as useful parameter priors for QNLP-style models under vocabulary shift. We do not claim quantum advantage or superiority over mature classical NLP systems. The contribution is a modular method for connecting pre-trained semantic representations with quantum-compatible language circuits, together with an explicit discussion of the limitations of the current SpiderReader/SpiderAnsatz prototype.

P10. Beyond SAT Traces: Representation as Inductive Bias for Neural Formal Reasoning in Regular Logic

Coecke-style process-theoretic semantics makes explicit a central QNLP principle: compositional structure is represented by wiring typed processes together, rather than by treating surface-string order as the primary object [2]. This poster brings that principle into a classical formal-reasoning setting. In bounded regular logic, variables behave as persistent wires, predicates as typed boxes, and repeated variable occurrences as shared connections. Although a left-to-right tokenization can encode the same formulas, it hides identity, sharing, and composition inside names, positions, and traversal order. We therefore study representation itself as a possible inductive bias for learning formal identity.

This work extends recent evidence that transformers can learn bounded propositional SAT traces while retaining only limited length generalization [1]. Moving beyond SAT introduces a qualitatively different burden: the model must preserve scoped objects across predicate roles. This burden aligns with Wang and Sun’s ICLR 2026 analysis of the Reversal Curse as a transformer binding problem [4]. We use bounded existential-conjunctive regular logic, rather than unrestricted first-order logic, because it includes predicates, variable reuse,

and relational composition while preserving exact symbolic checking. Its graphical reading is also mathematically natural, matching diagrammatic treatments of first-order logic [3].

We introduce a verifier-backed minimal-pair benchmark in which paired formulas differ only in which variable fills each predicate role, for example $P(v_0, v_1)$ versus $P(v_1, v_0)$. A small transformer trained from scratch maps formulas to canonical symbolic descriptions. We separately measure parseability, exact role identity, preference for the gold object among symbolically valid alternatives, and linear-probe recovery of the intended binding relation from hidden states. This decomposition separates grammar learning from identity preservation and from failures introduced by the output interface.

The diagnostic exposes a gap between formal appearance and formal identity. A compact-syntax baseline reaches 1.0 parseability but only 0.633 exact match on one held-out reversal panel, with complete collapse in a stricter high-index control. Candidate ranking among valid formulas closes only part of the gap, reaching 0.733 top-1 accuracy in the clean baseline. However, a linear probe on frozen final-layer hidden states recovers variable-to-slot assignments with 0.900 held-out accuracy. Thus, left-to-right generation can fail to externalize role information that remains substantially recoverable inside the model.

As a targeted pre-diagrammatic intervention, we add source-conditioned binding supervision: each output argument is attached to an explicit source-side candidate set rather than learned only as a repeated variable token. Across three seeds on the canonical held-out reversal panel, exact generation rises to 0.767/1.000/0.900, binding-only candidate top-1 accuracy rises to 0.800/1.000/0.967, and factorized hidden-state slot accuracy rises to 0.900/1.000/0.983. Thus we establish how to show that a role-sensitive failure mode of flat syntax can be exposed, localized, and substantially improved under verifier-backed controls.

This baseline now supports our main representation study. We have begun constructing linearized string-diagram-inspired proof-state/action encodings and matched symbolic transition rows, comparing compact syntax, source-conditioned binding, factorized frontier/action states, and structured state serializations under the same model family, data source, training budget, and verifier. Our working hypothesis is that syntactic bureaucracy can be a double-edged inductive bias: canonical surface conventions may improve in-distribution fitting by giving the transformer stable textual regularities, while increasing out-of-distribution fragility when those conventions entangle proof identity with arbitrary token order. These transition-level experiments explicit purpose is to test whether exposing proof state as wiring and action over wiring changes what a transformer learns compared with ordinary syntax.

[1] L. Pan, V. Ganesh, J. Abernethy, C. Esposito, and W. Lee, “Can Transformers Reason Logically? A Study in SAT Solving,” ICML, 2025. [2] B. Coecke, M. Sadrzadeh, and S. Clark, “Mathematical Foundations for a Compositional Distributional Model of Meaning,” Linguistic Analysis, 2010. [3] F. Bonchi, A. Di Giorgio, N. Haydon, and P. Sobociński,

“Diagrammatic Algebra of First Order Logic,” LICS, 2024. [4] B. Wang and H. Sun, “Is the Reversal Curse a Binding Problem? Uncovering Limitations of Transformers from a Basic Generalization Failure,” International Conference on Learning Representations (ICLR) Poster, 2026.

P11. Conformal Risk Control for Quantum Shot Selection

Variational quantum algorithms require choosing a measurement shot budget M —too few shots yield unreliable estimates, too many waste expensive quantum resources. We apply conformal risk control (CRC) to this problem, providing distribution-free, finite-sample guarantees on variational quantum eigensolver (VQE) energy estimation error. Through a fair comparison of five risk control methods—CRC, Gaussian t-intervals, bootstrap, grid search with Bonferroni correction, and no risk control—on H_2 VQE at four bond lengths (10 seeds each) and real IBM Quantum hardware (156-qubit Heron processor), we identify the regime where CRC’s distribution-free guarantee provides genuine value: at small quantum budgets ($B \leq 1000$ circuit executions), Gaussian and bootstrap methods select aggressive shot budgets but fail to maintain the target risk level 40% of the time. CRC correctly identifies that the budget is insufficient to certify any reduction, preventing invalid deployments. At larger budgets, CRC matches parametric methods while providing formal backing. On real hardware, CRC correctly issues a “do not deploy” verdict when hardware noise dominates shot noise.

P12. Learning When to Go Deeper: A Cost-Aware Adaptive Controller for Quantum Architecture Search and Noise Mitigation on NISQ Devices

The emergence of classical chaos from quantum dynamics remains a central problem in theoretical physics, particularly as one intends to simulate such complex systems on soon-to-be-available noisy quantum computers. Here, we consider the applicability of variational quantum simulation methods to simulate the dynamics of a quantum system that exhibits a transition from regular to chaotic motion: the quantum kicked top. By relating classical and quantum measures of chaos, we identify the requirements for variational quantum simulation methods to capture such chaos. Specifically, we find that existing approaches based on fixed-depth ansatzes are unable to capture the complexity of the chaotic regime. For this problem, we propose and implement an adaptive quantum circuit design method that increases the depth of the quantum circuit only when a specific signal indicates that the ansatz is no longer expressive enough to describe the quantum state. Building on this, we propose a noise-aware adaptive controller that adds a gatekeeper rule, ensuring depth is only increased when the expected accuracy improvement is worth the added noise and measurement overhead. Our findings illustrate that these adaptive algorithms are able to detect chaos accurately and provisionally assign computational effort. We show that in addition to being a numerical ingredient, the variational circuit depth can be interpreted as an operative probe of physical complexity. Methods allowing for adaptive depth allocation are therefore crucial when simulating non-integrable quantum dynamics. We further identify this adaptive controller as a concrete,

RL-ready testbed for AI-driven quantum architecture search, framing depth allocation as a sequential decision problem that connects directly to reinforcement learning for quantum circuit optimization.