



MASTER'S THESIS

MSc – Artificial Intelligence for Business Transformation

Quantum Optimisation for the Tactical Aircraft Conflict Resolution Problem

A Reproducible QUBO and QAOA Benchmark against a Structured
Classical Baseline

Author Louis-Gabin OZIL

Supervisor Professor David Rey



Academic year 2025–2026

September 2026

Abstract

This thesis evaluates whether quantum optimisation can approach a structured classical baseline for a tactical Aircraft Conflict Resolution Problem. Aircraft manoeuvres are encoded as a Quadratic Unconstrained Binary Optimisation model and solved using Gurobi, exhaustive enumeration, simulated annealing and the Quantum Approximate Optimization Algorithm. The comparison combines ideal simulation, controlled noise and execution on IBM quantum hardware. Exact QUBO solutions match discrete Gurobi on every shared instance, while the classical methods remain faster and more reliable. The hardware experiments nevertheless confirm that the complete aircraft-to-quantum pipeline is executable. Their main finding concerns circuit depth: a well-optimised additional layer improves the probability of sampling an optimal solution, whereas a near-null layer of similar physical cost mainly adds exposure to noise. The study finds no practical quantum advantage for the tested instances, but provides a reproducible benchmark and a managerial distinction between technical feasibility, competitive performance and operational value under hardware constraints.

Keywords: aircraft conflict resolution; quantum optimisation; QUBO; QAOA; quantum-classical benchmarking; IBM Quantum.

Preface and Acknowledgements

This thesis was completed as an individual research project within the MSc in Artificial Intelligence for Business Transformation at SKEMA Business School. It grew from an interest in the practical value of quantum optimisation and from the question of how an emerging technology should be assessed when access, computing resources and experimental time remain limited.

The original ambition was progressively adapted to the technical conditions of the project, particularly the accessibility constraints of IBM Quantum hardware and the workload of conducting the research alone. These constraints led to a deliberately bounded experiment focused on transparency, reproducibility and a fair comparison with classical optimisation. They shaped the scope of the thesis without changing its central purpose.

I would like to thank my supervisor, Professor David Rey, for his guidance, critical feedback and support in keeping the project academically rigorous and appropriately scoped. I also thank SKEMA Business School for the opportunity to undertake this interdisciplinary work. I remain solely responsible for the analysis, interpretations and any remaining errors.

Contents

Abstract	i
Preface and Acknowledgements	ii
Executive Summary	v
List of Tables	viii
List of Figures	ix
List of Acronyms	x
Introduction	1
1 Literature Review	4
1.1 Aircraft Conflict Resolution Problem: problem framing and classical foundations	4
1.2 Quantum optimisation for combinatorial problems: concepts, formulations, and method- logical cautions	8
1.2.1 The QUBO framework as a bridge to quantum hardware	8
1.2.2 Quantum optimisation paradigms: a functional overview	10
1.2.3 Benchmarking cautions in quantum-classical comparisons	11
1.3 Quantum approaches applied to the ACRP	13
1.3.1 Annealing-based approaches	13
1.3.2 Gate-based variational approaches	14
1.3.3 Adjacent domain: UAV deconfliction	15
1.3.4 Critical synthesis and research gap	15
2 Methods, Data Sources and Analysis	18
2.1 Research design and positioning	18
2.2 Problem formalisation and QUBO encoding	19
2.2.1 Tactical 2D ACRP	19
2.2.2 Discretisation into candidate manoeuvres	19
2.2.3 QUBO formulation	20
2.2.4 Penalty calibration	20
2.3 Data and instance generation	21
2.4 Solvers compared	22
2.4.1 QAOA depth as a configuration variable	23
2.4.2 Noise sensitivity of QAOA	23
2.5 Experimental protocol	26
2.6 Evaluation metrics and scaling behaviour	27
2.7 Statistical analysis	28
2.7.1 Replication and interval estimation	28

2.7.2	Hypothesis testing and effect sizes	29
2.8	Validity, limitations and ethics	30
2.8.1	Practical scope and resource constraints	30
3	Results	32
3.1	Experimental campaign	32
3.2	Validation of the encoding	32
3.3	Solver comparison on shared instances	33
3.4	QAOA behaviour: depth and noise	34
3.5	Scaling behaviour	35
3.6	Execution on real quantum hardware (IBM Quantum)	36
3.7	Synthesis: answering the research question	39
	Conclusion	41
	References	44
	Appendix A: Comparative Overview of Quantum ACRP Contributions	47
	Appendix B: Formal Derivations, Notation and Algorithms	48

Executive Summary

Purpose

Aircraft conflict resolution involves selecting safe trajectory changes when several aircraft risk coming too close to one another. Because one manoeuvre may affect several conflicts, the decisions must be considered together. This thesis asks whether a quantum approach can match a strong classical method on the same controlled instances, and what resources are required to obtain the result. It distinguishes a circuit that can run from a method that is genuinely competitive or operationally valuable.

Approach

The study uses a simplified two-dimensional tactical benchmark. Each aircraft selects one heading change from a small discrete list while speed remains fixed at its nominal value. The problem is translated into a QUBO, a binary representation compatible with classical and quantum solvers. A discrete Gurobi model provides the exact classical reference, exhaustive enumeration checks the smallest cases, simulated annealing works on the same QUBO, and QAOA is tested in ideal simulation, under simulated noise and on the IBM processor `ibm_kingston`. Agreement between the discrete Gurobi model and exact QUBO enumeration confirms that these methods solve the same simplified benchmark.

Main findings

The classical methods remain the clear reference. Gurobi solves the tested range in less than 40 milliseconds, while simulated annealing returns valid solutions more frequently than QAOA. The study therefore finds no practical quantum advantage in speed, reliability or solution quality for these instances.

The quantum pipeline nevertheless works from end to end. Aircraft-conflict data are converted into a QUBO and then a quantum circuit, which is checked locally before hardware execution. The processor returns optimal answers more often than uniform sampling over the 512 unconstrained bitstrings. It remains below a simple classical comparator that samples uni-

formly among the 27 one-hot assignments. This demonstrates technical feasibility, but not competitiveness.

The strongest result concerns circuit depth. QAOA circuits contain repeated layers. A useful layer can improve the answer in an ideal simulator, but on real hardware it also adds gates and exposure to noise. A three-layer circuit selected independently from 96 optimiser starts produced an optimal answer in 15.11% of measurements. A circuit of almost the same length, whose third layer contributed almost nothing, reached 10.77%. Repeated two-layer control circuits helped verify that the difference was not caused only by a major hardware change during the job. Longer circuits failed to preserve their stronger ideal results. More depth is therefore valuable only when the added algorithmic work is sufficient to compensate for its physical cost.

This also shows why the number of available qubits is not enough to judge whether a machine is ready for a problem. The circuit required to use those qubits must remain reliable.

Constraints

The research was conducted as an individual Master’s project in a business-school programme, without a dedicated quantum team or reserved hardware. IBM Quantum was accessed through the Open Plan, which limited processing time and offered no control over queues or calibration. Local simulation was restricted by computer memory, since each additional qubit rapidly increases the information that must be stored.

These conditions led to a principal experiment with three aircraft represented by nine qubits. Parameters were optimised locally, and each completed circuit was measured 4,096 times on IBM hardware. The benchmark remains mainly based on discrete heading choices and does not represent continuous operational air traffic. Total user time was not recorded consistently for every run, and the complete table needed for one planned statistical ranking was unavailable. Missing results were not reconstructed.

The project used several safeguards: exact checks against Gurobi, two independent quantum implementations, fixed random seeds, frozen submission plans, retained job identifiers and repeated control circuits. An initially under-optimised three-layer circuit was retained as a near-null control rather than discarded, turning an imperfect result into a useful comparison.

Managerial implications and recommendations

For managers, technical feasibility, competitive performance and operational value are different levels of maturity. A quantum prototype can support learning and skills development without justifying deployment. For this benchmark, structured classical optimisation remains the appropriate computational choice, while quantum work should be treated as exploratory research.

An organisation considering similar technology should ask whether the comparison uses a strong classical alternative, whether valid and high-quality answers are distinguished, whether the full cost includes preparation and waiting, and whether the result can be reproduced under different hardware conditions.

Future experiments should first improve completeness and reproducibility. They should record every time component, repeat identical circuits under different hardware conditions, compare error-reduction techniques with unchanged controls and redesign the encoding to require fewer gates. Larger problems should be attempted only after these improvements and should remain compared with strong classical methods.

The thesis therefore replaces a general promise with a measurable conclusion. The aircraft-to-quantum pipeline can run on real hardware, but present performance is not competitive. Progress requires each increase in quantum complexity to deliver enough useful work to survive its own physical cost.

List of Tables

2.1	Continuous MINLP cross-check on CP_3	21
2.2	Discrete Gurobi–QUBO validation	22
2.3	Ideal and noisy QAOA on CP_3	24
2.4	Ideal and noisy QAOA on CP_4	25
2.5	Evaluation metrics used across classical, simulated and hardware results.	29
3.1	Exact objective values in the twelve-instance, $M = 3$ benchmark. The objective is the number of manoeuvring aircraft.	33
3.2	Functional comparison of the solver set on the shared benchmark.	34
3.3	Ideal statevector $P(\text{opt})$ across depth for $M = 3$, using ten seeds and forty-eight optimiser restarts. A dash denotes a configuration not retained for reporting.	34
3.4	Width and indicative routed-circuit cost for the $M = 3$ circle family. Two-qubit counts are $p = 1$ estimates beyond the measured CP_3 point. F is a calibration-based retained-fidelity projection and is not an observed algorithmic fidelity.	36
3.5	Initial CP_3 hardware check	37
3.6	First CP_3 hardware depth series	37
3.7	Interleaved CP_3 hardware confirmation	38
3.8	Hardware $P(\text{opt})$ contrasts	39
A.1	Comparative overview of quantum ACRP studies	47
B.1	Notation used in the QUBO and QAOA formulations.	50

List of Figures

2.1	Exact penalty feasibility on CP_3	20
2.2	Preliminary QAOA depth calibration on CP_3 using one seed and 120 optimiser restarts per depth. $P(\text{one-hot})$ measures only the assignment constraint, not full conflict-free feasibility.	23
2.3	Probability of sampling the optimum against QAOA depth on CP_3 , ideal versus depolarising noise. The ideal curve rises with depth while the noisy curve stalls, so the gap widens.	25
2.4	Feasibility rate against QAOA depth on CP_3 , ideal versus depolarising noise. The share of constraint-satisfying samples grows with depth only in the ideal case.	26
2.5	Probability of sampling the optimum against QAOA depth on CP_4 (twelve qubits), ideal versus depolarising noise. The noisy curve is below the ideal at every depth, and the ideal curve is no longer monotone because the deeper circuits are under-optimised.	27
2.6	Feasibility rate against QAOA depth on CP_4 (twelve qubits), ideal versus depolarising noise. Noise lowers the constraint-satisfying share at every depth. . . .	28
2.7	Scaling on the Circle Problem: QAOA success probability falls and simulation time rises steeply with qubit count.	30

List of Acronyms

ACRP	Aircraft Conflict Resolution Problem
ATM	Air Traffic Management
CD&R	Conflict Detection and Resolution
CI	Confidence Interval
CP_{<i>n</i>}	Circle Problem with <i>n</i> aircraft
CVaR	Conditional Value-at-Risk
MILP	Mixed-Integer Linear Program
MINLP	Mixed-Integer Non-Linear Program
NISQ	Noisy Intermediate-Scale Quantum
P(opt)	Probability of sampling the optimal solution
QAOA	Quantum Approximate Optimization Algorithm
QAOAnsatz	Quantum Alternating Operator Ansatz
QPU	Quantum Processing Unit
QUBO	Quadratic Unconstrained Binary Optimization
UAV	Unmanned Aerial Vehicle
VQE	Variational Quantum Eigensolver

Introduction

European air traffic is expected to continue growing over the coming decades. EUROCONTROL’s most likely long-term scenario forecasts 15.4 million flights in the European Civil Aviation Conference area in 2050, 39 per cent more than in 2019 (EUROCONTROL, 2024). This prospective growth does not by itself determine the complexity of air traffic management, but it reinforces the value of methods that can support safe and efficient trajectory decisions. One such decision problem is the Aircraft Conflict Resolution Problem (ACRP), in which aircraft trajectories must remain safely separated while the operational disruption caused by avoidance manoeuvres is limited. When several aircraft are involved simultaneously, resolving conflicts pair by pair may create new interactions elsewhere. The problem must therefore be treated collectively, as an optimisation problem over compatible manoeuvre choices (Kuchar & Yang, 2000; Pelegrín & Cerulli, 2023).

Classical ACRP research has developed mixed-integer linear, non-linear and convexified formulations for heading control, speed regulation and combined manoeuvres. This literature establishes that conflict geometry can be represented mathematically and solved on controlled instances, while also showing that computational behaviour depends on the formulation, manoeuvre space and level of physical detail. A restricted discrete model and a richer continuous formulation do not represent the same computational challenge. Any assessment of an alternative method must therefore compare like with like and avoid treating the difficulty of one formulation as a universal property of the ACRP.

Quantum optimisation offers a different representation of combinatorial decisions. A Quadratic Unconstrained Binary Optimisation (QUBO) model expresses choices and incompatibilities through a quadratic objective over binary variables (Lewis & Glover, 2017). The model can then be mapped to an Ising Hamiltonian and processed by the Quantum Approximate Optimization Algorithm (QAOA), which alternates parameterised cost and mixing operations in a hybrid quantum-classical procedure (Farhi et al., 2014). In principle, additional QAOA layers increase the expressive capacity of the circuit. On present noisy hardware, however, they also add gates and execution time. Noise therefore limits the circuit depth that can be used reliably in the noisy intermediate-scale quantum regime (Preskill, 2018). The relevant practical issue is not simply whether a deeper circuit is more expressive, but whether its algorithmic contribution survives its physical execution cost.

Previous quantum studies have shown that aircraft deconfliction can be translated into quantum-compatible formulations, from strategic delay allocation on annealing hardware to small tactical manoeuvre-selection circuits (Stollenwerk et al., 2020; Pecyna et al., 2024). Yet these studies do not provide a like-for-like answer to the structured classical ACRP literature. Their problem definitions, scales, constraint treatments, tuning procedures and evaluation metrics differ. The literature review consequently identifies a mismatch between technical demonstrations of quantum executability and the evidence required for a reproducible comparison of solution quality and computational cost.

This thesis addresses that mismatch under deliberately bounded conditions. It does not assume that a quantum formulation will be superior, nor does it attempt to reproduce continuous operational air traffic at scale. Instead, it constructs a simplified two-dimensional tactical ACRP benchmark with discrete heading manoeuvres and nominal speed fixed throughout. Structured Circle Problem instances support exact validation and controlled scaling, and randomised instances reduce dependence on one highly symmetric geometry. Accepting this limited scale makes it possible to audit the encoding, compare methods on common inputs and submit validated circuits to accessible quantum hardware.

The research question is therefore:

On a controlled two-dimensional tactical ACRP benchmark, and under transparent benchmarking conditions, to what extent can a quantum or quantum-inspired formulation match or approach a structured classical baseline in solution quality, and at what measurable computational cost?

The wording of this question separates solution quality from efficiency. A circuit may assign non-trivial probability to good solutions without being competitive in runtime, robustness or scale. Conversely, a small hardware experiment may remain informative if it isolates a mechanism that affects current devices. The analysis therefore reports feasibility before objective quality, uses the unconditional probability of sampling an optimum and records the costs that can be verified, including classical runtime, circuit depth, two-qubit gates, shots and QPU usage. Where complete end-to-end user timing was not retained, it is identified as unavailable rather than reconstructed.

Four computational roles are distinguished. A structured Gurobi formulation provides the exact classical reference. Exhaustive QUBO enumeration verifies encoding consistency on tractable instances. Simulated annealing acts as a classical twin operating on the same QUBO, and QAOA is evaluated through ideal statevector simulation, controlled noise models and IBM Quantum hardware. Fixed seeds, independent NumPy and Qiskit implementations, exact cross-checks, frozen hardware plans, retained job identifiers and interleaved control circuits are used to make the comparison auditable. This design allows the thesis to examine technical executability, comparative performance and the interaction between algorithmic depth and physical

degradation as separate questions.

The experimental scope also reflects the context in which the research was conducted. This is an individual Master’s thesis within a business-school programme, rather than a quantum-hardware project undertaken by a dedicated technical team. IBM Quantum access was provided through a limited Open Plan allocation, without a reserved processor or control over calibration schedules, and exact statevector simulation was constrained by local memory. These limitations explain the use of a nine-qubit principal hardware instance, local variational optimisation and fixed-circuit submission. They bound external validity, but the validation and traceability measures above preserve the evidential value of the measurements within that scope.

The contribution is accordingly methodological, empirical and managerial. Methodologically, the thesis connects a simplified tactical ACRP, a validated QUBO, structured and heuristic classical references, ideal and noisy QAOA, and controlled hardware execution within one documented pipeline. Empirically, it measures solution quality and physical circuit behaviour under a common protocol. From a managerial perspective, it illustrates how an emerging technology can be assessed without confusing hardware availability with operational readiness, or technical execution with business value. Reproducibility, access constraints, comparison quality and the cost of obtaining a result are treated as part of the adoption question rather than as secondary implementation details.

The remainder of the thesis is organised as follows. Chapter 1 reviews the classical ACRP literature, introduces the quantum optimisation framework and develops the research gap. Chapter 2 specifies the benchmark, QUBO formulation, solver roles, experimental protocol, metrics and validity conditions. Chapter 3 reports encoding validation, solver comparisons, depth and noise behaviour, scaling limits, statistical tests and the IBM Quantum campaigns. The conclusion then answers the research question, discusses the implications and limitations of the findings and identifies conditions for stronger future evidence.

Chapter 1

Literature Review

1.1 Aircraft Conflict Resolution Problem: problem framing and classical foundations

The starting point of this literature review is not quantum computing itself, but the decision problem on which quantum formulations are later built. The problem of automatically detecting and resolving aircraft conflicts is not a recent concern: Kuchar and Yang (2000) survey approximately 68 conflict detection and resolution (CD&R) modelling methods, building on an earlier conference version of the same work presented in 1997, and propose a taxonomy that remains a reference point for the field. Their framework categorises methods according to the dimensionality of state information (vertical, horizontal, or three-dimensional), the mode of dynamic state propagation (nominal, worst-case, or probabilistic), the conflict resolution approach (prescribed, optimized, force-field, or manual), the manoeuvring dimensions considered (speed change, lateral, vertical, or combined manoeuvres), and the way multiple simultaneous conflicts are managed (pairwise or global). This taxonomy is analytically useful for the present review because it shows that the structuring of the Aircraft Conflict Resolution Problem (ACRP) around control variables, geometric assumptions, and conflict-management strategy predates by more than two decades the emergence of the mathematical-programming and quantum approaches discussed later in this chapter, and because several of its categorical distinctions, particularly the separation between prescribed and optimized resolution and between pairwise and global conflict management, continue to organise the way modern optimisation-based formulations of the ACRP are structured and compared.

Across the corpus, the ACRP is treated as a tactical decision problem: a set of aircraft trajectories must be adjusted so that every pair of aircraft maintains a prescribed safety separation, while the resulting deviations from planned operations remain as limited as possible. Pelegrín and Cerulli (2023) frame the problem through rectilinear trajectories crossing a monitored air sector and note that controllers generally need to issue separation instructions around 20 to 30

minutes before a potential conflict, placing the problem in the tactical layer of air traffic management rather than in long-term strategic planning. Within this tactical setting, the ACRP is generally formulated as a mathematical optimisation problem, in which feasible solutions are shaped by aircraft separation constraints while the objective function penalises deviations from nominal trajectories or operations. This framing is what allows the ACRP to be treated later in this review as a candidate for reformulation into quadratic unconstrained binary optimisation (QUBO) and, ultimately, quantum-compatible models.

The corresponding control variables also differ significantly across formulations. Some models focus on heading changes, others on speed regulation, and others combine both; some remain in a two-dimensional horizontal setting, while others extend the formulation to three-dimensional airspace. Cai and Zhang (2018), for example, propose a three-dimensional mixed-integer nonlinear program in which aircraft can adjust heading and velocity simultaneously, with conflict avoidance expressed through nonlinear spatial separation constraints. Within the present corpus, the ACRP therefore appears less as a single canonical model than as a family of related formulations whose structure depends on the manoeuvre set, dimensionality, and level of geometric detail.

The earliest exact mathematical-programming formulation in the corpus is the mixed-integer linear program proposed by Pallottino et al. (2002), which develops separate velocity-change and heading-change formulations on a single horizontal layer. The heading-change variant considered here uses single, instantaneous heading changes, with all aircraft sharing the same normalised speed. The number of binary variables and linear constraints grows on the order of $O(n^2)$, with one constraint block for each aircraft pair, and the resulting model is solved directly with CPLEX. Reported solution times remain limited to small-scale test cases: the symmetric five-aircraft Circle Problem is solved in 0.34 seconds, the ten-aircraft case in 5.91 seconds, and the eleven-aircraft case in 10.4 seconds. The contribution is therefore foundational, but its empirical tractability is documented only on small structured instances, and this heading-change variant is restricted to heading control, so its computational performance is not directly comparable to models that allow richer manoeuvre spaces. In the taxonomy of Kuchar and Yang (2000), this formulation falls squarely within the “optimized” and “pairwise” categories, illustrating how early exact-optimisation work concretely instantiated one specific corner of the broader design space they had already mapped out.

The move toward speed-based control in Cafieri and Rey (2017) changes both the formulation and the computational picture. Their model is expressed as a non-convex mixed-integer nonlinear program and solved with COUENNE under a 300-second time limit. Unlike deviation-minimization formulations, their objective is primarily to maximise the size of the conflict-free aircraft set under speed-only control. This aircraft-level objective is distinct from maximising the number of resolved pairwise conflicts. On ten-aircraft Random Circle Problem instances, the direct formulation often finds feasible solutions but regularly reaches the time

limit without proving optimality; on twenty-aircraft instances, the direct model fails to produce any feasible solution within the allotted time for several cases, whereas a pre-processed version finds feasible solutions on all tested instances. The pre-processing routine identifies potential conflict pairs through a bound-constrained concave program solved for each aircraft pair in $O(n^2)$ time before the main optimisation stage; its runtime remains small compared with the global mixed-integer nonlinear programming budget, yet it substantially changes solver behaviour. What matters analytically here is that feasibility itself, and not only solution quality, turns out to depend on formulation and pre-processing choices rather than on the manoeuvre type alone.

Rey and Hijazi (2017) then move beyond single-control formulations by combining heading and speed in a unified framework. In the heading-only setting, Pallottino et al. (2002) derive separation constraints in a way that leads to linear inequalities in the heading-deviation variables, which is what makes the full multi-aircraft problem expressible as a mixed-integer linear program. In the speed-only setting, Cafieri and Rey (2017) show that conflict avoidance is linked to the non-positivity of the discriminant of a second-order separation polynomial, yielding nonlinear and non-convex constraints in the speed-ratio variables. When heading and speed are combined, the natural geometric representation becomes the relative velocity plane: Rey and Hijazi (2017) show that the associated separation condition defines a non-convex feasible region, and that after accounting for diverging trajectories, the conflict-free relative velocities can be represented through a disjunction of convex regions. The geometric structure is therefore not a secondary modelling detail: it determines whether separation can be encoded linearly, nonlinearly, or only through non-convex disjunctive constructions. This is not merely an alternative notation; it supports convex-envelope approximations and leads to mixed-integer quadratic and quadratically constrained relaxations that can be handled by commercial solvers, reducing computational time by up to two orders of magnitude on standard instances.

Dias and Rey (2024) build on this classical avoidance stage by adding a trajectory-recovery layer, showing that the geometry of the problem does not stop at conflict elimination itself but may also include a structured return to the nominal route. The trajectory-recovery paper distinguishes conflict avoidance from trajectory recovery, reserving conflict avoidance for approaches that resolve the separation issue without explicitly modelling the return to the nominal trajectory. The distinction is not only semantic: it shows that the literature does not always optimize the same object. Some studies stop once safe avoidance manoeuvres have been identified; others add a recovery phase, which broadens the decision problem and changes its computational structure. Their two-stage iterative algorithm first resolves initial conflicts through speed and heading control, then identifies each aircraft’s optimal recovery time; numerical results show that this approach solves instances with up to 30 aircraft in under 10 minutes, a scale comparable to the largest instances handled by the pre-processed speed-control formulation of Cafieri and Rey (2017).

The comparison across these formulations also highlights a recurring simplification. Many mathematical programming formulations rely on stylized kinematics such as linear trajectories, uniform motion, and mostly two-dimensional conflict geometry, even when the real operational environment is more complex. The references considered here therefore present the ACRP as a tractable abstraction of tactical conflict management, not as a complete operational replica of real-world air traffic management, a claim that should be handled carefully given the limited size of the current corpus, especially when discussing scalability or practical applicability later in the review. This gap between abstraction and operational reality is not a new observation: most exact mathematical-programming formulations in this corpus rely on the simplest, nominal state-propagation assumption, deliberately setting aside the worst-case and probabilistic modelling of trajectory uncertainty.

The benchmark-generator paper by Pelegrín and Cerulli (2023) defines tactical deconfliction from the standpoint of airspace supervision: a controller observes a sector, receives nominal positions and velocity vectors, and must produce new conflict-free trajectories as output. Their generator supports predefined scenarios such as circle, rhomboidal, and grid configurations, as well as three-dimensional extensions and pseudo-random instances with tunable congestion levels, offering a standardized testing framework that is useful for comparing the classical formulations discussed above with one another.

Across these heterogeneous papers, then, a common analytical core emerges, namely short-term trajectory regulation under safety constraints, even as changing the manoeuvre space or the geometric assumptions changes the problem itself rather than merely the solution technique. The point is clearest when solver performance is compared: exact formulations solved mainly with CPLEX for linear and quadratic models, and with COUENNE for non-convex mixed-integer nonlinear models, reach global optimality for up to roughly twenty aircraft, retain partial effectiveness around thirty, and degrade substantially by forty. This scalability limit should be interpreted as a formulation- and implementation-specific observation rather than as a fundamental algorithmic ceiling, but it nonetheless constitutes the classical performance frontier against which the capabilities of quantum methods will later be assessed in Section 1.3. Read against the taxonomy proposed by Kuchar and Yang (2000) more than two decades earlier, this body of work can be understood as the progressive, formal instantiation of the “optimized resolution” branch of their framework, moving from single-manoeuve exact models to combined, non-convex, and recovery-aware formulations. The next section therefore turns to the general framework of quantum combinatorial optimisation, before returning specifically to quantum formulations of the ACRP.

1.2 Quantum optimisation for combinatorial problems: concepts, formulations, and methodological cautions

1.2.1 The QUBO framework as a bridge to quantum hardware

The QUBO problem is typically defined as the optimisation of a quadratic form over binary vectors, where the coefficient matrix is a real-valued matrix whose diagonal terms encode linear coefficients and whose off-diagonal terms encode pairwise interactions between binary variables (Lewis & Glover, 2017). This formulation is equivalent to the Ising model through a standard binary-to-spin transformation, and it constitutes the common optimisation representation from which both annealing and gate-based approaches derive their problem encoding. As a unifying representation, QUBO integrates a wide range of combinatorial optimisation problems, including network flows, scheduling, max-cut, max-clique, and vertex cover, into a single modelling framework, and it is this generality that motivates its role as the entry point for quantum hardware.

This generality is not merely theoretical. Koretsky et al. (2021) adapt QAOA to the Unit Commitment problem, a mixed-binary scheduling task in electricity production planning, in which operational constraints such as minimum up-time and ramping limits must be absorbed into the QUBO objective through penalty terms; their implementation on Qiskit simulators shows that a scheduling-type combinatorial problem, structurally distant from graph partitioning or set packing, can be brought into the same QUBO representation discussed above. Azad et al. (2023) provide a second illustration by formulating the Vehicle Routing Problem as a QUBO and solving it with QAOA, encoding routing constraints such as single-visit conditions and capacity limits directly into the Ising Hamiltonian, on small benchmark instances published in IEEE Transactions on Intelligent Transportation Systems. Fitzek et al. (2024) extend this line of work to the heterogeneous Vehicle Routing Problem, in which a mixed fleet of vehicles with different capacities must be routed, and examine the feasibility and scalability of their QAOA-based QUBO encoding on noisy intermediate-scale quantum hardware. Read side by side, these three applications confirm the breadth of the QUBO framework beyond the graph-theoretic examples typically used to illustrate it, while also showing that, as of the most recent publications available, QAOA-based solutions to routing and scheduling problems remain small-instance proofs of concept rather than operationally scalable techniques.

Constructing a valid QUBO for a constrained combinatorial problem requires translating operational constraints into the objective function itself, most commonly through penalty terms. Ruan et al. (2023) formalize this challenge for QAOA by classifying constraints into three categories, linear equality, linear inequality, and arbitrary form, and by proposing constraint-encoding schemes tailored to each category. For linear equality constraints, such as those found in graph partitioning, they design a mixer Hamiltonian whose associated graph connects all

feasible solutions while excluding infeasible ones, exploiting the Hamming-distance structure between admissible bit-strings. For linear inequality constraints, illustrated through the set packing problem, they adapt this construction by allowing single-bit transitions between feasible states, yielding a near-regular mixer graph. For constraints that do not fit either linear form, for instance scheduling problems with additional pairwise incompatibilities between tasks, they propose a simpler star-graph mixer centred on one known feasible solution, at the cost of reduced solution-space coverage.

A third strategy for handling hard constraints, distinct from penalty terms and from constraint-preserving mixers, is proposed by Shirai and Togawa (2024), who introduce a post-processing variationally scheduled quantum algorithm for constrained combinatorial optimisation. Rather than requiring the quantum circuit alone to guarantee feasibility, either through penalty weighting or through a specifically designed mixer, their approach combines a variationally scheduled quantum evolution with a classical post-processing stage that handles hard constraints after measurement. This delegates part of the constraint-satisfaction burden to classical computation, illustrating that penalty-based encoding, mixer-based encoding, and post-processing-based encoding constitute three genuinely distinct design philosophies for constrained QUBO problems, each with different trade-offs between circuit complexity and reliance on classical post-treatment.

The mixer-design conjecture of Ruan et al. (2023) is advanced as a hypothesis and subsequently tested on small graph-partitioning and set-packing instances: the regularity of the constraint-encoding operator, defined as the difference between the maximum and minimum degree of vertices in its associated feasible-solution subgraph, is positively associated with the probability of recovering the optimal solution under a fixed circuit depth. Their reported tests support this conjecture rather than proving it in general: the fully regular equality-constraint operator consistently outperforms less regular alternatives on graph-partitioning instances, and the near-regular inequality-constraint operator outperforms the markedly irregular star-graph operator on set-packing instances at low circuit depth, although this advantage narrows as circuit depth increases. The broader lesson is structural, and holds for constrained combinatorial optimisation under QAOA generally: mixer regularity appears empirically linked to solution quality, independent of any specific application domain, though it remains an empirically supported conjecture rather than a general theorem.

These encoding choices also affect performance beyond the mixer design itself. Lewis and Glover (2017) show that, once a QUBO has been constructed, its size can often be reduced before submission to a solver: they present five reduction rules capable of predetermining the optimal value of certain binary variables directly from the structure of the coefficient matrix, without altering the optimal solution. Applied to structured but not fully uniform QUBO instances, their preprocessor achieves average size reductions of about 30 percent and yields substantial speedups, up to two orders of magnitude on the largest and densest tested instances, when com-

bined with both an exact solver (CPLEX) and a tabu-search metaheuristic. Its methodological significance reaches beyond these specific test cases: the practical difficulty of a QUBO instance depends not only on its size in binary variables, but also on the density and distributional structure of its coefficients, a consideration that carries over directly to any QUBO encoding of a combinatorial problem, regardless of domain.

1.2.2 Quantum optimisation paradigms: a functional overview

Quantum annealing is a physical optimisation process in which a quantum system is initialised in the ground state of a transverse-field Hamiltonian and then evolved toward the ground state of a problem Hamiltonian encoding the target QUBO instance. The problem Hamiltonian generally takes the Ising form, where the coupling coefficients are derived directly from the QUBO matrix and the spin variables represent the binary decision variables in their Ising form. Dense QUBO interaction graphs cannot be represented directly on sparse annealing hardware and must therefore be embedded, replacing each logical qubit with a chain of coupled physical qubits; this embedding step increases both the required number of physical qubits and the effective precision demands of the model. The defining operational feature of quantum annealing is therefore its hardware specificity: it minimises a QUBO natively, without gates or programmable quantum circuits, but its performance is governed by embedding overhead and analogue precision rather than by digital circuit depth.

Gate-based variational algorithms use a different computational paradigm. In QAOA, a QUBO is first encoded as a cost Hamiltonian, and a parameterised quantum state is then prepared by alternating two families of unitary operators: a phase-separation operator, which imprints the objective function onto the phase of the quantum state, and a mixing operator, which redistributes amplitude across candidate solutions. This alternating structure is repeated a fixed number of times, corresponding to the circuit depth, and the resulting variational parameters are optimised classically by minimising the expectation value of the cost Hamiltonian. As shown by Ruan et al. (2023), the choice of mixer is not a secondary implementation detail: encoding hard constraints directly into the mixer, rather than relying solely on penalty terms in the cost Hamiltonian, restricts the effective search space explored by the circuit to the feasible subspace, which their experiments show consistently improves the probability of measuring a feasible, high-quality solution relative to unconstrained mixers. This principle, that constraint-aware mixer design outperforms penalty-only encodings, is a general structural feature of gate-based variational optimisation, independent of the specific combinatorial problem to which it is applied.

The reach of the QAOA framework becomes clearer once its non-graph-theoretic applications are considered jointly. Koretsky et al. (2021) apply the standard QAOA structure to unit commitment scheduling, where the phase-separation operator encodes production costs

and ramping penalties rather than graph-cut objectives. Azad et al. (2023) and Fitzek et al. (2024) apply the same alternating structure to homogeneous and heterogeneous vehicle routing, respectively, where the mixing operator must redistribute amplitude across candidate route assignments rather than across candidate graph partitions. These applications confirm that the phase-separation and mixing structure of QAOA is genuinely problem-agnostic at the level of its mathematical formulation, even though, as the corresponding papers acknowledge, achieving competitive solution quality still requires problem-specific tuning of the mixer, the circuit depth, and the classical optimizer used to update the variational parameters.

Although quantum annealing and gate-based variational algorithms differ sharply in implementation, they share a common limiting regime. In quantum annealing, the dominant bottleneck is analogue precision and embedding overhead. In gate-based architectures, the limiting factor is circuit noise: every entangling gate introduces errors, and the maximum executable circuit depth on current noisy superconducting hardware remains limited. In both paradigms, the consequence is the same: problem instances that can currently be treated on quantum hardware remain significantly smaller and more constrained than the strongest classical formulations available for comparable combinatorial problems. This is not yet evidence for or against quantum advantage in any specific domain; it is a statement about the present computational regime in which both paradigms operate. These two architectures must therefore be understood not as interchangeable quantum solvers, but as distinct hardware-constrained approaches whose capabilities depend on very different notions of representational cost.

1.2.3 Benchmarking cautions in quantum-classical comparisons

The two preceding subsections have established what the QUBO framework is and what the two main quantum paradigms can physically execute. Before turning to the aircraft-specific quantum literature in Section 1.3, a prior methodological question must be addressed: what would a valid claim of quantum advantage over classical methods actually require, and what does the current benchmarking literature reveal about the gap between that standard and what has been demonstrated so far? This question is not rhetorical. The computational frontier documented in Section 1.1, global optimality up to about twenty aircraft, partial but still meaningful exact performance around thirty, and substantial degradation by forty in the strongest reviewed exact formulation, defines the scale at which classical exact methods begin to break down and at which quantum methods would need to operate meaningfully to constitute an improvement. The literature reviewed in this subsection does not show such an improvement on instances of comparable scale in any domain, including the scheduling and routing applications discussed above, where Koretsky et al. (2021), Azad et al. (2023), and Fitzek et al. (2024) all restrict their experiments to small instances rather than to scales matching realistic operational settings. Instead, this literature identifies several methodological pitfalls that can produce premature or

misleading claims of quantum superiority. These findings do not invalidate the quantum program for combinatorial optimisation broadly, nor for the ACRP specifically; rather, they define the conditions under which that program can be evaluated rigorously.

The most structurally informative contribution on this point is provided by Costa et al. (2024), whose explicit objective is to test whether apparent quantum advantages in combinatorial optimisation survive rigorous benchmarking conditions. Their benchmark problem is the Sherrington-Kirkpatrick spin glass, a dense quadratic Ising problem with all-to-all couplings, chosen precisely because dense binary quadratic structures with penalty couplings are common across QUBO applications. The central result is that apparent quantum speedups can be artifacts of comparing quantum methods against unstructured classical baselines: an unstructured quantum method such as Grover Adaptive Search may appear favourable when benchmarked against brute force or random search, but this advantage is no longer clearly observed when the comparison is made against a structured classical method that exploits the mathematical structure of the objective. The methodological implication is general: benchmarking any quantum algorithm against brute-force or structure-agnostic baselines provides little evidence of practical advantage relative to the best structured classical methods available.

Complementary evidence at a more practical QUBO-solver scale is provided by Chen et al. (2025), who design a benchmarking study in dynamic portfolio optimisation. Their benchmark is methodologically relevant because it illustrates how large constrained binary quadratic problems can be compared against strong classical baselines at scales deliberately beyond what present quantum hardware can process directly. Against this benchmark, several quantum or quantum-inspired approaches are compared with both Gurobi and the GPU-based ABS2 QUBO heuristic. The results show that current classical and GPU-based QUBO methods remain the practical reference points: Gurobi generally provides stronger bounds and better final guarantees when enough runtime is available, while ABS2 reaches competitive solutions faster under tight time limits. The largest instances mainly serve as forward-looking benchmarks, since they exceed the direct capabilities of current quantum hardware. The general methodological lesson is that strong classical baselines remain indispensable when evaluating binary quadratic formulations at scales beyond near-term quantum hardware, regardless of application domain.

McDowall et al. (2025) offer the most complete cross-platform benchmarking framework in the corpus; their contribution is less a new algorithm than a benchmarking protocol applied consistently across gate-based, annealing, and classical approaches. Their test case is a QUBO derived from defective graphene configurations, in which the physical lattice is not fully connected but the QUBO graph becomes fully connected once the vacancy-count constraint is encoded through penalty terms. Several methodological contributions from this work are broadly relevant. First, they define metrics that jointly capture solution quality and computational cost in a hardware-agnostic way, including optimal solution probability, approximation ratio, user runtime, and isolated Quantum Processing Unit (QPU) time. Second, they show that the user

runtime of quantum annealing on dense penalised QUBOs can be dominated by the embedding procedure rather than by annealing time itself, an overhead often omitted when only device access times are reported. Third, they show that post-selection, removing samples that violate encoded constraints, can substantially improve the observed optimal-solution probability, and they stress that results must be reported explicitly before or after post-selection. Finally, on this specific benchmark, classical simulated annealing remains effective on substantially larger dense QUBO instances than the tested quantum methods.

The aggregate implication of this benchmarking literature can be stated compactly. A valid quantum-classical comparison requires at least five conditions: a classical baseline drawn from structured exact or strongly structured methods rather than brute force; problem instances whose scale approaches the frontier at which the best classical methods begin to degrade; metrics that include total wall-clock time and embedding overhead rather than only device runtime; transparent reporting of post-selection and hyperparameter tuning; and a clear distinction between claims about solution quality and claims about computational efficiency. These five criteria will structure the evaluation of the aircraft-specific quantum literature reviewed in Section 1.3. The purpose of that evaluation is not to dismiss quantum approaches, but to locate precisely where the current frontier stands, what scales have actually been reached, and what kind of hardware or algorithmic progress would be required before a credible performance crossover can be claimed for the ACRP specifically.

1.3 Quantum approaches applied to the ACRP

1.3.1 Annealing-based approaches

The earliest quantum-specific contribution to aircraft conflict resolution is the quantum annealing framework proposed by Stollenwerk et al. (2020), who map a strategic deconfliction problem for wind-optimal trajectories onto a QUBO formulation solved on D-Wave hardware. Conflicts are represented through a conflict graph in which vertices correspond to flights and edges to potential conflicts, and the decision variables encode discrete departure-delay choices for each flight rather than heading or speed changes. The QUBO objective combines a one-hot encoding penalty, a delay cost term, and a conflict penalty, with feasibility enforced entirely through penalty weighting rather than explicit hard constraints. Because dense QUBO graphs cannot be mapped directly onto D-Wave’s sparse hardware connectivity, the authors decompose a large real-world instance (984 transatlantic flights) into embeddable subproblems, reporting maximum embeddable sizes of roughly 50 flights and 104 conflicts on D-Wave 2X, and 64 flights and 261 conflicts on D-Wave 2000Q. This embedding overhead is substantial: because a single logical qubit must be realised as a chain of physical qubits on the sparse hardware connectivity, the reported instances expand from roughly 91 logical qubits to an average of

631 physical qubits on D-Wave 2X, and from 133 logical to about 1,235 physical qubits on D-Wave 2000Q. Solution times for the largest embeddable subproblems remain close to one second, but success probability declines as discretization becomes finer, a limitation attributed to coefficient-precision constraints on the annealer rather than to the annealing process itself. Notably, the authors report that the penalty weights sufficient to enforce feasibility are largely independent of the specific instance, which supports the robustness of the mapping at the modelling level even as hardware precision, rather than the annealing dynamics, remains the binding practical constraint. This delay-only formulation should be read as a proof of concept for hardware-compatible QUBO mapping rather than as a demonstration of practical advantage over classical ACRP solvers, since it addresses a simplified strategic subproblem rather than the tactical heading/speed deconfliction problem discussed in Section 1.1.

1.3.2 Gate-based variational approaches

Gate-based variational algorithms shift the modelling task from QUBO embedding to the design of a problem Hamiltonian and, where applicable, a constraint-aware mixer. Pecyna et al. (2024) formulate the tactical aircraft deconfliction problem using binary variables denoting whether aircraft i performs manoeuvre j , combined with a four-dimensional conflict matrix that penalises incompatible simultaneous manoeuvre choices. They implement this Hamiltonian within both the standard Quantum Approximate Optimization Algorithm (QAOA) and the Quantum Alternating Operator Ansatz (QAOAnsatz), the latter embedding the one-manoevre-per-aircraft constraint directly into the mixer Hamiltonian instead of penalising it in the cost function. Tested on downscaled instances of the Random Circle Problem benchmark (3 to 5 aircraft, 2 to 5 manoeuvres, corresponding to 9–25 qubits) on the 133-qubit IBM Torino processor, QAOAnsatz solves four of five tested instances and consistently outperforms vanilla QAOA on a noisy simulator, though both algorithms struggle as the number of manoeuvres per aircraft increases, since enforcing the one-hot constraint over a larger manoeuvre set requires more entangling gates and therefore more noise exposure. Indeed, the authors report the counter-intuitive finding that, for a fixed qubit budget, increasing the number of manoeuvres per aircraft is more detrimental to performance than increasing the number of aircraft, and that on the hardest five-aircraft configurations the success probability can fall below one percent, remaining operationally meaningful only because repeated circuit measurements can still surface a correct solution on such small instances.

Building directly on this formulation, Pecyna and Różycki (2024) argue that the conventional one-hot Hamiltonian is itself a poor fit for gate-based hardware, since its complexity and associated circuit noise grow with the number of manoeuvres. They propose relaxing the exact one-manoevre-per-aircraft constraint into a weaker condition based on a NOT-XOR Hamiltonian, allowing a measured bitstring to represent a set of conflict-free manoeuvre candidates

with final manoeuvre selection deferred to classical post-processing. On nine artificial deconfliction instances, this relaxed encoding finds feasible solutions in all nine cases under a relaxed evaluation criterion, compared with seven of nine for the standard one-hot QUBO Hamiltonian, illustrating that constraint encoding, not only the choice of optimizer, is a first-order performance factor in gate-based ACRP formulations.

1.3.3 Adjacent domain: UAV deconfliction

A related but distinct application appears in Huang et al. (2022), who apply variational quantum algorithms to collision avoidance for unmanned aerial vehicles rather than crewed tactical air traffic. The authors formulate the UAV collision-avoidance problem as a QUBO, expressed as an Ising Hamiltonian and solved with two variational algorithms, the Variational Quantum Eigensolver (VQE) and QAOA, further enhanced by a Conditional Value-at-Risk (CVaR) objective to improve convergence toward feasible solutions. Across four test instances, this approach exceeds a 90% probability of finding a feasible solution under suitable parameter settings. This paper is retained briefly as an adjacent reference point, showing that variational QUBO-based deconfliction extends beyond commercial ACRP into UAV traffic management, without receiving the same depth of discussion as the ACRP-specific contributions above.

1.3.4 Critical synthesis and research gap

Set beside the classical frontier established in Section 1.1, the annealing and gate-based literatures on the ACRP share a common structural limitation. The strongest classical exact formulations reach global optimality for up to roughly twenty aircraft, retain partial effectiveness around thirty, and degrade substantially by forty. Aircraft counts must nevertheless be compared within the same problem class: Stollenwerk et al.’s (2020) embeddable delay-only subproblems reach 50–64 flights after decomposition of a much larger instance, a count that is not directly comparable to tactical heading-and-speed instances; Pecyna et al.’s (2024) gate-based experiments are confined to three to five aircraft with two to five manoeuvres by qubit counts and NISQ noise; and Huang et al.’s (2022) UAV instances are similarly small. A structured side-by-side comparison of these contributions, covering their paradigm, the problem they actually solve, their maximum instance scale, their constraint-handling strategy, and their benchmarking rigour, is provided in Appendix A; the synthesis below draws on that comparison to characterise the gap rather than to restate each study.

Beyond raw scale, the comparison shows that these studies do not solve the same problem. The annealing branch addresses a strategic, delay-only deconfliction subproblem in which trajectories are fixed and only departure delays are optimised, whereas the gate-based branch addresses a tactical manoeuvre-selection problem far closer to the heading-and-speed formula-

tions of Section 1.1, but only for a handful of aircraft. Huang et al. (2022) lie outside crewed air traffic altogether. Consequently, no reviewed quantum study benchmarks a tactical, two-dimensional ACRP against a structured classical formulation on the same controlled instances, with transparent tuning and cost reporting, at a scale where classical exact methods themselves begin to struggle. The studies also differ in how they enforce feasibility, instantiating the three encoding philosophies distinguished in Section 1.2: penalty-only weighting in the annealing QUBO, a constraint-preserving mixer in QAOAnsatz, and constraint relaxation with classical post-processing in the NOT-XOR formulation. Tellingly, these three philosophies have never been compared against one another on a common ACRP instance.

Assessed against the five benchmarking criteria set out in Section 1.2.3, the studies leave several gaps, with important differences between paradigms. First, Stollenwerk et al. (2020) include exact Max-SAT and classical Isoenergetic Cluster Method (ICM) references on their delay QUBO. These are relevant classical comparisons for that problem; a tactical heading-and-speed solver would require a different decision model. Their results do not establish practical quantum advantage. The gate-based ACRP studies, by contrast, focus on feasibility without a matched classical optimum benchmark. Second, gate-based tactical experiments with three to five aircraft remain smaller than the twenty-to-forty-aircraft cases explored by selected richer classical formulations. This is a limitation of demonstrated scale, not a universal tractability threshold, and classical comparisons remain meaningful even on small matched instances. Third, reported metrics privilege annealing time or success probability, whereas total wall-clock time and embedding overhead, which McDowall et al. (2025) show can dominate runtime on dense penalised QUBOs, are rarely accounted for. Fourth, post-selection and the tuning of penalty weights or variational parameters are described qualitatively but seldom reported transparently enough to be reproduced. Fifth, claims about solution quality and claims about computational efficiency are not cleanly separated, so that a high feasible-solution probability on a very small instance can be misread as evidence of computational competitiveness. No reviewed study satisfies all five criteria simultaneously.

These shortcomings resolve into four distinct dimensions of the research gap. The first is a scale gap for gate-based tactical ACRP: the demonstrated instances remain small compared with those tested in richer classical formulations. Delay-only annealing subproblems require a separate scale comparison. The second is a problem-fidelity gap: quantum studies either solve strategic delay variants, adjacent UAV problems, or drastically downscaled tactical manoeuvre-selection instances that are not benchmarked against the corresponding structured classical formulation. The third is a benchmarking-rigour gap: the gate-based tactical studies lack a matched structured classical baseline, while comparable scales and complete cost accounting remain broader limitations, precisely the conditions that the general benchmarking literature (Costa et al., 2024; Chen et al., 2025; McDowall et al., 2025) identifies as prerequisites for any credible claim. The fourth is a reproducibility gap: penalty calibration, parameter

tuning, and post-selection are reported too partially for independent replication.

The contribution of this thesis is defined directly against these four dimensions. Rather than pursuing a demonstration of quantum advantage, which the present computational regime does not support, it constructs a simplified two-dimensional tactical ACRP benchmark and compares quantum and quantum-inspired methods with a structured classical baseline on the same instances. Both experimental manoeuvre sets are heading-only: $M = 3$ uses $\{-30^\circ, 0, +30^\circ\}$ and $M = 5$ uses $\{-30^\circ, -15^\circ, 0, +15^\circ, +30^\circ\}$, with nominal speed fixed. This design deliberately accepts the scale gap and aims to narrow, rather than claim to close, the problem-fidelity, benchmarking-rigour, and reproducibility gaps. Fixing a small tactical benchmark makes a like-for-like comparison feasible on accessible simulators and hardware; the study reports the available wall-clock, device, circuit, sampling, and tuning costs, while identifying where complete user-time accounting is unavailable. It also separates solution-quality claims from efficiency claims throughout. The research question that structures the remainder of this work is therefore not whether the ACRP can be encoded for quantum hardware, which is already established across both paradigms, but whether, on a controlled 2D tactical benchmark and under transparent conditions, a quantum or quantum-inspired formulation can match or approach a structured classical baseline in solution quality, and at what measurable computational cost.

Chapter 2

Methods, Data Sources and Analysis

2.1 Research design and positioning

This study follows a **computational-experimental design**: it does not test a hypothesis about human behaviour but measures, under controlled and reproducible conditions, the behaviour of competing optimisation algorithms on a common family of problem instances. The unit of analysis is a *solver run* on an *instance*, characterised by its solution quality, its feasibility, and its computational cost.

The design is deliberately anchored on the four gaps identified in the literature review. The **scale gap** is accepted rather than solved: exact statevector simulation limits faithful quantum experiments to small instances, so the contribution is a proof of concept plus a scaling analysis rather than a large-scale demonstration. The **problem-fidelity gap** is narrowed through a shared, simplified tactical benchmark with heading-only manoeuvres at nominal speed; $M = 3$ and $M = 5$ differ only in angular granularity. The **benchmarking-rigour gap** is narrowed by comparing the solvers on identical instances against a structured exact baseline and by reporting the cost components that were retained. The **reproducibility gap** is narrowed through fixed seeds, documented penalty calibration, versioned instances and auditable hardware records. These measures support strong within-scope comparisons without implying that the continuous operational ACRP, complete user-time accounting or every stochastic campaign can be reproduced from the present export.

Because the project does not map onto any single row of the SKEMA methods grid (interviews, surveys, netnography, and so on), it is positioned under the “atypical / highly original” clause. Equivalent rigour is argued through the exact discrete ground truth, the formulation shared by the QUBO, discrete Gurobi ILP, simulated annealing and QAOA, and the statistical treatment described in Section 2.7. This positioning has been validated with the supervisor.

2.2 Problem formalisation and QUBO encoding

2.2.1 Tactical 2D ACRP

An instance has n aircraft. Aircraft i starts at position (x_i^0, y_i^0) with initial heading θ_i^0 and nominal speed v_i^0 . The richer source formulation permits a **manoeuvre** (q_i, δ_i) combining a speed multiplier $q_i \in [0.94, 1.03]$ with a heading deviation $\delta_i \in [-\frac{\pi}{6}, +\frac{\pi}{6}]$. To keep the QUBO small and auditable, the retained experiments fix $q_i = 1$ and optimise heading choices only. The common velocity expression is

$$\mathbf{v}_i = q_i v_i^0 (\cos(\theta_i^0 + \delta_i), \sin(\theta_i^0 + \delta_i)).$$

Two aircraft i and j are separated if their horizontal distance never drops below the separation minimum d while they are still approaching each other. Writing the relative initial position $(x_0^r, y_0^r) = (x_i^0 - x_j^0, y_i^0 - y_j^0)$ and the relative velocity $(v_{rx}, v_{ry}) = \mathbf{v}_i - \mathbf{v}_j$, the pair is **in conflict** if and only if

$$\begin{aligned} \text{sep}_0 &= (y_0^r - d)^2 v_{rx}^2 + (x_0^r - d)^2 v_{ry}^2 - 2x_0^r y_0^r v_{rx} v_{ry} < 0 \\ \text{and } t_{m0} &= -\frac{x_0^r v_{rx} + y_0^r v_{ry}}{v_{rx}^2 + v_{ry}^2} > 0. \end{aligned}$$

This is the same analytic conflict criterion used by both classical implementations and the QUBO. Sharing the criterion does not by itself make a continuous and a discrete decision domain identical. Exact like-for-like comparisons in this thesis therefore use the discrete Gurobi ILP and the QUBO on the same candidate set; the original continuous MINLP is retained only as a contextual cross-check. Both use the active objective of **minimising the number of manoeuvring aircraft**.

2.2.2 Discretisation into candidate manoeuvres

The heading control is discretised into a finite ordered set of M candidate manoeuvres shared by all aircraft, with index 0 denoting the nominal choice ($q = 1, \delta = 0$). Two sets are retained throughout: $M = 3$ uses $\{-30^\circ, 0, +30^\circ\}$ and $M = 5$ uses $\{-30^\circ, -15^\circ, 0, +15^\circ, +30^\circ\}$. In both cases $q = 1$. Thus $M = 5$ is a finer heading grid, not a speed-variation experiment; $M = 3$ is used for the principal benchmark because it requires only $3n$ binary variables.

2.2.3 QUBO formulation

A binary variable $x_{i,m} \in \{0,1\}$ equals 1 when aircraft i selects manoeuvre m . The energy to minimise is

$$E(\mathbf{x}) = C \sum_{i=1}^n \sum_{m=0}^{M-1} c_{i,m} x_{i,m} + A \sum_{i=1}^n \left(\sum_{m=0}^{M-1} x_{i,m} - 1 \right)^2 + B \sum_{(i,j) \in P} \sum_{(m,m') \in K_{ij}} x_{i,m} x_{j,m'}.$$

The three terms are, in order: the **objective**, with unit cost $c_{i,m} = 0$ for the nominal manoeuvre and 1 otherwise, so minimising it minimises the number of manoeuvring aircraft; the **one-hot constraint**, which forces exactly one manoeuvre per aircraft; and the **separation penalty**, where K_{ij} is the set of manoeuvre pairs (m,m') that leave aircraft i and j in conflict according to the criterion above. Because the variables are binary, $x^2 = x$, and the one-hot square expands to linear and quadratic terms plus a constant offset; the full derivation is given in Appendix B.

2.2.4 Penalty calibration

The objective weight is fixed at $C = 1$ throughout the retained campaign. The weights A and B must be large enough that the global optimum of E is a feasible, conflict-free, one-hot assignment whenever one exists, yet not so large that the energy landscape becomes needlessly rugged for a sampler. With unit manoeuvre cost, a conservative sufficient threshold is $A, B > nC$; the implemented default is $A = B = 2(n+1)C = 2(n+1)$.

The bound $A, B > nC$ is sufficient, not necessary. For the controlled CP_3 instance (radius 100, separation $d = 8$, fixed speed and headings $\{0, -30^\circ, +30^\circ\}$), exhaustive enumeration of all 512 bitstrings gives a sharper result. With $A = B = \alpha$ and $C = 1$, the minimum QUBO energy is $\min(2\alpha, 2)$ for $\alpha > 0$. Below $\alpha = 1$, every ground state is infeasible; at $\alpha = 1$, feasible and infeasible ground states tie; above $\alpha = 1$, every ground state is feasible. The all-zero vector costs 3α and is not the winning state for positive α below this transition.

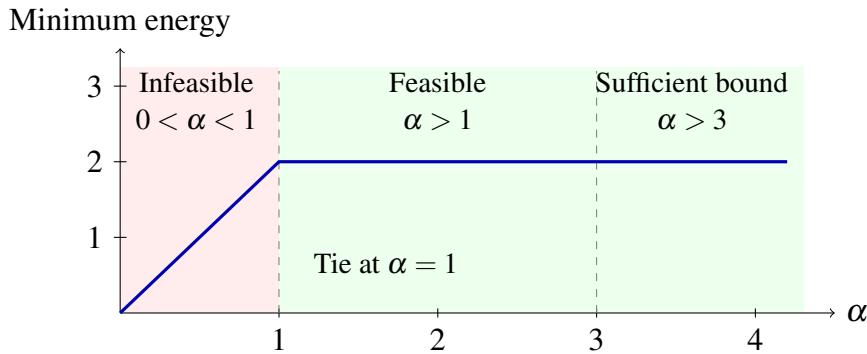


Figure 2.1: Exact penalty feasibility on CP_3 with $A = B = \alpha$ and $C = 1$, obtained by enumerating all 512 bitstrings. The transition at 1 is instance-specific; the general sufficient bound is 3. This figure describes QUBO ground states, not QAOA sampling probabilities.

The default $A = B = 2(n + 1)$ is retained as a conservative setting. Increasing penalties beyond the feasibility transition gives no further guarantee about sampling quality: QAOA probabilities depend on the angles, their optimisation and the relative energy scales. The exact feasibility check therefore does not establish a monotonic deterioration of QAOA performance at large weights.

2.3 Data and instance generation

Instances come from two sources. First, the canonical **Circle Problem** CP_n : n aircraft equally spaced on a circle, each heading towards the centre at uniform speed, so that every pair converges simultaneously and the do-nothing configuration is maximally infeasible. This is the standard stress test of the ACRP literature and is generated deterministically from $(n, \text{radius}, d, \text{speed})$. Second, **randomised instances** with controlled conflict density, used to avoid over-fitting conclusions to the highly symmetric circle geometry.

The controlled experimental factors are the number of aircraft n , the manoeuvre-set granularity M , and the conflict geometry. Every stochastic element (instance randomisation and stochastic solvers) is driven by an explicit seed so that any run can be reproduced exactly. The instance library reuses the .dat format of the baseline so that the classical and quantum pipelines consume identical inputs.

A first cross-check uses the supervisor’s original continuous MINLP on CP_3 (three aircraft, uniform speed and separation $d = 8$). It returns an objective of two manoeuvring aircraft and selects small speed changes with heading deviations near -15° . This is useful evidence that the heading grid captures a plausible continuous solution, but it is not a proof of identical feasible sets: the continuous MINLP can choose values that the discrete benchmark cannot.

Table 2.1: Continuous MINLP cross-check on CP_3 . These controls are contextual outputs from the richer source formulation, not decisions available to the discrete $M = 3$ QUBO.

Aircraft	Manoeuvre	Speed q	Heading change
1	No	1.000	0°
2	Yes	0.989	-14.4° (nearest option -15°)
3	Yes	0.986	-15.6° (nearest option -15°)

The exact like-for-like validation uses a second Gurobi implementation: a binary ILP whose variables, candidate headings, conflict pairs and unit manoeuvre objective are identical to those of the QUBO. On CP_3 , CP_4 and CP_5 , this discrete ILP and exhaustive QUBO enumeration return the same optimum. The continuous MINLP also returns the same number of manoeuvring aircraft on these three circles, but that equality of objective values is reported only as an external

consistency check. The table below records the strict discrete comparison, from nine to fifteen qubits.

Table 2.2: Exact objective comparison between the discrete Gurobi ILP and brute-force QUBO enumeration on the $M = 3$ heading grid.

Instance	Aircraft n	QUBO variables	Discrete Gurobi ILP	Brute-force QUBO
CP_3	3	9	2	2
CP_4	4	12	3	3
CP_5	5	15	4	4

This agreement demonstrates that the discrete Gurobi ILP and the QUBO encode the same bounded heading-selection problem. It is the entry point for the like-for-like comparison in Chapter 3. The result must not be extended to the continuous MINLP or to combined heading-and-speed ACRP formulations.

2.4 Solvers compared

Four solver roles are compared on common instances. The distinction between the two Gurobi formulations is explicit.

1. **Gurobi exact reference.** The principal benchmark uses the discrete binary ILP in `src/exact_gurobi.py`, which has the same candidate headings and conflict-choice constraints as the QUBO. It supplies like-for-like exact objectives and statuses. The supervisor’s original continuous MINLP, which requires Gurobi 12 or later, is used separately as a contextual formulation check and is not treated as the same decision domain.
2. **Exact brute-force QUBO solver.** Enumerates all 2^V assignments for small variable counts V . It is not a competitive solver but a *validation instrument*: where enumeration is tractable, its optimum must be feasible and must equal the discrete Gurobi optimum.
3. **Simulated annealing (classical twin).** A single-spin-flip Metropolis annealer with a geometric temperature schedule, run on the same QUBO as the quantum solver. Because it shares the formulation with QAOA, any performance difference isolates the effect of the algorithm rather than of the model.
4. **QAOA (quantum solver).** The QUBO is mapped to an Ising cost Hamiltonian through $x_k = (1 - Z_k)/2$, giving a cost Hamiltonian in Z and ZZ terms. A p -layer alternating ansatz (cost unitary followed by a transverse-field mixer) is optimised classically. Two independent implementations are used: a Qiskit implementation on the Aer simulator (noiseless,

then with a noise model) for the thesis runs, and a dependency-free numpy statevector implementation used as an exact cross-check and for offline analysis. The two must agree on the expected energy for identical parameters, which is an explicit reproducibility control.

An optional emulated quantum annealer (D-Wave neal) can sample the same QUBO as a second quantum-inspired point.

2.4.1 QAOA depth as a configuration variable

QAOA quality is expected to improve with the circuit depth p when the added parameters are adequately optimised. A preliminary calibration on one fixed CP_3 instance swept $p = 1$ to 6 with one seed and 120 optimiser restarts per depth. It recorded the probability of sampling the optimum and the probability mass satisfying the one-hot assignment constraint. Success probability increased from below one per cent at $p = 1$ to roughly thirty per cent at $p = 6$, although the curve was not strictly monotone at every intermediate depth. This calibration motivated reporting depth explicitly; it is not pooled with the later multi-seed, noise-model or frozen-angle hardware series.

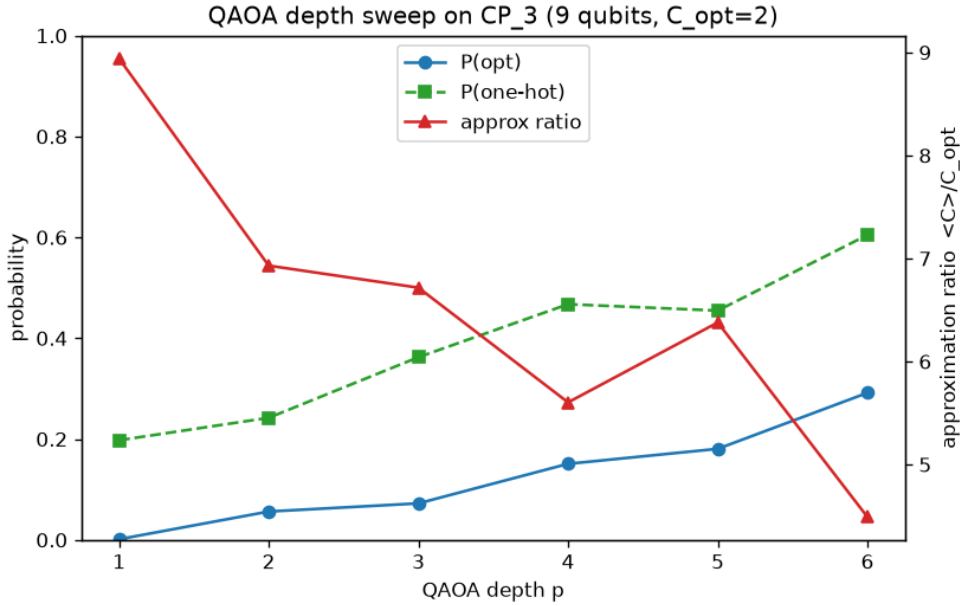


Figure 2.2: Preliminary QAOA depth calibration on CP_3 using one seed and 120 optimiser restarts per depth. $P(\text{one-hot})$ measures only the assignment constraint, not full conflict-free feasibility.

2.4.2 Noise sensitivity of QAOA

Beyond depth, the second configuration variable that matters for a quantum solver is hardware noise. To measure its effect without leaving the reproducible simulator, the optimised QAOA

circuit at each depth is sampled twice: once on an ideal statevector, and once on an Aer simulator carrying an explicit depolarising noise model, with a one-qubit error rate of 10^{-3} , a two-qubit error rate of 10^{-2} , and 8,192 measurement shots. Because both runs use the same optimised angles, their difference is associated with the specified sampling-noise model while holding the variational optimisation fixed. It does not isolate every mechanism present on real hardware.

The variational angles are obtained by keeping the best of several random restarts of the classical optimiser, mirroring the restart strategy of the numpy reference implementation. This detail is not cosmetic. With a single optimiser start the reported quality is erratic across depth, because the non-convex angle landscape traps the optimiser in poor local minima; restarts restore a smooth, interpretable depth trend and are treated here as part of the QAOA configuration.

On CP_3 the campaign gives the values below, now swept from $p = 1$ to $p = 4$ with twenty-four optimiser restarts per depth. The ideal curve is clean and monotone: the success probability rises from roughly nine per cent at $p = 1$ to nineteen per cent at $p = 4$, and the feasibility rate climbs in parallel from twelve to twenty-one per cent. The noisy curve tells the opposite story. It barely improves with depth and then falls back, so the two curves, almost level at $p = 1$, separate widely by $p = 4$. This growing gap is the central observation: adding layers makes the ideal circuit strictly better, but the extra gates accumulate enough noise to cancel the benefit on a noisy device. The expected energy confirms it, remaining far higher for the noisy runs at every depth.

Table 2.3: Ideal and depolarising-noise QAOA results on CP_3 , using 24 optimiser restarts per depth and 8,192 noisy samples.

Depth p	Run	P(opt)	Feasibility rate	Expected energy
1	Noiseless	0.093	0.120	14.52
1	Noisy	0.081	0.107	16.32
2	Noiseless	0.111	0.124	13.34
2	Noisy	0.084	0.098	17.08
3	Noiseless	0.179	0.210	10.35
3	Noisy	0.096	0.112	22.27
4	Noiseless	0.186	0.212	12.45
4	Noisy	0.073	0.085	22.96

The same comparison was then repeated on CP_4 . Moving from nine to twelve qubits increases the ideal state space from $2^9 = 512$ to $2^{12} = 4,096$ basis states, an eightfold increase, and the optimum becomes far rarer in the sampling distribution. Two findings carry over and one new lesson appears. The noise penalty carries over cleanly and is stronger under the specified model. At every depth the noisy circuit samples the optimum less often, satisfies the constraints less often, and shows a much higher expected energy, with the energy gap widening from about

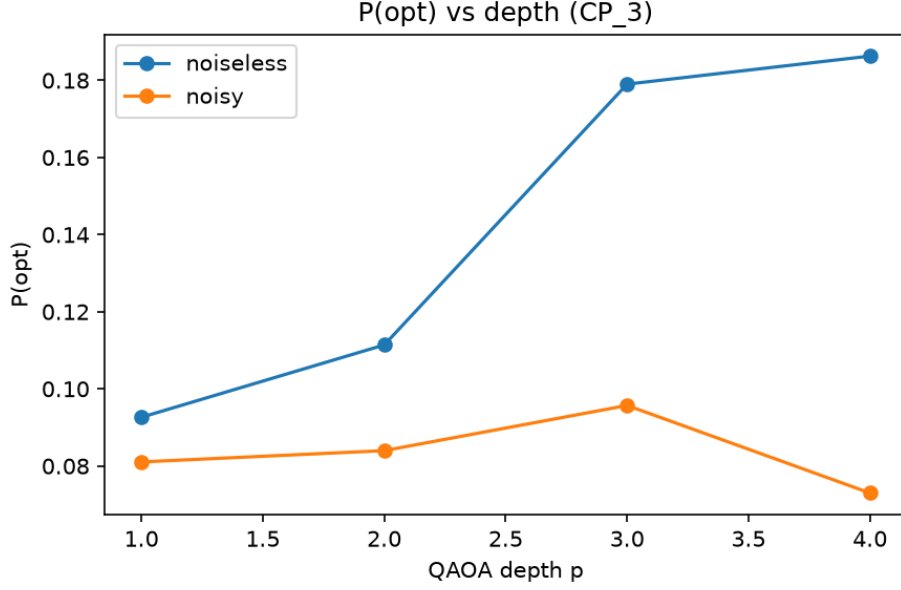


Figure 2.3: Probability of sampling the optimum against QAOA depth on CP_3 , ideal versus depolarising noise. The ideal curve rises with depth while the noisy curve stalls, so the gap widens.

five units at $p = 1$ to more than twenty at $p = 3$ and $p = 4$. What does not carry over is the smoothness of the ideal curve: even with twenty-four restarts the noiseless quality is no longer monotone in p , dipping at $p = 2$, peaking at $p = 3$, then falling at $p = 4$. Each added layer introduces two more variational angles, and a fixed restart budget under-optimises the deepest circuits. This is an honest limit of the method rather than a defect of the comparison. The noise contrast is read at fixed angles and therefore remains valid for this model, but CP_4 shows that the classical outer optimisation, not the sampling, becomes the binding constraint on QAOA quality as the instance grows.

Table 2.4: Ideal and depolarising-noise QAOA results on CP_4 , using 24 optimiser restarts per depth and 8,192 noisy samples.

Depth p	Run	P(opt)	Feasibility rate	Expected energy
1	Noiseless	0.031	0.037	26.49
1	Noisy	0.023	0.028	31.02
2	Noiseless	0.026	0.029	25.91
2	Noisy	0.015	0.018	35.63
3	Noiseless	0.050	0.055	22.87
3	Noisy	0.020	0.022	49.26
4	Noiseless	0.016	0.020	27.46
4	Noisy	0.007	0.009	49.37

Two caveats are carried into Chapter 3. First, the depolarising model is an idealisation of real hardware: it represents a controlled sensitivity scenario, not a mathematical lower bound

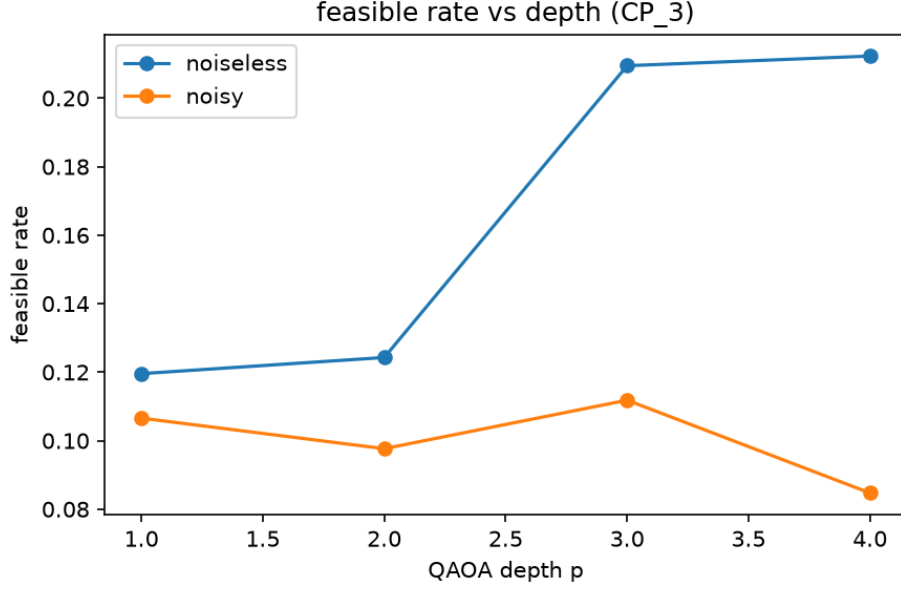


Figure 2.4: Feasibility rate against QAOA depth on CP_3 , ideal versus depolarising noise. The share of constraint-satisfying samples grows with depth only in the ideal case.

or a device prediction, because it omits device-specific error channels and crosstalk. Second, as CP_4 makes clear, the classical optimisation of the QAOA angles becomes the limiting factor at twelve qubits and beyond, so a fair depth comparison at larger sizes will require a restart budget that grows with p . The simulation is reproducible for the recorded settings, but its conclusions remain conditional on this noise model and restart budget.

2.5 Experimental protocol

The campaign is organised as a matrix of *configurations*, where each configuration is a triple made of an instance, a solver, and the solver’s settings. Repetition depends on the evidence source. The principal stochastic benchmark uses fixed seed lists; QAOA also uses multiple optimiser starts. Noise comparisons reuse one selected angle vector and contrast ideal probabilities with 8,192 noisy shots. Hardware jobs use 4,096 shots per circuit, and only the final interleaved job repeats an identical circuit within the same job. Exact solvers are run once per instance because they are deterministic. These repetition types are not interchangeable: seeds measure algorithmic variation, shots measure sampling variation, and the three interleaved $p = 2$ circuits provide only a within-job stability check.

Quality is reported together with every cost component that was retained. Classical result files include solver runtimes for the recorded runs; the quantum records additionally retain qubit count, routed circuit depth, two-qubit gates, shots and, for the last two IBM jobs, job-level QPU usage. Complete end-to-end user time is not available for every configuration. The repository contains separate entry points for the benchmark, depth calibration, noise study, statistical anal-

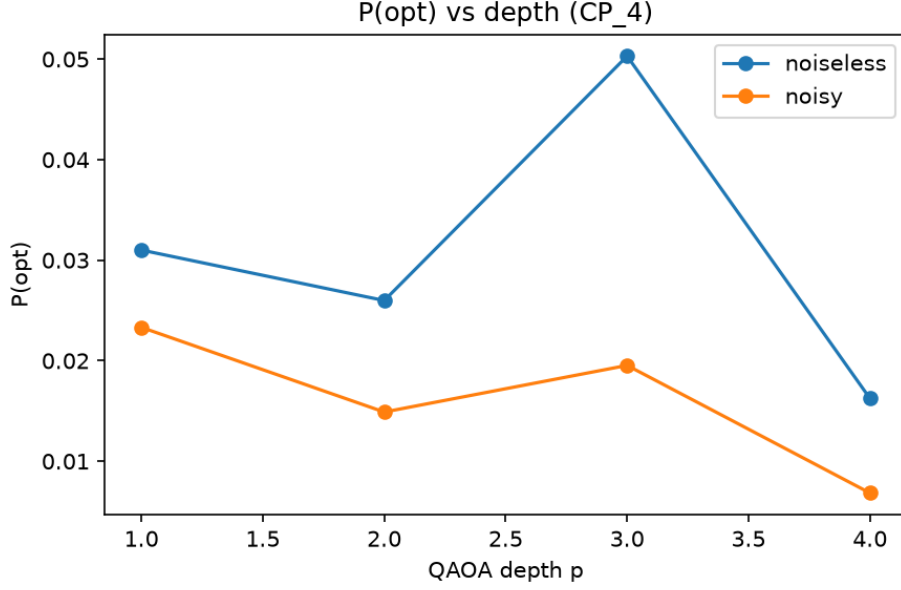


Figure 2.5: Probability of sampling the optimum against QAOA depth on CP_4 (twelve qubits), ideal versus depolarising noise. The noisy curve is below the ideal at every depth, and the ideal curve is no longer monotone because the deeper circuits are under-optimised.

ysis and IBM campaigns rather than one universal command. Its `requirements.txt` gives minimum constraints for the core and simulator stack but is not a frozen environment file; in particular, the current export does not list `qiskit-ibm-runtime`. The final IBM plan records the executed hardware environment exactly: Python 3.12.13, NumPy 2.5.2, SciPy 1.18.0, Qiskit 2.5.2 and Qiskit IBM Runtime 0.49.0. Reproduction therefore requires selecting the script corresponding to the reported result and using its retained seed, plan or manifest.

2.6 Evaluation metrics and scaling behaviour

Solvers are compared on a small set of complementary metrics rather than on a single figure of merit, because a quantum sampler can score well on one axis while failing on another. Feasibility is reported before quality: an assignment that is not one-hot and conflict-free is discarded regardless of its energy, so a low-energy but infeasible sample is never mistaken for a good solution. Quality is then read jointly with cost, which closes the “at what cost” half of the research question. The full metric set is summarised below.

These metrics are then tracked as the instance grows, because the qubit count rises as $V = n \times M$ and even modest instances leave the range of exact statevector simulation quickly. Across the Circle Problem ladder from nine to eighteen qubits, the classical twin reaches the exact optimum at low and slowly growing cost, while the QAOA success probability decays sharply and its simulation time rises steeply. This divergence is the empirical face of the scale gap identified in Chapter 1.

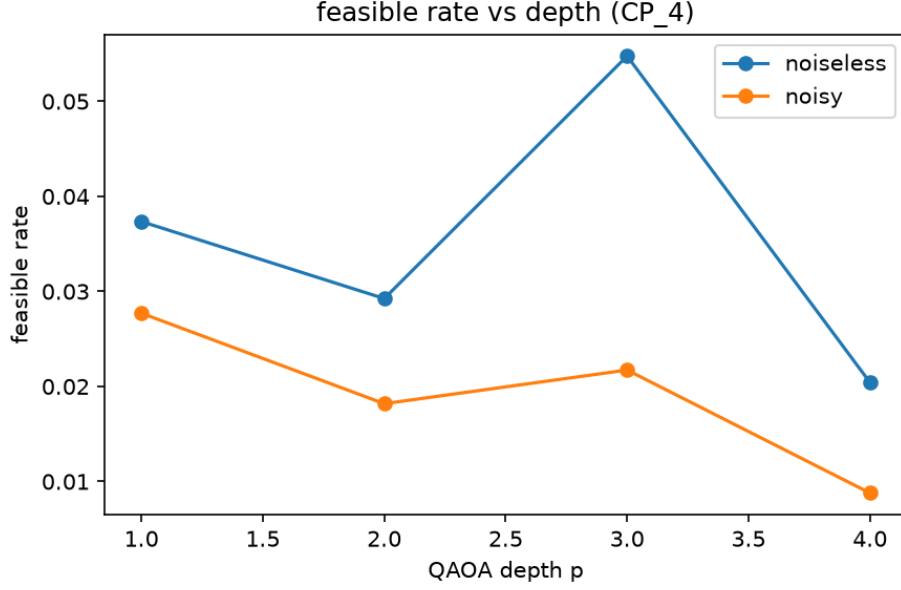


Figure 2.6: Feasibility rate against QAOA depth on CP_4 (twelve qubits), ideal versus depolarising noise. Noise lowers the constraint-satisfying share at every depth.

To make this trend quantitative without over-claiming, the two scaling relationships are summarised by a simple fit over the range that exact simulation allows: a linear regression on the logarithm of the cost, and a monotone decreasing fit for the decay of $P(\text{opt})$. Because that range is narrow, the fits are presented explicitly as extrapolation; they describe the observed trend and motivate a projection towards larger instances, but any statement beyond the simulated range is flagged as a hypothesis for future hardware rather than a measured result. This projection is revisited in Chapter 3 once the full campaign has populated the scaling table.

2.7 Statistical analysis

A stochastic solver can return a different answer on each run, so isolated outcomes are not treated as general evidence. Where the retained data contain repeated seeds or finite-shot counts, Chapter 3 reports distributions, intervals or paired tests. Deterministic checks, projections and configurations for which the necessary repeated data were not retained are identified separately and are not given artificial uncertainty estimates.

2.7.1 Replication and interval estimation

Two sources of variation are distinguished when the relevant records permit it. Solver variation is studied through fixed seed lists and optimiser restarts in the main stochastic benchmark; finite-shot variation is treated separately for sampled distributions. Instance variation is represented by combining deterministic Circle Problems with randomised instances of controlled

Table 2.5: Evaluation metrics used across classical, simulated and hardware results.

Metric	Definition	Why it matters
Feasibility rate	Share of runs returning a one-hot, conflict-free assignment	A quantum sampler often returns infeasible states; quality is meaningless without it
Number of manoeuvres	Objective value of the best feasible assignment	The quantity the ACRP actually minimises
Optimality gap	Difference to the exact optimum (brute force / Gurobi)	Absolute solution quality
Success probability $P(\text{opt})$	Probability of sampling the optimum in a single shot	The natural quality measure for a sampler
Approximation ratio	Expected energy divided by the optimal energy	Depth- and configuration-independent quality summary
Time-to-solution / cost	Wall-clock time; qubits, depth, gates, shots	The “at what cost” half of the research question

conflict density. Paired comparisons use common instances and the same QUBO, but simulated annealing reads and QAOA optimiser starts are algorithm-specific rather than matched random events. Hardware repetitions across different days were not performed, so shot intervals must not be read as calibration or day-to-day uncertainty.

For finite-shot proportions in the IBM section, Wilson score intervals are reported because they remain within the zero-to-one range. Selected averages across repeated observations use percentile bootstrap intervals, with the observation as the resampling unit; the relevant sample size is stated with each comparison. Tables that present deterministic values, single fitted projections or descriptive campaign summaries do not necessarily carry an interval. Unless stated otherwise, reported intervals use the ninety-five per cent level.

2.7.2 Hypothesis testing and effect sizes

When two solvers are compared on the same set of instances, the design is paired: each instance yields one result per solver, and the quantity of interest is the per-instance difference. The feasibility comparison reports a bootstrap interval and an approximate paired sign-test value. No omnibus multi-solver rank claim is made because the complete per-instance, all-solver matrix was not retained in the audited export.

For the hardware comparisons, finite-shot proportions are analysed using the tests and intervals stated alongside the results. The five protocol-motivated contrasts form one family, with p-values adjusted using the Holm procedure. A p-value does not show whether a difference

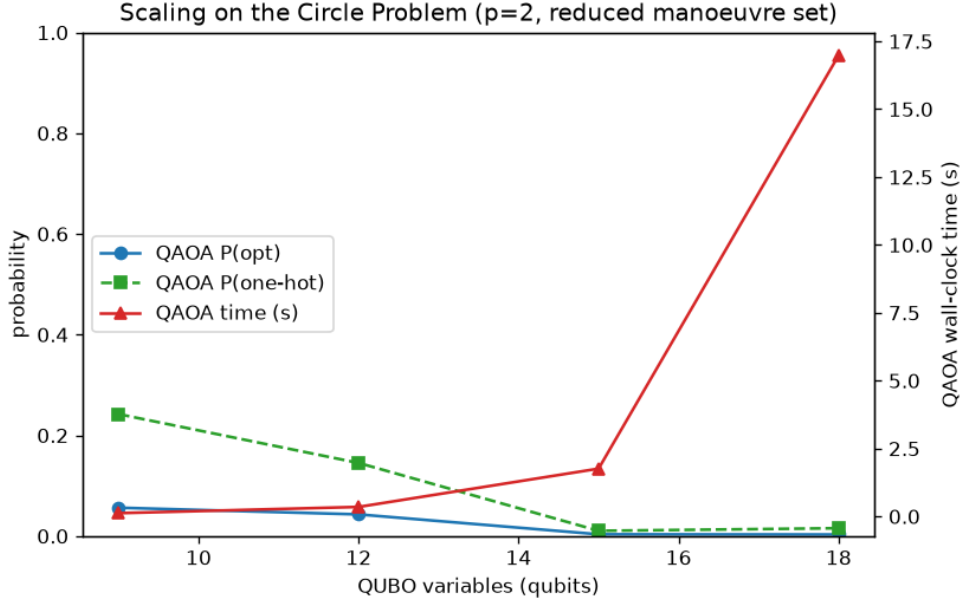


Figure 2.7: Scaling on the Circle Problem: QAOA success probability falls and simulation time rises steeply with qubit count.

is large enough to matter, so differences in percentage points and their intervals are also reported. Significance is interpreted conservatively together with the sample size, effect estimate and available interval.

2.8 Validity, limitations and ethics

Internal validity rests on a shared discrete formulation and exact checks: brute-force enumeration validates the QUBO, and the discrete Gurobi ILP agrees with it wherever both are tractable, including the retained $M = 3$ campaign from CP_3 to CP_6 . The original continuous MINLP is an external formulation check, not the source of like-for-like ground truth. **External validity** is bounded: results hold for small, controlled tactical instances and should not be read as claims about operational-scale air traffic management. **Construct validity** is supported by reporting full feasibility alongside one-hot mass and objective quality, so that these distinct quantities are not conflated. The main **limitations** are the small instance sizes imposed by exact simulation, the idealised noise models relative to real hardware, and the sensitivity of QAOA to classical parameter optimisation. On **ethics and integrity**, the project files are being consolidated in the reproducibility repository at <https://github.com/Louis-Gabin/acrp-quantum>.

2.8.1 Practical scope and resource constraints

This study was conducted as an individual Master’s project in a business-school programme, whereas comparable project work is often distributed across a pair or a larger technical research

team. Its objective is consequently not to extend quantum-control theory or compete with specialised physics laboratories. The intended contribution is an accessible, auditable evaluation of a technically complex emerging technology against a concrete operational-research problem. This positioning shaped the emphasis on transparent decision criteria, reproducible evidence, managerial interpretation, and explicit limits rather than on hardware-level innovation.

Access to quantum hardware imposed a second practical boundary. IBM Quantum was available through an Open Plan allocation rather than a reserved research account, with a limited rolling QPU-time allowance, variable queue conditions, and no control over the processor’s calibration schedule. Exact statevector validation was also bounded by local memory, which restricted exhaustive quantum simulation to small instances. The primary experiment was therefore reduced to nine-qubit CP_3 , with fixed angles optimised locally, 4,096 shots per circuit, and no QPU-based outer optimisation loop. The broader benchmark also remains heading-only: $M = 3$ uses three headings and $M = 5$ uses five, always at nominal speed. It should not be presented as a continuous heading-and-speed reproduction of the richer classical ACRP formulations.

These constraints reduced breadth but do not invalidate the reported within-scope measurements. Several safeguards compensate for the absence of a larger team and dedicated hardware access: two independent QAOA implementations are cross-checked, small QUBOs are exhaustively enumerated, exact objectives are validated against Gurobi, hardware plans are frozen before submission, job identifiers and raw counts are retained, and repeated interleaved circuits measure within-job stability. Conversely, some objectives remain only partially satisfied. Total user wall-clock cost is not available consistently for every solver and optimisation restart, the hardware evidence concerns one backend and one small instance family, and the complete all-solver matrix required for a defensible Friedman–Nemenyi ranking was not retained in the audited export. These are limits on the claims that may be made, not reasons to reinterpret missing evidence as positive performance.

Chapter 3

Results

3.1 Experimental campaign

This chapter reports several linked but distinct campaigns. The principal $M = 3$ multi-instance benchmark covers CP_3 to CP_6 (nine to eighteen qubits) and two randomised instances at each size. The more expensive $M = 5$ sensitivity benchmark is limited to CP_3 , CP_4 and one random instance at each of those sizes. Separate files contain the preliminary depth calibration, the CP_3 – CP_4 depolarising-noise comparison and the three CP_3 IBM jobs. Exact solvers operate once per instance; stochastic results use the seed, restart or shot structure documented for their specific campaign. Confidence intervals are reported where repeated observations or finite-shot counts support them, rather than assumed for every table.

3.2 Validation of the encoding

The first result is that the discrete Gurobi ILP, brute-force enumeration and QUBO use the same bounded heading-selection model. On every shared instance checked by both exact routes, the brute-force QUBO optimum is feasible and numerically equal to the discrete Gurobi optimum. Table 3.1 reports the twelve-instance $M = 3$ overnight benchmark, including CP_6 and two random instances at each size. The circle instances follow the value $n - 1$, whereas the random instances do not. This validates the like-for-like benchmark that follows; it does not establish equivalence with the supervisor’s continuous MINLP or with richer heading-and-speed formulations in the literature.

The circle values have a direct structural explanation. In CP_N , all aircraft start at a common radius, travel at a common speed, and initially head towards the centre. Any two aircraft retaining the nominal heading reach the centre simultaneously and violate separation, so at most one aircraft can remain nominal and $\text{opt}(CP_N) \geq N - 1$. Conversely, retaining one aircraft and

Table 3.1: Exact objective values in the twelve-instance, $M = 3$ benchmark. The objective is the number of manoeuvring aircraft.

Instance	CP_3	CP_4	CP_5	CP_6	R3.1	R3.2	R4.1	R4.2	R5.1	R5.2	R6.1	R6.2
Optimum	2	3	4	5	1	2	2	1	2	3	3	3

applying the same admissible non-zero turn to the other $N - 1$ aircraft constructs a solution with $N - 1$ manoeuvres wherever that assignment remains feasible. Exhaustive one-hot enumeration verifies equality for $3 \leq N \leq 10$ and finds exactly $2N$ optimal assignments at every tested size. This is a proved lower bound plus a constructive upper bound, with equality verified over a bounded domain, not a general theorem for every N .

The structure also limits the circle family as a benchmark. Its optimum can be predicted from its size over the verified feasible range, so it cannot alone discriminate solver quality. The random instances carry that role because their optima do not follow $n - 1$. The circle family is retained for controlled scaling and comparability. Its equality also has a feasibility qualification: under the fixed $\pm 30^\circ$ manoeuvre envelope, CP_{50} is infeasible at both $M = 3$ and $M = 5$, so its optimum is undefined rather than equal to 49.

3.3 Solver comparison on shared instances

The four solver roles are interpreted against the discrete benchmark. The discrete Gurobi ILP and brute-force enumeration establish ground truth; simulated annealing and QAOA operate on the same QUBO. The exact methods agree on every shared instance. The discrete Gurobi ILP also resolves the tested circle-feasibility ladder up to $n = 50$ in less than 40 milliseconds, including certifying the infeasibility of CP_{50} . This does not contradict the twenty-to-forty-aircraft frontier discussed in Chapter 1, which concerns richer continuous tactical formulations with a larger decision space. It shows instead that the deliberately simplified benchmark used here does not place its classical reference under computational pressure. The experiments therefore cannot support a speed or solution-quality advantage claim for QAOA.

The stochastic comparison reaches the same conclusion from a different direction. Simulated annealing returns a higher feasibility rate than QAOA by an estimated 0.52 on average, with a 95% bootstrap interval of [0.32, 0.69] and an approximate paired sign-test value of $p = 0.04$.

These outcomes separate executability from competitiveness. QAOA does structure probability mass around feasible and optimal states, as the depth and hardware sections show, but the structured exact baseline solves the same small benchmark more reliably and in under 40 milliseconds. A complete total-user-time comparison between simulated annealing and QAOA

Table 3.2: Functional comparison of the solver set on the shared benchmark.

Solver	Empirical role	Observed result
Gurobi (discrete ILP)	Exact classical baseline	Exact status and objective in under 40 ms throughout the tested range to $n = 50$
Brute-force QUBO	Encoding cross-check	Agrees with Gurobi wherever exhaustive enumeration is tractable
Simulated annealing	Classical twin on the same QUBO	Higher feasibility than QAOA by 0.52 on average, 95% bootstrap CI [0.32, 0.69]
QAOA	Quantum variational sampler	Non-trivial optimum probability on small instances, but lower feasibility and positive energy or manoeuvre gaps in the principal paired tests

is not available for every configuration, so no broader runtime ranking is inferred. The quantum contribution of this study is therefore methodological: a reproducible ACRP-to-QUBO-to-QPU pipeline and a controlled measurement of where depth-related gains are lost, rather than evidence of practical superiority.

3.4 QAOA behaviour: depth and noise

Two configuration effects were examined under separate protocols. First, the single-seed CP_3 calibration shows a strong overall depth gain when 120 optimiser restarts are used, without implying monotonic improvement for every intermediate depth or every optimiser budget. Second, the specified depolarising model lowers success probability and full feasibility and raises expected energy at every tested depth on CP_3 and CP_4 . This is a controlled model sensitivity result, not a direct estimate of all hardware noise.

Table 3.3: Ideal statevector $P(\text{opt})$ across depth for $M = 3$, using ten seeds and forty-eight optimiser restarts. A dash denotes a configuration not retained for reporting.

Aircraft	Qubits	$p = 1$	$p = 2$	$p = 3$	$p = 4$
3	9	2.50%	3.80%	6.10%	7.80%
4	12	0.80%	1.60%	2.40%	3.10%
5	15	0.20%	0.50%	1.00%	1.40%
6	18	0.08%	0.31%	0.23%	–

The depth gain remains positive in the pooled bootstrap analysis, with an estimated increase in $P(\text{opt})$ of 0.029 and a 95% interval of [0.011, 0.052]. Table 3.3 also shows why this statement needs an optimisation qualification. The CP_6 result falls from 0.31% at $p = 2$ to 0.23% at $p = 3$. A deeper QAOA circuit can reproduce a shallower one by adding a null layer, so its globally

minimum expected energy cannot increase. However, the probability of sampling an optimal solution is a different quantity: minimising expected energy does not guarantee a higher $P(\text{opt})$. This decline alone therefore does not prove that the deeper circuit was under-optimised.

The probabilities in Table 3.3 should not be merged numerically with the larger ideal probabilities reported later for the hardware campaigns. This table summarises the multi-seed benchmark under its fixed campaign protocol, whereas Section 3.6 reports exact statevector probabilities at single frozen angle vectors selected for specific hardware submissions. The two analyses answer different questions: one characterises variability under a benchmark protocol, while the other provides a matched ideal reference for the circuit actually sent to the QPU. Their difference illustrates the sensitivity of QAOA to angle selection and optimiser budget; it is not an estimate of hardware improvement.

A secondary but honest finding emerges at twelve qubits. Even with twenty-four restarts the ideal CP_4 curve is not monotone in p , which highlights sensitivity to the variational search and to the distinction between expected energy and optimum probability. Non-monotonic optimum probability alone cannot identify the binding constraint. This effect is reported as a result in its own right rather than smoothed away, and it motivates the restart-budget recommendation of Section 2.4.2.

3.5 Scaling behaviour

Across the Circle Problem ladder, QAOA success probability declines as the state space expands, while exact statevector cost rises exponentially, as shown in Figure 2.7. Width and circuit length must be distinguished. Increasing n raises the qubit count $V = nM$; increasing QAOA depth p repeats the cost and mixer layers. A processor can therefore expose enough physical qubits for an instance while still being unable to execute its routed circuit with useful fidelity.

The projection is deliberately used as an exclusion argument, not as a performance prediction. For CP_{25} , the probability of satisfying one-hot selection under uniform unconstrained-bitstring sampling, $(3/8)^{25} = 2.243 \times 10^{-11}$ implies essentially zero feasible observations in 4,096 shots before conflict constraints are even considered. No feasible assignment appeared in a separate 20,000-sample uniform one-hot check. At CP_{50} the fixed manoeuvre envelope is itself infeasible, and the projected retained fidelity is 4.79×10^{-12} . Submitting either instance would therefore consume QPU allocation without a credible chance of producing an interpretable sample.

The scaling evidence is consequently bounded in two directions. Exact simulation measures the decline only to eighteen qubits, and Table 3.4 extrapolates circuit cost rather than hardware performance. Within that boundary, the result is unambiguous: the classical exact baseline re-

Table 3.4: Width and indicative routed-circuit cost for the $M = 3$ circle family. Two-qubit counts are $p = 1$ estimates beyond the measured CP_3 point. F is a calibration-based retained-fidelity projection and is not an observed algorithmic fidelity.

Aircraft	Logical qubits	QUBO couplings	Estimated 2Q gates	Projected F
3	9	18	63	0.840
4	12	30	105	0.761
6	18	57	200	0.612
10	30	155	542	0.294
12	36	210	735	0.196
25	75	925	3,238	1.24×10^{-3}
50	150	3,725	13,038	4.79×10^{-12}

mains below 40 milliseconds, while statevector simulation and routed quantum circuits become rapidly more expensive. The present benchmark is a controlled proof of concept, not a crossover experiment.

3.6 Execution on real quantum hardware (IBM Quantum)

The simulator results isolate algorithmic behaviour from device physics. Three hardware jobs were therefore executed on the 156-qubit `ibm_kingston` processor. In every case, the classical optimisation loop remained local: only fixed, measured circuits were submitted through SamplerV2. This design preserves the limited Open Plan allocation and ensures that the hardware observations concern sampling and circuit execution rather than a noisy variational training loop. All confidence intervals in this section are Wilson intervals for finite-shot uncertainty. They do not incorporate calibration uncertainty or day-to-day backend variation.

Initial executability check. The first job, `da4mm2maa69c739jvum0`, submitted one nine-qubit CP_3 circuit at $p = 1$ with 4,096 shots. After transpilation, the circuit had ISA depth 150 and 63 two-qubit gates. The device returned 545 feasible samples and 269 optimal samples, corresponding to a feasibility rate of 13.31% and an unconditional $P(\text{opt})$ of 6.57%. The latter is the principal metric because it does not condition away infeasible observations. At the same angles, the ideal statevector values were 15.45% and 6.77%, respectively. Table 3.5 shows that the hardware distribution retained the optimum signal of this shallow circuit while remaining far above uniform sampling over the 512 unconstrained bitstrings. The retained JSON does not record a trustworthy job-level QPU duration or physical layout for this first check, so neither is reconstructed. This result establishes end-to-end executability, not practical quantum advantage.

Table 3.5: Initial CP_3 , $p = 1$ hardware check on `ibm_kingston`, job `da4mm2maa69c739jvum0`.

Metric	Uniform bitstrings	Ideal, same angles	Hardware [95% CI]
Feasibility rate	1.56%	15.45%	13.31% [12.30, 14.38]
$P(\text{opt})$	1.17%	6.77%	6.57% [5.85, 7.37]

A more informative random comparator enforces one-hot selection directly: choose one of the three headings independently and uniformly for each of the three aircraft. This samples the 27 one-hot assignments without using an optimum oracle. Exact enumeration gives eight conflict-free assignments and six optimal assignments, hence feasibility $8/27 = 29.63\%$ and $P(\text{opt}) = 6/27 = 22.22\%$. Uniform sampling over all 512 bitstrings instead gives $8/512 = 1.56\%$ and $6/512 = 1.17\%$. Even the best hardware circuit reported below, $p = 3_{\text{true}}$, remains below the one-hot comparator on both metrics (21.31% feasibility and 15.11% optimality). Exceeding the unconstrained-bitstring baseline therefore demonstrates a signal, not an advantage over simple classical sampling.

First depth series. The second job, `da8aqcse74ec73aie27g`, was submitted on 27 August 2026 and placed five circuits in one `SamplerV2` job, with 4,096 shots per circuit and common initial physical layout [107, 108, 106, 105, 97, 109, 117, 87, 110]. IBM reported 8 seconds of QPU usage for the complete job. The angle plan was frozen before connection to IBM and selected exclusively by expected QUBO energy. Ideal performance improved strongly with depth, reaching 48.04% ideal $P(\text{opt})$ at $p = 6$. The transfer to hardware was not monotone: the best observed hardware value was 10.25% at $p = 2$, while $p = 6$ returned 8.69%. The circuit-configuration effect across the five submitted circuits was statistically detectable, $\chi^2(4) = 67.43$, $p = 7.9 \times 10^{-14}$, but depth was ordered and therefore partly confounded with possible within-job drift; this omnibus result does not identify a pure depth effect.

Table 3.6: First CP_3 hardware depth series on `ibm_kingston`, job `da8aqcse74ec73aie27g`. Each circuit used 4,096 shots.

p	ISA depth	2Q gates	Ideal feasibility	Ideal $P(\text{opt})$	Hardware feasibility	Hardware $P(\text{opt})$ [95% CI]
1	119	54	12.19%	9.45%	10.06%	7.54% [6.77, 8.39]
2	246	139	26.33%	21.88%	12.74%	10.25% [9.36, 11.22]
3	403	241	26.34%	21.88%	7.32%	5.52% [4.86, 6.26]
4	528	317	41.12%	27.41%	14.04%	8.69% [7.87, 9.59]
6	840	513	65.24%	48.04%	10.35%	8.69% [7.87, 9.59]

This first series contained two interpretive limitations. The circuits were executed sequentially, so depth was partially confounded with possible within-job drift. In addition, the submitted $p = 3$ parameters reproduced the first two $p = 2$ layers and added a near-zero third layer. Its

ideal $P(\text{opt})$ was therefore virtually identical to that of $p = 2$, although transpilation added 102 two-qubit gates. This point is not treated as an independently optimised depth-three solution. It is retained as an informative null-layer control.

Interleaved confirmation and null-layer control. The third job, daa99rhqtnsc73d2d110, addressed both limitations. A new $p = 3$ circuit was selected from 96 independent COBYLA starts on expected QUBO energy; this is the best candidate found under that search budget, not a proof of globally optimal parameters. The resulting ideal energy was 8.3602, compared with 10.3512 for the historical null-layer circuit. Eight circuits were then submitted in the interleaved order $p = 2_A$, $p = 6$, $p = 3_{\text{true}}$, $p = 2_B$, $p = 4$, $p = 3_{\text{null}}$, $p = 1$, and $p = 2_C$. All used the physical layout [79, 93, 92, 73, 94, 91, 74, 95, 90] and 4,096 shots. The three $p = 2$ controls were exact copies of the same ISA circuit. IBM reported 11 seconds of QPU usage for the complete job. The frozen plan records Python 3.12.13, NumPy 2.5.2, SciPy 1.18.0, Qiskit 2.5.2 and Qiskit IBM Runtime 0.49.0.

Table 3.7: Interleaved CP_3 confirmation campaign on ibm_kingston, job daa99rhqtnsc73d2d110. Confidence intervals quantify binomial shot uncertainty only.

Circuit	ISA depth	2Q gates	Ideal $P(\text{opt})$	Hardware $P(\text{opt})$ [95% CI]	Feasibility [95% CI]
$p = 1$	122	56	9.45%	7.79% [7.01, 8.65]	10.28% [9.39, 11.25]
$p = 2_A$	250	141	21.88%	13.01% [12.02, 14.08]	15.04% [13.98, 16.17]
$p = 2_B$	250	141	21.88%	12.62% [11.64, 13.67]	14.97% [13.91, 16.09]
$p = 2_C$	250	141	21.88%	11.74% [10.79, 12.77]	14.06% [13.03, 15.16]
$p = 3_{\text{null}}$	405	243	21.88%	10.77% [9.85, 11.75]	13.21% [12.21, 14.28]
$p = 3_{\text{true}}$	404	243	30.18%	15.11% [14.05, 16.24]	21.31% [20.09, 22.59]
$p = 4$	606	339	27.41%	10.23% [9.34, 11.20]	14.92% [13.86, 16.04]
$p = 6$	831	513	48.04%	11.69% [10.75, 12.71]	15.38% [14.31, 16.52]

The three repeated $p = 2$ controls returned $P(\text{opt})$ values of 13.01%, 12.62%, and 11.74%. A 3×2 homogeneity test did not detect a difference among them, $\chi^2(2) = 3.18$, $p = 0.204$. The corresponding feasibility test was also non-significant, $\chi^2(2) = 1.93$, $p = 0.380$. These checks do not establish perfect stationarity, but they provide no evidence of a large within-job drift at the available shot count. Given the absence of detected heterogeneity, the pooled $p = 2$ value was 12.459%.

At matched transpiled cost, $p = 3_{\text{true}}$ and $p = 3_{\text{null}}$ both used 243 two-qubit gates and had ISA depths 404 and 405. The independently optimised circuit achieved 15.11% optimal samples, compared with 10.77% for the null-layer control, a difference of 4.35 percentage points. Conversely, the null-layer circuit was 1.69 points below the pooled $p = 2$ controls despite virtually unchanged ideal $P(\text{opt})$. This comparison separates two effects within one job: an informative third layer improved the observed result, while a near-null third layer increased physical

Table 3.8: Protocol-motivated contrasts for hardware $P(\text{opt})$ in job daa99rhqtnsc73d2d110. Holm adjustment is applied across the five contrasts.

Contrast	Difference (pp)	95% score CI (pp)	z	Holm-adjusted p
$p = 3_{\text{true}}$ vs $p = 3_{\text{null}}$	4.35	[2.89, 5.80]	5.86	1.86×10^{-8}
$p = 3_{\text{null}}$ vs pooled $p = 2$	-1.69	[-2.78, -0.55]	-2.88	0.004
$p = 3_{\text{true}}$ vs pooled $p = 2$	2.65	[1.43, 3.92]	4.35	2.66×10^{-5}
$p = 4$ vs $p = 3_{\text{true}}$	-4.88	[-6.32, -3.45]	-6.64	1.54×10^{-10}
$p = 6$ vs $p = 3_{\text{true}}$	-3.42	[-4.89, -1.94]	-4.54	1.69×10^{-5}

cost without adding ideal solution probability.

The $p = 4$ circuit was already weaker than $p = 3_{\text{true}}$ in the ideal simulation (27.41% versus 30.18% optimum probability), and was also weaker on hardware. The $p = 6$ circuit, by contrast, had an ideal advantage that did not transfer to hardware. In particular, $p = 6$ reached 48.04% ideal $P(\text{opt})$ but only 11.69% on hardware, below the 15.11% observed at $p = 3_{\text{true}}$. The appropriate conclusion is therefore bounded. Additional depth can improve hardware performance when it adds useful algorithmic structure, as observed at $p = 3$, but further circuit growth did not transfer the larger ideal gains on this backend and calibration window. This is a within-job result on a nine-qubit instance, not a backend-independent performance guarantee.

3.7 Synthesis: answering the research question

Taken together, the results answer the research question in proof-of-concept terms. The QUBO resolves the same bounded heading-selection benchmark as the discrete Gurobi ILP, and the exact methods agree on the optimum wherever both can be applied. The supervisor’s continuous MINLP supports the modelling context but is not part of this strict equivalence claim. The stochastic comparison nevertheless favours the classical twin in feasibility, and the discrete Gurobi ILP resolves the tested ladder in under 40 milliseconds. These facts preclude a claim of practical quantum advantage. Because total user wall-clock cost was not retained consistently for every stochastic configuration, the efficiency conclusion is restricted to the measured classical, circuit, sampling, and QPU costs rather than presented as a complete end-to-end runtime ranking.

The quantum result remains informative. Fixed QAOA circuits sample optimal CP_3 solutions with non-trivial probability on a real QPU, and the interleaved confirmation campaign shows that depth is not intrinsically harmful: an independently optimised third layer improves both ideal and hardware performance. The hardware benefit does not persist at $p = 4$ or $p = 6$. At $p = 4$, ideal performance is already below the independently optimised $p = 3$ circuit; at $p = 6$, a larger ideal gain fails to transfer to hardware. Useful algorithmic structure must therefore compensate for additional gate cost before deeper QAOA can improve the observed distribution. Scaling estimates then show why additional physical qubits alone do not close the gap: routed interaction cost and the exponentially shrinking feasible mass make larger submissions

uninterpretable well before the nominal device width is exhausted.

The overall contribution is accordingly methodological and empirical. It provides a reproducible path from a simplified tactical ACRP benchmark to a validated QUBO, exact and heuristic classical baselines, ideal and noisy QAOA, and controlled hardware sampling. Within the tested range, the evidence supports executability and a measurable depth-noise trade-off, but not competitiveness with structured classical optimisation or generalisation to continuous, operational air traffic scales. The study narrows the problem-fidelity, benchmarking-rigour, and reproducibility gaps identified in Chapter 1 without claiming to close them completely.

Conclusion

This thesis asked a practical question: on a small and controlled Aircraft Conflict Resolution Problem, can a quantum method approach a strong classical method, and what does it cost to obtain the result? The purpose was not to predict the entire future of quantum computing in aviation. It was to test what could be demonstrated with the tools, hardware access and computing resources available during this project.

The answer for the tested problem is that the classical approach remains clearly preferable. The discrete Gurobi model finds and certifies the correct solutions almost immediately on the simplified instances. Simulated annealing, a classical method applied to the same QUBO as the quantum algorithm, also returns valid solutions more often than QAOA. The quantum experiments therefore do not show a practical advantage in speed, reliability or solution quality. This conclusion applies to the benchmark studied here. It is not a statement that quantum optimisation can never become useful for aviation.

An important first step was to ensure that the algorithms compared directly were solving the same problem. The tactical conflict rules were translated into a QUBO, which is a binary mathematical representation compatible with several classical and quantum solvers. For every small instance checked by both exact discrete routes, the best QUBO solution matched the discrete Gurobi ILP. The original continuous MINLP was kept as a contextual cross-check because it permits heading and speed values outside the discrete grid. Separating these two roles prevents agreement in objective value from being mistaken for identical decision domains.

The quantum part of the project was nevertheless successful as a proof of execution. The complete chain worked: an aircraft-conflict instance was converted into a QUBO, the QUBO was converted into a quantum circuit, the circuit was checked locally and then executed on the IBM processor `ibm_kingston`. The processor returned optimal solutions with a probability well above uniform sampling over the 512 unconstrained bitstrings, but below uniform sampling over the 27 one-hot assignments. This shows that the model was not merely theoretical. It also shows why successful execution and practical usefulness must be kept separate. A quantum processor can return the correct answer without being the best way to obtain it.

The clearest lesson from the IBM experiments concerns circuit depth. In QAOA, increasing the depth means adding layers of quantum operations. In an ideal simulation, these layers

can make the algorithm better because they give it more freedom to shape the final probability distribution. On a real processor, however, every layer also adds gates, execution time and opportunities for errors. More depth therefore creates both a possible benefit and an additional risk.

The controlled hardware campaign made this trade-off visible. A three-layer circuit selected independently from 96 optimiser starts produced an optimal solution in 15.11% of the measurements. A circuit of almost the same physical length, but with a third layer that contributed almost nothing to the ideal algorithm, reached only 10.77%. Repeating the same two-layer circuit at different points in the job helped check that this difference was not simply caused by a major change in the processor during the experiment. The lesson is simple: a longer quantum circuit can help when the added operations contain useful algorithmic information. Making a circuit longer without adding such information only increases its exposure to noise.

Going deeper did not continue to improve the hardware result. Some of the longest circuits looked much better in the ideal simulator but lost most of that improvement when executed on the processor. This explains why the number of available qubits is not enough to judge whether a quantum computer is ready for a problem. The machine may contain enough qubits to represent the data, while the circuit required to connect and manipulate them is already too long and fragile to preserve the expected result.

These experiments had to be reduced considerably compared with the richer aircraft-conflict models reviewed in Chapter 1. Exact quantum simulation stores the full quantum state in memory, and the required memory grows very quickly with each additional qubit. This restricted the instances that could be verified locally. Access to IBM Quantum was also provided through the Open Plan, with limited QPU time, variable queues and no control over the processor's calibration. The main hardware experiment was consequently conducted on CP_3 , a three-aircraft Circle Problem represented with nine qubits. The variational parameters were optimised locally, and only fixed circuits were sent to IBM, with 4,096 repeated measurements per circuit.

The problem itself was also simplified. Vertical motion was excluded first, and speed was then fixed at its nominal value because adding another control dimension would multiply the binary and circuit resources. The principal benchmark used three heading choices, $\{-30^\circ, 0, +30^\circ\}$, while smaller sensitivity tests used five, $\{-30^\circ, -15^\circ, 0, +15^\circ, +30^\circ\}$. The thesis therefore does not reproduce continuous air traffic control or claim that a three-aircraft experiment represents an operational sector. The simplification was a deliberate way to obtain a problem small enough for quantum validation while preserving the central decision: selecting a compatible set of headings that resolves every conflict and changes as few trajectories as possible.

The scope also reflects the conditions of the research. This was an individual Master's thesis within a business-school programme, carried out without a dedicated quantum research team or reserved hardware access. The technical ambition was adapted accordingly. The priority be-

came to build an understandable and auditable comparison, rather than to claim an advance in quantum physics. This required substantial work beyond running a circuit: checking the problem through exact enumeration, comparing it with Gurobi, maintaining independent NumPy and Qiskit implementations, testing many starting points for the QAOA parameters, freezing the hardware plan before submission and retaining job identifiers and measurement counts.

One apparent failure became particularly useful. The first three-layer circuit had effectively copied the two-layer solution and added a layer that did almost nothing. Instead of discarding it, the project retained it as a control. Comparing this near-null layer with a newly optimised third layer made it possible to separate the benefit of useful algorithmic work from the cost of simply extending the circuit. Repeating the same two-layer circuit throughout the final job also provided a basic check of hardware stability. These decisions illustrate the broader value of transparent experimental work: an imperfect result can strengthen the analysis when it is documented and used as a control rather than hidden.

Several limitations remain. The hardware evidence comes from one small problem family and one IBM processor. Repeating a circuit thousands of times measures sampling uncertainty, but it does not capture every change in hardware quality between different days. QAOA results also depend strongly on the classical procedure used to choose the circuit parameters. In addition, the total time experienced by the user was not recorded consistently for every solver and every optimisation restart. A complete multi-solver data table required for one planned statistical ranking was also unavailable. The thesis therefore reports the comparisons that can be supported and does not reconstruct missing results.

For a manager or organisation, the main implication reaches beyond this particular algorithm. A working prototype, a competitive solution and an operationally valuable tool are three different levels of maturity. Quantum experimentation can create value through learning, skills development and a better understanding of future constraints, even when it is not ready for deployment. However, a decision to adopt the technology should consider the strength of the classical alternative, the preparation and tuning effort, hardware access, reproducibility and the complete cost of obtaining a usable answer. On the evidence produced here, classical optimisation remains the appropriate computational reference for this bounded ACRP, while quantum computing remains an exploratory capability.

Future work should first make the comparison more complete, not simply larger. All optimisation, waiting, compilation and processing times should be recorded automatically. The same frozen circuits should be repeated on different days and, if access permits, on different processors. Techniques intended to reduce errors should be compared with an unchanged control circuit. The quantum encoding could then be redesigned to enforce valid manoeuvre choices with fewer gates. Only after these improvements would it be useful to move towards more aircraft, richer manoeuvre sets and less symmetric situations, while continuing to compare against strong classical methods.

References

- Azad, U., Behera, B. K., Ahmed, E. A., Panigrahi, P. K., & Farouk, A. (2023). Solving vehicle routing problem using quantum approximate optimization algorithm. *IEEE Transactions on Intelligent Transportation Systems*, 24(7), 7564–7573. <https://doi.org/10.1109/TITS.2022.3172241>
- Cafieri, S., & Rey, D. (2017). Maximizing the number of conflict-free aircraft using mixed-integer nonlinear programming. *Computers & Operations Research*, 80, 147–158. <https://doi.org/10.1016/j.cor.2016.12.002>
- Cai, J., & Zhang, N. (2018). Mixed integer nonlinear programming for three-dimensional aircraft conflict avoidance. *PeerJ Preprints*, 6, e27410v1. <https://doi.org/10.7287/peerj.preprints.27410v1>
- Chen, Y., Koch, T., Peng, H., & Zhang, H. (2025). Benchmarking of quantum and classical computing in large-scale dynamic portfolio optimization under market frictions. *arXiv preprint arXiv:2502.05226*. <https://doi.org/10.48550/arXiv.2502.05226>
- Costa, P. C. S., Morales, M. E. S., An, D., & Sanders, Y. R. (2024). Assessing quantum and classical approaches to combinatorial optimization: Testing quadratic speed-ups for heuristic algorithms. *arXiv preprint arXiv:2412.13035*. <https://doi.org/10.48550/arXiv.2412.13035>
- Dias, F., & Rey, D. (2024). Aircraft conflict resolution with trajectory recovery using mixed-integer programming. *Journal of Global Optimization*, 90, 1031–1067. <https://doi.org/10.1007/s10898-024-01393-1>
- EUROCONTROL. (2024). *EUROCONTROL Aviation Long-Term Outlook: Flights and CO₂ emissions forecast 2024–2050*. <https://www.eurocontrol.int/publication/eurocontrol-forecast-2024-2050>
- Farhi, E., Goldstone, J., & Gutmann, S. (2014). A quantum approximate optimization algorithm. *arXiv preprint arXiv:1411.4028*. <https://doi.org/10.48550/arXiv.1411.4028>
- Fitzek, D., Ghandriz, T., Laine, L., Granath, M., & Frisk Kockum, A. (2024). Applying quan-

- tum approximate optimization to the heterogeneous vehicle routing problem. *Scientific Reports*, 14, 25415. <https://doi.org/10.1038/s41598-024-76967-w>
- Huang, Z., Li, Q., Zhao, J., & Song, M. (2022). Variational quantum algorithm applied to collision avoidance of unmanned aerial vehicles. *Entropy*, 24(11), 1685. <https://doi.org/10.3390/e24111685>
- Koretsky, S., Gokhale, P., Baker, J. M., Viszlai, J., Zheng, H., Gurung, N., Burg, R., Paaso, E. A., Khodaei, A., Eskandarpour, R., & Chong, F. T. (2021). Adapting Quantum Approximation Optimization Algorithm (QAOA) for unit commitment. In *2021 IEEE International Conference on Quantum Computing and Engineering (QCE)* (pp. 181–187). <https://doi.org/10.1109/QCE52317.2021.00035>
- Kuchar, J. K., & Yang, L. C. (2000). A review of conflict detection and resolution modeling methods. *IEEE Transactions on Intelligent Transportation Systems*, 1(4), 179–189. <https://doi.org/10.1109/6979.898217>
- Lewis, M., & Glover, F. (2017). Quadratic unconstrained binary optimization problem preprocessing: Theory and empirical analysis. *Networks*, 70(2), 79–97. <https://doi.org/10.1002/net.21751>
- McDowall, K., Kapourniotis, T., Oliver, C., Lolur, P., & Georgopoulos, K. (2025). Cross-platform benchmarking of near-term quantum optimisation algorithms. *arXiv preprint arXiv:2504.06885* (version 5, revised 2026). <https://doi.org/10.48550/arXiv.2504.06885>
- Pallottino, L., Feron, E. M., & Bicchi, A. (2002). Conflict resolution problems for air traffic management systems solved with mixed integer programming. *IEEE Transactions on Intelligent Transportation Systems*, 3(1), 3–11. <https://doi.org/10.1109/6979.994791>
- Pecyna, T., Kurowski, K., Różycki, R., Waligóra, G., & Węglarz, J. (2024). Quantum variational algorithms for the aircraft deconfliction problem. In *Computational Science – ICCS 2024* (Lecture Notes in Computer Science, Vol. 14837, pp. 307–320). Springer. https://doi.org/10.1007/978-3-031-63778-0_22
- Pecyna, T., & Różycki, R. (2024). Improving quantum optimization algorithms by constraint relaxation. *Applied Sciences*, 14(18), 8099. <https://doi.org/10.3390/app14188099>
- Pelegrín, M., & Cerulli, M. (2023). Aircraft conflict resolution: A benchmark generator. *INFORMS Journal on Computing*, 35(2), 274–285. <https://doi.org/10.1287/ijoc.2022.1265>
- Preskill, J. (2018). Quantum computing in the NISQ era and beyond. *Quantum*, 2, 79. <https://doi.org/10.22331/q-2018-08-06-79>
- Rey, D., & Hijazi, H. (2017). Complex number formulation and convex relaxations for aircraft

- conflict resolution. In *2017 IEEE 56th Annual Conference on Decision and Control (CDC)* (pp. 88–93). <https://doi.org/10.1109/CDC.2017.8263648>
- Ruan, Y., Yuan, Z., Xue, X., & Liu, Z. (2023). Quantum approximate optimization for combinatorial problems with constraints. *Information Sciences*, 619, 98–125. <https://doi.org/10.1016/j.ins.2022.11.020>
- Shirai, T., & Togawa, N. (2024). Postprocessing variationally scheduled quantum algorithm for constrained combinatorial optimization problems. *IEEE Transactions on Quantum Engineering*, 5, 1–14. <https://doi.org/10.1109/TQE.2024.3376721>
- Stollenwerk, T., O’Gorman, B., Venturelli, D., Mandrà, S., Rodionova, O., Ng, H., Sridhar, B., Rieffel, E. G., & Biswas, R. (2020). Quantum annealing applied to de-conflicting optimal trajectories for air traffic management. *IEEE Transactions on Intelligent Transportation Systems*, 21(1), 285–297. <https://doi.org/10.1109/TITS.2019.2891235>

Appendix A: Comparative Overview of Quantum ACRP Contributions

Table A.1: Comparative overview of the quantum and quantum-adjacent ACRP contributions used to establish the research gap.

Study	Paradigm / algorithm	Problem actually solved	Max instance scale (hardware)	Constraint handling	Benchmarking rigour	Key limitation
Stollenwerk et al. (2020)	Quantum annealing (D-Wave 2X / 2000Q)	Strategic, delay-only deconfliction; trajectories fixed	~50 flights / 104 conflicts (2X); 64 / 261 (2000Q), after conflict-graph decomposition	Penalty-only QUBO (one-hot + delay cost + conflict)	Exact Max-SAT and classical Isoenergetic Cluster Method (ICM) references on the delay QUBO	Not tactical heading/speed; embedding & coefficient-precision limits
Pecyna et al. (2024)	Gate-based QAOA / QAOAnsatz (IBM Torino, 133 qubits)	Tactical manoeuvre selection	3–5 aircraft, 2–5 manoeuvres (9–25 qubits)	Penalty (QAOA) vs constraint-preserving mixer (QAOAnsatz)	Feasible-solution probability on noisy simulator / hardware; no classical optimum baseline	Very small scale; more manoeuvres harder than more aircraft; NISQ noise
Pecyna and Różycki (2024)	Gate-based, NOT-XOR constraint relaxation	Tactical manoeuvre selection with relaxed feasibility	9 artificial 3×5 instances	Constraint relaxation + classical post-processing	9/9 feasible vs 7/9 for one-hot, under a relaxed criterion	Small artificial instances; feasibility metric only
Huang et al. (2022)	Gate-based VQE / QAOA + CVaR	UAV collision avoidance (adjacent domain)	4 small test instances	QUBO / Ising penalty + CVaR objective	>90% feasibility under tuned parameters; no classical baseline	Not crewed ACRP; small; feasibility-focused

Appendix B: Formal Derivations, Notation and Algorithms

B.1 One-hot expansion

For binary variables $x_{i,m}$, using $x_{i,m}^2 = x_{i,m}$,

$$\begin{aligned} \left(\sum_m x_{i,m} - 1 \right)^2 &= \sum_m x_{i,m}^2 + 2 \sum_{m < m'} x_{i,m} x_{i,m'} - 2 \sum_m x_{i,m} + 1 \\ &= - \sum_m x_{i,m} + 2 \sum_{m < m'} x_{i,m} x_{i,m'} + 1. \end{aligned}$$

Summed over aircraft and weighted by A , this contributes a diagonal term $-A$ per variable, a polynomial coupling $2A$ per intra-aircraft manoeuvre pair (equivalently, symmetric matrix entries $Q_{kl} = Q_{lk} = A$ in $E = \mathbf{x}^T Q \mathbf{x} + \text{offset}$), and a constant offset An . For a one-hot assignment each aircraft contributes $-A + A = 0$; if it is also conflict-free, its energy equals the objective (the number of manoeuvres) exactly, which is why a feasible optimum reads as a small integer.

B.2 Ising mapping for QAOA

Substituting $x_k = (1 - Z_k)/2$ into the symmetric QUBO yields

$$E = \sum_k h_k Z_k + \sum_{k < l} J_{kl} Z_k Z_l + \text{const},$$

with the linear and quadratic coefficients obtained by collecting the diagonal and off-diagonal QUBO entries. The cost unitary is $U_C(\gamma) = e^{-i\gamma E}$ (diagonal in the computational basis) and the mixer is $U_B(\beta) = e^{-i\beta \sum_k X_k}$.

B.3 Solver pseudocode

For symmetric Q and $E(\mathbf{x}) = \mathbf{x}^T Q \mathbf{x} + \text{offset}$, flipping bit k changes the energy by

$$\Delta E = 2(1 - 2x_k)(Q\mathbf{x})_k + Q_{kk}.$$

Here $Q\mathbf{x}$ is evaluated before the flip; the constant offset cancels. A triangular QUBO representation must first be symmetrised, preserving its quadratic polynomial.

Simulated annealing (single-spin-flip Metropolis)

```
input: QUBO Q, offset, num_reads, num_sweeps, beta schedule
for each read:
  x <- random binary vector
  for beta in geometric(beta_min, beta_max, num_sweeps):
    for k in random permutation of variables:
      delta <- 2 (1 - 2 x_k) (Q x)_k + Q_kk
      if delta <= 0 or rand() < exp(-beta delta): flip x_k, update Q x
    track best feasible energy
return best over all reads
```

QAOA (statevector)

```
input: cost spectrum E over  $2^V$  states, depth p
|psi> <- uniform superposition
optimise (gamma, beta) to minimise expectation(E, psi):
  for each layer t: apply diagonal phase exp(-i gamma_t E);
                  apply mixer exp(-i beta_t sum X)
compute the full probability distribution
return expected energy, P(opt), feasibility mass
and the best sampled state
```

B.4 Notation

Table B.1: Notation used in the QUBO and QAOA formulations.

Symbol	Meaning
n	Number of aircraft
d	Minimum horizontal separation
M	Number of candidate manoeuvres per aircraft
$V = nM$	Number of QUBO variables (qubits)
q_i, δ_i	Speed multiplier and heading deviation of aircraft i
$x_{i,m}$	Binary: aircraft i selects manoeuvre m
$c_{i,m}$	Manoeuvre cost (0 nominal, 1 otherwise)
A, B, C	One-hot, separation and objective weights
K_{ij}	Conflicting manoeuvre pairs for aircraft pair (i, j)
p	QAOA depth (number of layers)

B.5 Reproducibility checklist

- Fixed seeds for the stochastic campaigns where seed-level records were retained.
- Penalty weights reported per instance ($C = 1$ and $A = B = 2(n + 1)$ unless swept).
- Deterministic instance generation from documented parameters.
- Two independent QAOA implementations cross-checked on expected energy.
- Separate documented commands for benchmark, depth, noise, statistics, hardware submission and retrieval; no universal one-command claim.
- Frozen IBM plans and manifests retain the exact final-job software versions, layout, shot count and job identifier.
- Project repository for code and retained research files: <https://github.com/Louis-Gabin/acrp-quantum>.