



Kuberina: Maritime Stowage-Inspired Combinatorial Optimization for Pre-deployment Scheduling in Heterogeneous Kubernetes Clusters

Authors: Dinh Tan Dung (ORCID: <https://orcid.org/0009-0003-1374-7525>)

Affiliation: Independent Researcher, Ho Chi Minh City, Vietnam

Date: 26th July 2026

Abstract

The default Kubernetes scheduler (`kube-scheduler`) makes millisecond-latency placement decisions using a first-come, first-served heuristic optimized for homogeneous, stateless microservices. As clusters become increasingly heterogeneous — incorporating GPUs, TPUs, and memory-optimized nodes — this reactive approach produces severe resource fragmentation, with industry analyses consistently reporting 30–40% average CPU utilization across cloud environments. Existing solutions such as Volcano, Kueue, and Google Autopilot address scheduling fairness or vertical scaling but do not solve the underlying combinatorial packing problem offline. We present **Kuberina**, an offline, pre-deployment CLI engine that reformulates Kubernetes pod scheduling as a Multi-Dimensional Bin Packing Problem (MDBP) — drawing a structural isomorphism from the Container Stowage Planning Problem (CSPP) used by mega-vessel shipping lines. Kuberina employs a three-phase hybrid pipeline: (1) Vector Packing First-Fit Decreasing (FFD) warm-start, (2) Genetic Algorithm (GA) optimization with gang-aware repair, and (3) Constraint Satisfaction Problem (CSP) enforcement with Forward Checking — all integrated tightly rather than executed sequentially. On a

synthetic benchmark modeled after the MSC Irina mega-vessel (186 nodes, 2,714 pods, 5,128 affinity constraints), Kuberina achieves 100% scheduling success with zero constraint violations, consolidates workloads onto 152 of 186 nodes (18.3% node reduction), reaches 88.7% average CPU utilization, and produces a mathematically verified feasible solution with approximation ratio $\alpha = 1.34$ relative to the LP lower bound — all computed in under 44 seconds. Monte Carlo testing confirms the result is statistically significant ($p < 10^{-4}$). The output is a declarative YAML blueprint that can be reviewed, version-controlled, and applied via `kubect1` without touching the live cluster.

1. Introduction

The evolution of cloud-native infrastructure has positioned Kubernetes as the de facto standard for orchestrating distributed systems. However, as clusters grow increasingly heterogeneous — integrating specialized hardware such as NVIDIA A100 GPUs, TPUs, and memory-optimized nodes — the default scheduling mechanisms reveal deep limitations. The built-in `kube-scheduler` is designed for latency optimization: it makes dynamic, millisecond-scale placement decisions on a first-come, first-served basis whenever resource gaps appear [1]. While this approach suffices for homogeneous, stateless microservices, it fundamentally fails under the geometric complexity of diverse hardware topologies and advanced AI workloads. The consequence is severe resource fragmentation, with industry analyses consistently reporting that cloud environments operate at only 30–40% CPU utilization [1, 35].

The root cause is temporal: `kube-scheduler` evaluates pods individually as they arrive in the queue. When it detects an available slot on any node, it binds the pod immediately. Although this first-fit heuristic is computationally trivial, it inevitably leads to spatial fragmentation over time. High-density workloads submitted later frequently find the cluster's aggregate resources shattered across partially-filled nodes, rendering them unschedulable despite sufficient total cluster capacity [1]. This fragmentation creates a dependency on cluster autoscalers to continuously provision new infrastructure, inflating cloud costs without proportional workload gains.

To address the limitations of dynamic runtime scheduling, a paradigm shift toward **offline static planning** is necessary. This shift draws a direct architectural lineage from the domain of maritime logistics — specifically, the Container Stowage Planning Problem (CSPP) employed by mega-vessel shipping lines [1]. Modern supercomputers optimize the physical placement of shipping containers on mega-vessels long before vessels arrive at port, evaluating millions of spatial scenarios [1]. By transforming Kubernetes scheduling from a dynamic, opaque process into a pre-deployment CLI engine grounded in mathematical optimization, organizations can solve the Multi-Dimensional Bin Packing Problem (MDBP) offline [1]. This approach produces declarative placement blueprints that completely separate the optimization mathematics from the live `kube-apiserver`, thereby eliminating resource fragmentation without interfering with the running cluster state.

Contributions

This paper makes the following contributions:

1. **(Algorithmic)** A hybrid three-phase pipeline combining FFD warm-start, Genetic Algorithm optimization, and CSP Forward Checking for offline Kubernetes scheduling on heterogeneous clusters, inspired by maritime stowage planning.
2. **(Practical)** A CLI tool that generates pre-deployment blueprints directly applicable via `kubectl apply`, requiring zero interaction with the live cluster — no controllers, no daemons, no API throttling risks.
3. **(Methodological)** A demonstration that the structural isomorphism between container stowage planning and Kubernetes pod scheduling is both valid and productive: every constraint in the maritime domain maps 1:1 to a Kubernetes scheduling primitive, yielding a complete constraint taxonomy.
4. **(Process)** A proposition that the value of offline scheduling optimization lies not only in solution quality, but in producing an **auditable, iterable artifact** — a blueprint computed through combinatorial optimization across thousands of evolutionary generations, with mathematical justification for every placement decision. This replaces scheduling decisions based on architect intuition that cannot be audited, reproduced, or challenged — analogous to how Git transformed code deployment into code review, and Terraform transformed infrastructure provisioning into reviewable plans.

2. Related Work

2.1. Kubernetes Scheduling Optimizers

The Kubernetes ecosystem has produced several scheduling extensions that address specific limitations of `kube-scheduler`. **Volcano** [36] is a batch scheduling system designed for high-performance computing and AI workloads on Kubernetes, providing gang scheduling, fair-share queuing, and preemption policies. **Kueue** [39] manages job queuing with `WorkloadPriorityClass` and `ClusterQueue` abstractions, enforcing all-or-nothing admission semantics. **YuniKorn** [37] offers gang scheduling with hierarchical resource fairness for Spark-on-Kubernetes deployments. **Descheduler** reactively evicts and re-schedules pods to rebalance utilization, but operates post-hoc rather than proactively. **Trimaran** extends the scheduler with real-time load-aware scoring.

All of these tools operate **within the dynamic scheduling paradigm** — they improve scheduling decisions at runtime but remain fundamentally reactive. None of them solve the offline combinatorial optimization problem of finding a globally optimal placement before any pod is deployed. Kuberina occupies a complementary position: it computes the placement plan offline and exports it as a static blueprint, which the dynamic schedulers then execute.

2.2. Bin Packing and Cloud Resource Management

The problem of placing workloads onto servers is a well-studied variant of the Multi-Dimensional Bin Packing Problem (MDBP), known to be NP-hard [Garey & Johnson, 1979]. **Google Borg** [Google, 2015] manages cluster resources at planet-scale using a combination of priority-based preemption and equivalence classes, but its scheduling logic is tightly coupled to Google's internal infrastructure.

Microsoft Tetris applies multi-resource packing heuristics for VM placement in Azure, optimizing for utilization alignment across dimensions. **Alibaba Sigma** uses a two-level scheduling architecture to manage container placement across massive data centers. **SAGE** [20] proposes an optimization model for Kubernetes

deployments but focuses on deployment configuration tuning rather than the combinatorial placement problem.

These systems demonstrate the practical importance of bin packing in production environments. Kuberina extends this line of work by applying a maritime-inspired decomposition specifically designed for heterogeneous Kubernetes clusters with GPU constraints, gang scheduling requirements, and topology-aware affinity rules.

2.3. Maritime Stowage Planning

The Container Stowage Planning Problem (CSPP) is a well-established NP-hard combinatorial optimization problem in operations research [2, 29]. **Pacino et al.** [13, 17] developed fast generation methods using mixed-integer programming for master bay planning and constraint-based approaches for slot planning. **Avriel et al.** established foundational models considering vessel stability, stack weight limits, and port rotation sequences. **Delgado et al.** [6, 14] introduced accurate models incorporating ballast tank optimization for seaworthiness constraints. Recent work has applied genetic algorithms [11, 31, 32, 34], many-objective evolutionary algorithms (NSGA-III) [3], and deep reinforcement learning [2] to various formulations of the stowage problem.

The hierarchical decomposition of CSPP into Master Bay Plan and Slot Plan sub-problems [12] directly informs Kuberina's approach: workloads are first assigned to node pools (bays) and then packed into specific nodes (slots). The constraint taxonomy of maritime stowage — stability, segregation, destination grouping, reefer connectivity — provides a complete 1:1 mapping to Kubernetes scheduling primitives, as we demonstrate in Section 4.1.

2.4. Google Autopilot and the Resource Canal Effect

Google Kubernetes Engine (GKE) Autopilot [23, 24] abstracts node management entirely, provisioning infrastructure dynamically and billing per-pod resource requests rather than per-VM. It employs historical metrics, exponentially-smoothed sliding windows, and reinforcement learning through the Vertical Pod Autoscaler (VPA) to automatically adjust resource allocations and resolve OOM failures [24].

However, Autopilot and Kuberina solve fundamentally different problems along orthogonal axes:

Property	Google Autopilot (Time-Bounded)	Kuberina (Resource-Bounded)
Operating domain	Runtime	Pre-deployment
Reference frame	Temporal axis	Spatial axis
Input data	Historical consumption metrics	Static declarative requests/limits
Error handling	Learn from failure in next cycle	Prune infeasible branches before deployment
Objective	Track actual load, auto-adjust limits	Maximize pod packing within fixed node capacity

Autopilot assumes effectively infinite resources — if a container needs more, the control plane provisions new nodes. This assumption breaks for AI-intensive clusters requiring strictly bounded hardware such as NVIDIA L40S, A100, or T4 GPUs, where physical boundaries are absolute and cannot be abstracted away by software [22].

We propose that the two systems are complementary. Kuberina establishes fixed physical boundaries (the "canal banks") using MDBP optimization, specifying exactly where workload archetypes reside. Dynamic scalers like Autopilot then operate within these boundaries, adjusting resource consumption based on real-time traffic (the "water level"). We term this symbiosis the **Resource Canal Effect**. When Autopilot operates within a Kuberina-defined canal, its reinforcement learning cost function no longer explores blindly — Kuberina has already eliminated the extremes of both overrun and underrun, allowing the RL algorithm to converge orders of magnitude faster [1].

3. Problem Formulation

3.1. Structural Isomorphism: Maritime Stowage and Kubernetes Scheduling

The structural correspondence between stowage planning on mega-vessels such as the MSC Irina and workload scheduling on Kubernetes clusters is functionally **isomorphic**. Both environments represent physically bounded infrastructures in which multi-dimensional cargo must be packed tightly while satisfying a strict set of hard constraints (mandatory) and soft constraints (preferred) [1].

3.1.1. The Scale of Mega-Vessel Stowage

Modern mega-vessels such as the MSC Irina class carry over 24,300 Twenty-foot Equivalent Units (TEU) [7]. At full load, the containers placed end-to-end would stretch 147.5 km, equivalent to the volume of 322 Olympic swimming pools or the weight of 52 Eiffel Towers [8]. Without a standardized unit of measurement and a robust mathematical foundation, unifying these variables into a single stowage plan would be intractable [10].

3.1.2. Physical Dynamics: Stability, Shear Forces, and Bending Moments

Stowage plans are governed by two frequently conflicting objectives: ensuring vessel stability and minimizing unnecessary container relocations (restows) [3]. Stability is defined by hydrostatic principles — particularly the initial metacentric height (GM) — computed from the vessel's heel angle, draft, and trim [3]. Total weight and buoyancy must be distributed precisely to prevent transverse bending moments (torsion) and longitudinal shear forces from endangering the hull [6]. To correct unavoidable load imbalances, engineers use ballast tanks distributed along the hull, pumping water in or out to modify displacement and the longitudinal center of gravity (LCG) [6]. This establishes a baseline weight before any cargo is loaded — a principle that maps directly to cloud computing (see DaemonSet pre-deduction in Section 3.2).

3.1.3. Cargo Classification and Segregation Regulations

Beyond structural stability, the stowage engine must compute positions based on container-specific classifications. Standard containers share space with flat-rack containers (ISO type 1B) for oversized equipment and tank containers (ISO type 1T) for pressurized liquids and gases [15]. Reefer containers require specific slots equipped with electrical grid connections, while hazardous materials (hazmat) must be physically segregated according to strict safety regulations [1].

3.1.4. Hierarchical Decomposition

Because simultaneously optimizing 24,000+ individual units is computationally intractable, maritime researchers decompose the CSPP into two hierarchical phases [3, 12]. The first phase — the **Master Bay Plan Problem** — distributes container groups at a macroscopic level into specific longitudinal sections (bays). The second phase — the **Slot Plan Problem** — assigns individual containers to precise grid coordinates within those bays [12]. This decomposition mirrors Kubernetes scheduling: workloads are first assigned to node pools before being scheduled onto specific CPU sockets within a node.

3.1.5. Constraint Mapping Table

The following table establishes the complete 1:1 constraint mapping that structures Kuberina's pre-deployment scheduling engine:

Maritime Stowage Context	Kubernetes Scheduling Context	Mathematical Technique
Container Dimensions (20ft, 40ft, High-Cube)	Resource Requests/ Limits (CPU, RAM, GPU, VRAM)	Multi-Dimensional Bin Packing (MDBP)
Reefer Containers (require powered slots)	AI Compute Workloads (require NVIDIA A100, T4 GPUs)	CSP Hard Constraints (NodeSelector, NodeAffinity)
Hazmat Segregation	Pod Anti-Affinity / Taints & Tolerations	Conflict Graph & Graph Coloring

Maritime Stowage Context	Kubernetes Scheduling Context	Mathematical Technique
Destination Port Grouping (LIFO rotation)	Pod Affinity / Network Topology (service co-location)	Fitness Function for Soft Constraints
Vessel Trim & Stability (GM height, weight balance)	Resource Utilization Balancing (even node loading)	Variance Minimization
Block Booking / Slot Charter	Gang Scheduling (distributed AI training)	Coupled Variables in CSP, Forward Checking
Lashing & Securing (storm safety)	QoS Classes (Guaranteed vs. Burstable)	Knapsack Problem with Strict Capacity Bounds
Hatch Covers (under-deck / on-deck separation)	Topology Spread Constraints / Availability Zones	Distribution Constraints (min/max per zone)
Ballast Water	DaemonSets (core system pods: CNI, CSI, kube-proxy)	Fixed Variables in ILP
Costly Restows	Pod Preemption & Eviction	Heavy Penalty Functions in Fitness Evaluation
BAPLIE / Final Stowage Plan	Pre-deployment Blueprint (Kuberina output YAML)	Final State Matrix (GA output)

3.2. Formal Definition

Notation

Symbol	Definition
$\mathcal{N} = \{n_1, \dots, n_m\}$	Set of nodes in the cluster
$\mathcal{P} = \{p_1, \dots, p_k\}$	Set of pods to schedule (excluding DaemonSet pods)
$\mathcal{R} = \{\text{CPU}, \text{RAM}, \text{GPU}, \dots\}$	Set of resource dimensions
$\mathcal{G} = \{G_1, \dots, G_q\}$	Set of pod groups (gangs)
$\mathcal{D} = \{d_1, \dots, d_h\}$	Set of DaemonSets
$x_{ij} \in \{0, 1\}$	Decision variable: 1 if pod p_i is assigned to node n_j
$y_j \in \{0, 1\}$	1 if node n_j has at least one pod assigned
req_i^r	Resource request of pod p_i for resource $r \in \mathcal{R}$
C_j^r	Allocatable capacity of node n_j for resource r , after DaemonSet pre-deduction
U_j^r	Utilization of node n_j for resource r : $U_j^r = \sum_i x_{ij} \cdot \text{req}_i^r / C_j^r$

Phase 0: DaemonSet Pre-deduction (Ballast Water Analogy)

DaemonSets are not decision variables — they are the ship's own systems (ballast, monitoring, communications), pre-deducted before optimization begins:

$$C_j^r = C_{j,\text{raw}}^r - \sum_{d \in \mathcal{D}} 1[\text{eligible}(d, n_j)] \cdot \text{res}_d^r$$

where $1[\text{eligible}(d, n_j)]$ is 1 if DaemonSet d runs on node n_j (based on `nodeSelector` and tolerations). After this step, $\mathcal{P}$ and C_j^r are the only inputs to the optimizer. This guarantees that the optimization space represents only the net

allocatable resources, preventing phantom capacity overflow at the end of the algorithmic process.

Decision Variables (Chromosome Encoding)

Each solution (blueprint) is encoded as a pod-level assignment vector:

$$\mathbf{s} = [x_1, x_2, \dots, x_k] \quad \text{where } x_i \in \{1, \dots, m\} \text{ is the node index for pod } p_i$$

Gang pods are **not** aggregated into macro-blocks. Each pod in a gang remains an individual decision variable (coupled variable in CSP), because each pod independently consumes resources on its assigned node.

Objective Function (Single-Objective, Weighted Sum)

$$\min F(\mathbf{s}) = w_1 \cdot f_{\text{nodes}}(\mathbf{s}) + w_2 \cdot f_{\text{frag}}(\mathbf{s}) + w_3 \cdot f_{\text{affinity}}(\mathbf{s}) + w_4 \cdot f_{\text{var}}(\mathbf{s}) + \Phi(\mathbf{s})$$

where:

Component	Formula	Maritime Analogy
f_{nodes}	$\sum_{j=1}^m y_j$ (number of active nodes)	Minimize number of bays used
f_{frag}	$\sum_{j:y_j=1} \sum_r \max(0, C_j^r - \sum_i x_{ij} \cdot \text{req}_i^r)$ (wasted capacity)	Minimize empty slots in used bays
f_{affinity}	Number of soft affinity/anti-affinity rule violations	Destination port grouping violations
f_{var}	$\text{Var}(\{U_j^r : y_j = 1\})$ (utilization variance across active nodes)	Vessel trim & stability
$\Phi(\mathbf{s})$	Hard constraint penalty: $+\infty$ if any hard constraint violated	Immediate rejection of illegal stowage

Unlike a pure bin packing formulation that only minimizes the number of bins used, this maritime-inherited model evaluates overall fitness across multiple competing objectives via a weighted sum.

Hard Constraints (CSP — must not violate)

1. **Capacity:** No node exceeds allocatable resources on any dimension.

$$\forall j, \forall r \in \mathcal{R} : \sum_{i=1}^k x_{ij} \cdot \text{req}_i^r \leq C_j^r$$

2. **Assignment:** Every pod is assigned to exactly one node.

$$\forall i : \sum_{j=1}^m x_{ij} = 1$$

3. **Taint/Toleration:** A pod can only be placed on a tainted node if it has the matching toleration.

$$\forall i, j : x_{ij} = 1 \implies \text{Taints}(n_j) \subseteq \text{Tolerations}(p_i)$$

4. **NodeSelector / NodeAffinity (required):** A pod can only be placed on nodes matching its selector.

$$\forall i, j : x_{ij} = 1 \implies \text{Labels}(n_j) \supseteq \text{Selector}(p_i)$$

5. **Gang All-or-Nothing (Block Booking):** For each pod group $G_q = \{p_{q_1}, \dots, p_{q_t}\}$, either all pods are feasibly placed, or none.

$$\forall G_q \in \mathcal{G} : \sum_{i \in G_q} 1[\text{feasible}(p_i)] = |G_q| \quad \text{or} \quad 0$$

This is a coupled constraint — each $x_{q_i,j}$ is a separate decision variable, but the group constraint binds them. The maritime analogy is Block Booking: individual containers with a commercial all-or-nothing commitment, where each container still has its own weight, type, and stability impact.

Soft Constraints (Fitness — optimize but don't reject)

1. **Pod Affinity (preferred):** Reward co-locating communicating pods on same node/zone.
2. **Pod Anti-Affinity (preferred):** Penalize co-locating conflicting pods.
3. **Topology Spread:** Penalize uneven distribution across zones/racks.

4. **Utilization Balance:** Minimize variance of utilization across active nodes (vessel stability analogy).

Complexity

The problem is a Multi-Dimensional Bin Packing Problem (MDBP), known to be **NP-hard** (Garey & Johnson, 1979). The search space is:

$$|\mathcal{S}| = m^k$$

For a medium cluster ($m = 100$, $k = 500$): $|\mathcal{S}| = 10^{1000}$ — brute-force enumeration is infeasible. This motivates the hybrid FFD (warm-start) + GA (heuristic optimization) + CSP (constraint enforcement) approach.

4. Proposed Method

4.1. System Overview: Offline CLI Engine Architecture

To implement maritime stowage heuristics in a cloud environment, the scheduling mechanism must be surgically decoupled from the live cluster. The system design requires a CLI tool that functions as an independent algorithmic architect — a simulation engine producing static pre-deployment scheduling blueprints [1].

Zero-Touch Interference. The defining characteristic of this engine is absolute separation from the `kube-apiserver`. The CLI operates entirely offline, typically integrated within a CI/CD pipeline or executed locally on an engineer's workstation [1]. Because it does not run as a controller or daemon inside the live cluster, it eliminates all risks of consuming precious control plane resources, causing API throttling, or triggering infinite race conditions during scheduling [1].

Input. The engine accepts two declarative YAML configuration files: (1) an infrastructure specification describing nodes, their resource capacities, labels, taints, and DaemonSet definitions; and (2) a workload specification describing pods, their resource requests, node selectors, affinity rules, and gang group

memberships.

Output. The engine produces a YAML blueprint mapping each pod to a specific node — a declarative artifact that can be version-controlled, code-reviewed, and applied via `kubect1 apply`.

4.2. Phase 1: Initialization via Vector Packing First-Fit Decreasing (FFD)

Motivation. A purely random initialization for the Genetic Algorithm in a highly constrained space (such as heterogeneous Kubernetes scheduling) results in an initial population composed almost entirely of infeasible solutions that violate capacity constraints. Correcting these violations takes the GA an exorbitant number of generations.

The FFD Warm-Start. We apply a greedy First-Fit Decreasing algorithm to generate a set of *feasible* initial blueprints, accelerating GA convergence by 3–5×.

1. **Synthetic Volume Calculation.** We calculate a scalar weight V_i for each pod based on normalized resource scarcity:

$$V_i = \alpha \cdot \text{CPU}_i + \beta \cdot \text{RAM}_i + \gamma \cdot \text{GPU}_i$$

where α, β, γ are tunable parameters reflecting the relative cost or scarcity of resources in the specific cluster. Because pod resources are multi-dimensional (RAM cannot compensate for CPU), a simple size-based sort is insufficient — the synthetic volume unifies dimensions into a single comparable metric.

2. **Decreasing Sort.** Pods are sorted in descending order of V_i . Maritime analogy: stow the heaviest and largest containers first.
3. **First-Fit Placement.** The algorithm iterates through the sorted pods and places each pod into the first node that has sufficient residual capacity across all dimensions.

This fast $O(k \log k + k \cdot m)$ heuristic produces the seed population for the GA, ensuring that subsequent evolutionary steps start from physically valid chromosomes.

4.3. Phase 2: Optimization via Genetic Algorithm (GA)

The GA optimizes the soft constraints (affinity, resource balancing, fragmentation) using the FFD output as its starting point.

1. **Population & Parallelism.** The population size scales with the problem size (e.g., $|Pop| = 512$ for a medium cluster of 100 nodes and 500 pods). Because fitness evaluation for each individual is completely independent, we implement an embarrassingly parallel evaluation model using Rust threads, achieving evaluation times of under 10 milliseconds per generation on a modern multi-core CPU.
2. **Selection.** We use Tournament Selection with tournament size $k_{\text{tour}} = 3$ to maintain high selection pressure while preserving diversity.
3. **Crossover with Gang Repair.** We apply Uniform Crossover. However, standard crossover can break the feasibility of gang scheduling (Block Booking). If a crossover operation splits a gang (e.g., pods 1–4 inherit from parent A, pods 5–8 inherit from parent B) and violates a node's capacity, a **Gang Repair Mechanism** is triggered: the algorithm rolls back the entire gang's assignment to match the parent that yielded a feasible placement for that gang.
4. **Mutation with Forward Checking.** We apply a random reset mutation with rate $p_m \approx 0.05$. Crucially, mutation is deeply integrated with the CSP solver. Before a pod is moved to a new node, the solver performs a forward capacity check. If the mutation violates hard constraints (or breaks a gang's all-or-nothing constraint), the mutation is rejected (rolled back). This prevents computational effort from being wasted on dead-end branches.
5. **Termination.** The GA employs an early stopping criterion. If the best fitness score in the population does not improve for N_{stop} consecutive generations (e.g., 200 generations), the algorithm assumes convergence to a near-optimal solution and halts.

4.4. Phase 3: Constraint Enforcement via CSP Solver with Forward Checking

Unlike traditional pipelines where the solver is a separate sequential step, Kuberina

tightly integrates the CSP solver *into* the FFD and GA operators (mutation and repair).

- **Hard Constraint Filtering.** Every placement decision (FFD insertion or GA mutation) is pre-screened by the CSP solver against Taints, Tolerations, NodeSelectors, and exact resource capacities. If an assignment is invalid, it is pruned immediately, saving the computational cost of full fitness evaluation.
- **Forward Checking for Block Booking.** When evaluating a placement for a pod belonging to a gang G_q , the CSP solver employs Forward Checking. It does not merely check if the target node has room for the *single* pod; it verifies whether the target node (or set of eligible nodes) possesses enough total residual capacity to accommodate the *entire* group G_q . If the collective requirement cannot be met, the branch is discarded instantly. This prevents the optimizer from wandering into deep infeasible regions of the search space — analogous to how a maritime planner would never begin stowing a block booking if the vessel cannot accommodate the full lot.

4.5. Gang Scheduling via the Block Booking Model

One of the most complex challenges in modern computing orchestration is managing distributed machine learning (ML) and AI training workloads [22]. These workloads depend entirely on specialized accelerator hardware and must satisfy a brutal operational requirement known as **Gang Scheduling** [36].

The All-or-Nothing Problem. Gang scheduling enforces an all-or-nothing lifecycle [36]. Large language models (LLMs) trained via data parallelism require the simultaneous presence of dozens of distributed pods. If a training job requires 64 GPUs simultaneously but the cluster can only aggregate 60 at that moment, partially filling 60 GPUs is catastrophically worse than doing nothing at all — the training process cannot start without the final 4 workers, leaving 60 expensive GPUs stranded indefinitely [22]. In conventional dynamic scheduling, this pattern frequently produces irrecoverable deadlocks [22].

Projects such as Volcano [36] and Kueue [39] were designed as overlay layers on Kubernetes to hold these jobs in queues until sufficient resources accumulate.

However, even with all-or-nothing admission semantics, hardware fragmentation remains the core barrier [22]. If an NVIDIA DGX Cloud architecture with hundreds of L40S GPUs experiences fragmentation, multi-node training jobs stall permanently — even though mathematically, the total free GPU count satisfies the requirement, but physically the GPUs are scattered across nodes where high-bandwidth interconnects like NVLink cannot reach [22].

The Maritime Solution: Block Booking and Gang Repair. The maritime logistics world has conquered this deadlock dynamic through two structural concepts: Block Booking and Slot Chartering. When a major logistics partner books a 500-TEU lot on a vessel, the stowage algorithm never fragments these containers across the hull; they are locked into a macroscopic geometric structure that cannot be split [1].

Kuberina's offline blueprint engine resolves Kubernetes gang scheduling deadlocks by explicitly modeling pod groups G_q as coupled variables within the CSP loop [1]. Forward Checking consolidates the required GPUs into algorithmic blocks. During the GA phase, if crossover inadvertently fractures a gang across invalid node boundaries, the Gang Repair mechanism reverses the chromosome to a clean inherited state [1]. Because the entire packing process executes before cloud deployment, it provides an absolute guarantee: once the blueprint is applied, exactly the required pods fill exactly the right GPU slots with 100% scheduling success — completely eliminating the scenario of expensive hardware idling due to the myopia of real-time scheduling.

5. Security and Pre-deployment Auditability

The shift to offline pre-deployment planning delivers substantial downstream benefits for the security posture of organizations, particularly those subject to strict compliance regimes such as FinTech or defense environments [35].

5.1. Dynamic Risk Analysis

In environments relying on automated dynamic scheduling, the velocity of entity creation and destruction (e.g., hundreds of nodes or 5,000 pods cycling per minute) generates an extremely large attack surface [35]. The 2024 Verizon Data Breach Investigation Report found that 15% of security breaches involved vulnerabilities in software supply chains and orchestration misconfigurations [35]. When placement decisions are fully delegated to dynamic scheduling, safety validation — such as ensuring a pod processing personally identifiable information (PII) never co-locates with a public-facing web application — rests entirely on Admission Controllers [40]. If these controllers fail, crash, or are bypassed, vulnerable software reaches production infrastructure unchecked [26].

5.2. Deterministic Pre-deployment Validation

With an offline static planning architecture, the final cluster topology becomes **deterministic at the earliest possible moment** [20]. When optimization completes, it delivers a declarative blueprint providing a comprehensive, container-centric view of risk [40]. Security teams can launch automated safety checks against the artifact repository: base image vulnerability scanning, static YAML analysis, namespace isolation verification, taint/toleration inspection, and global network policy enforcement [26, 35]. If the blueprint reveals logic errors that expose vulnerabilities, the CI/CD pipeline halts immediately — blocking the deployment before any flawed configuration touches running infrastructure [20].

This transforms infrastructure security review from a reactive runtime process into a proactive, auditable gate — analogous to code review before merge.

6. Experimental Setup

6.1. Testbed Description

We evaluate Kuberina on a synthetic benchmark designed to mirror the scale and heterogeneity of the MSC Irina mega-vessel. The testbed is generated using

research/gen_irina_testdata.py and comprises:

Parameter	Value
Total nodes	186
Node types	Standard (64-core, 256 GiB), Memory-optimized (32-core, 512 GiB), GPU (48-core, 192 GiB, 8× GPU)
Total pods	2,714
Constraint count	4,632 anti-affinity + 496 affinity = 5,128 total
Cluster CPU (raw / net)	10,272 / 9,853.5 cores
Cluster RAM (raw / net)	54,912 / 54,400.5 GiB
Cluster GPU	240 units
Pod CPU demand	7,198 cores (73.1% fill)
Pod RAM demand	27,840 GiB (51.2% fill)
Pod GPU demand	152 units (63.3% fill)
DaemonSets	4 (CNI, CSI, logging, monitoring — pre-deducted)

The "net" capacity reflects post-DaemonSet pre-deduction (Phase 0), ensuring the optimizer operates on physically allocatable resources only.

6.2. Experimental Configurations

We evaluate two configurations:

1. **Full Packing (100%):** No capacity cap — the optimizer packs pods as tightly

as possible to minimize active nodes.

2. **Pareto 80/20**: Node capacities are artificially capped at 80% to leave headroom for runtime bursts, simulating a production-realistic Resource Canal configuration.

6.3. Baselines

- **Random Uniform Placement**: Each pod is assigned to a uniformly random node (10,000 Monte Carlo trials).
- **Selector-Aware Random Placement**: Each pod is assigned to a uniformly random *eligible* node (respecting NodeSelector constraints only, 10,000 Monte Carlo trials).
- **FFD-only**: Phase 1 output without GA optimization (the seed fitness).
- **Theoretical LP Lower Bound**: Computed via the Coffman-Garey-Johnson (1978) homogeneous bound and a heterogeneous utilization bound.

6.4. Metrics

Metric	Definition
Active Nodes	Number of nodes with ≥ 1 pod assigned
Node Reduction	$(m - \text{active})/m \times 100\%$
Avg CPU Utilization	Mean of U_j^{CPU} across active nodes
Fragmentation	Total wasted capacity across active nodes: $\sum_{j:y_j=1} \sum_r (C_j^r - \sum_i x_{ij} \cdot \text{req}_i^r)$
Constraint Violations	Count of capacity, selector, and gang violations
Scheduling Success	$\text{placed pods} / \text{total pods} \times 100\%$
Utilization Variance	$\text{Var}(\{U_j^r : y_j = 1\})$

Metric	Definition
Wall-Clock Time	Solver execution time in seconds
Approximation Ratio (α)	active nodes/LP lower bound

6.5. Verification Methodology

All results are independently verified by an external Python validator (`research/inspector.py`) that re-reads the infrastructure, workload, and solution YAML files and checks every constraint from scratch. Additionally, `research/mathematical_proof.py` performs:

1. **Feasibility Proof:** Verifies all hard constraint predicates (capacity, assignment, node selector).
 2. **Optimality Bound:** Computes LP relaxation lower bounds and the approximation ratio.
 3. **Statistical Significance:** Runs 10,000 Monte Carlo random trials to compute p -values.
-

7. Results and Analysis

7.1. Full Packing Configuration (100% Capacity)

Metric	Kuberina	Best Random (Selector-Aware)
Pods placed	2,714 / 2,714 (100%)	—
Active nodes	152 / 186	—
Node reduction	18.3% (34 nodes freed)	—
Avg CPU	88.7%	—

Metric	Kuberina	Best Random (Selector-Aware)
utilization		
Capacity violations	0	76 (best of 10,000 trials)
Selector violations	0	0
Gang violations	0	—
Fragmentation	18,418.00	—
Affinity violations	643	—
Utilization variance	0.0374	—
Fitness	23,153.07	—
Wall-clock time	43.67 s	—

Key findings:

- **100% scheduling success** with zero hard constraint violations across all three dimensions (CPU, RAM, GPU).
- **34 nodes freed** for shutdown, representing direct infrastructure cost savings.
- **88.7% average CPU utilization** — more than double the industry average of 30–40%.
- FFD alone found the optimal seed (fitness did not improve after 199 GA generations), indicating that for this workload mix, the greedy warm-start was already near-optimal and the GA served primarily as a verification layer.

7.2. Pareto 80/20 Configuration (Resource Canal Mode)

Metric	Pareto 80%	Full Packing 100%
Pods placed	2,714 (100%)	2,714 (100%)
Active nodes	182 / 186	152 / 186
Node reduction	2.2% (4 nodes freed)	18.3% (34 nodes freed)
Avg CPU utilization	74.6%	88.7%
Max node utilization	79%	100%
Fragmentation	15,627.60	18,418.00
Affinity violations	549	643
Utilization variance	0.0256	0.0374
Fitness	20,192.65	23,153.07
Wall-clock time	43.71 s	43.67 s

Key findings:

- Even at 80% capacity cap, 100% of pods are successfully placed with zero violations.
- The 80% cap produces **lower utilization variance** (0.0256 vs 0.0374) — more evenly balanced nodes, directly analogous to better vessel stability (lower GM deviation).
- **Fewer affinity violations** (549 vs 643) because the optimizer has more room to satisfy soft constraints when not packing at maximum density.
- No node exceeds 79% utilization, leaving 20%+ headroom for runtime autoscaling — the "canal banks" for Autopilot to operate within.

7.3. Mathematical Verification

The external verifier (`mathematical_proof.py`) confirms:

Proof 1 — Constraint Satisfaction (Feasibility):

Predicate	Result
$\forall j \in \mathcal{N} : \sum_i x_{ij} \cdot \text{req}_i^r \leq C_j^r$ (Capacity)	✓ Satisfied (0 overflow on CPU, RAM, GPU)
$\forall i \in \mathcal{P} : \sum_j x_{ij} = 1$ (Assignment)	✓ Satisfied (0 missing pods)
$\forall i : x_{ij} = 1 \implies \text{Selector}(p_i) \subseteq \text{Labels}(n_j)$ (NodeSelector)	✓ Satisfied (0 violations)

Proof 2 — Optimality Bound (LP Relaxation):

Bound	Value
$L^{\text{CPU}} = \lceil 7198/61.75 \rceil$	117
$L^{\text{RAM}} = \lceil 27840/509.25 \rceil$	55
$L^{\text{GPU}} = \lceil 152/8.00 \rceil$	19
Homogeneous lower bound $L = \max(L^r)$	117
Heterogeneous utilization bound $L_{\text{het}} = \max(\lceil \rho^r \cdot m \rceil)$	136
Kuberina active nodes (Pareto 80%)	182
Approximation ratio $\alpha = 182/136$	1.3382

The approximation ratio $\alpha = 1.34$ exceeds the theoretical $11/9 \text{ OPT} + 6/9$ guarantee of FFD in one dimension, which is expected for multi-dimensional bin packing where dimensional conflicts prevent achieving the 1D bound.

Proof 3 — Statistical Significance (Monte Carlo):

Trial Type	Zero-Violation Rate	Avg Violations
Random Uniform (10,000 trials)	0 / 10,000	125.7 ± 5.8
Selector-Aware Random (10,000	0 / 10,000	89.1 ± 3.3 (best:

Trial Type	Zero-Violation Rate	Avg Violations
trials)		76)

Kuberina achieves 0 violations; the best random trial (even with selector awareness) achieves 76. The probability of random placement matching Kuberina's result is $p < 10^{-4}$, confirming statistical significance.

7.4. Computational Performance

Both configurations complete in under 44 seconds wall-clock time on a consumer-grade laptop (solver implemented in Rust with release-mode optimizations). The GA detects datacenter-scale input (2,714 pods) and automatically increases population size and generation budget, yet converges via early stopping at generation 199 — indicating that the FFD warm-start provides a strong initial solution that the GA efficiently validates.

8. Discussion

8.1. Limitations

Static vs. Dynamic. Kuberina is an offline static planner. When workloads change at runtime — due to autoscaling, pod crashes, or traffic spikes — the blueprint may become stale. Organizations must determine an appropriate re-planning frequency: per-deployment (CI/CD trigger), periodic (e.g., daily), or event-driven (when utilization deviation exceeds a threshold).

No Runtime Feedback Loop. Unlike Autopilot, Kuberina does not observe actual resource consumption. Its placement decisions are based solely on declared requests/limits, which may diverge from real-world usage patterns.

8.2. Scalability

The GA's search space grows as m^k , but the combination of FFD warm-start and

CSP pruning keeps practical runtime manageable. Our benchmark (186 nodes, 2,714 pods) completes in under 44 seconds. Scaling to 5,000+ pods on 1,000+ nodes would require profiling to determine whether the current early-stopping heuristic remains effective or whether adaptive population sizing and parallelism adjustments are needed.

8.3. Sensitivity to Hyperparameters

The weighting coefficients w_1, w_2, w_3, w_4 in the objective function, as well as FFD scarcity parameters α, β, γ and GA parameters $(|Pop|, p_m, k_{\text{tour}}, N_{\text{stop}})$, influence solution quality. In our experiments, the FFD seed was already near-optimal, suggesting that for well-structured workloads, the heuristic initialization dominates and GA hyperparameters have limited marginal impact. A formal sensitivity analysis across diverse workload profiles remains future work.

8.4. Practical Deployment

Kuberina is designed for integration into CI/CD pipelines: an infrastructure change or workload manifest update triggers `kuberina plan`, the blueprint undergoes code review (security audit, topology inspection), and upon approval, `kubectrl apply` deploys it. This workflow mirrors Terraform's `plan` → `apply` cycle, making infrastructure scheduling decisions reviewable, reproducible, and auditable.

8.5. Threats to Validity

Synthetic workloads. Our benchmark uses synthetic data generated to mirror mega-vessel-scale heterogeneity. While the constraint structure and resource profiles are realistic, validation on real-world cluster traces (e.g., Google Cluster Trace, Alibaba Cluster Trace) would strengthen external validity.

No gang scheduling in benchmark. The current benchmark loads 0 pod groups (gangs). While the gang scheduling machinery (coupled CSP variables, forward checking, gang repair) is implemented and formally specified, its effectiveness under load has not been empirically evaluated in this experiment.

9. Conclusion and Future Work

Conclusion

We presented Kuberina, an offline pre-deployment scheduling engine for heterogeneous Kubernetes clusters that draws a structural isomorphism from maritime container stowage planning. By reformulating pod scheduling as a Multi-Dimensional Bin Packing Problem and applying a three-phase hybrid pipeline — FFD warm-start, Genetic Algorithm optimization with gang-aware repair, and CSP Forward Checking — Kuberina produces declarative placement blueprints that are mathematically verified, auditable, and directly deployable via `kubectl apply`.

On a benchmark modeled after the MSC Irina mega-vessel (186 nodes, 2,714 pods, 5,128 constraints), Kuberina achieves 100% scheduling success with zero constraint violations, consolidates workloads from 186 to 152 active nodes (18.3% reduction), and reaches 88.7% average CPU utilization — more than doubling the industry average. The solution is statistically significant ($p < 10^{-4}$) and computes in under 44 seconds on a consumer laptop.

Beyond solution quality, Kuberina's primary contribution is the auditable blueprint artifact itself — a scheduling decision computed through thousands of evolutionary generations with mathematical justification for every placement, replacing opaque runtime decisions that cannot be reviewed, reproduced, or challenged.

Future Work

1. **Online/Incremental Re-planning.** Extend the engine to perform incremental re-optimization on deltas rather than re-solving the entire problem when workloads change.
2. **Multi-Objective Optimization.** Replace the weighted-sum objective with a Pareto front approach (e.g., NSGA-III) to expose the full trade-off surface between node count, fragmentation, affinity, and balance.
3. **Kubernetes Scheduler Extender Integration.** Develop a scheduler extender that automatically applies the blueprint as scoring preferences within `kube-scheduler`, bridging the gap between offline planning and runtime execution.

4. **Multi-Cluster / Federation Scheduling.** Extend the model to optimize placement across federated clusters with inter-cluster network latency constraints.
 5. **Reinforcement Learning Augmentation.** Investigate whether RL agents can replace or augment the GA for workload profiles with temporal patterns, using the FFD+CSP framework as the constraint backbone.
 6. **Gang Scheduling Empirical Evaluation.** Design benchmarks with realistic distributed AI training jobs (64–256 GPU pod groups) to empirically validate the Block Booking and Gang Repair mechanisms under load.
-

10. Author Contributions

Using the CRediT (Contributor Roles Taxonomy) framework, the author's contributions are defined as follows:

- **Conceptualization:** Dinh Tan Dung formulated the original research idea, discovering and defining the structural isomorphism between maritime container stowage (CSPP) and Kubernetes pod scheduling.
- **Methodology & Data Curation:** Dinh Tan Dung designed the constraint mapping taxonomy (e.g., Gang Scheduling as Block Booking, DaemonSets as Ballast Water) and designed the synthetic MSC Irina benchmark parameters.
- **Writing - Original Draft:** Dinh Tan Dung authored the initial conceptual narrative, framing the "Resource Canal" effect and the limitations of dynamic schedulers like Google Autopilot.
- **Formal Analysis, Software, & Validation:** Artificial Intelligence agents (Claude Opus and Gemini) were utilized as computational research assistants to formulate the mathematical LP relaxation proofs, implement the hybrid FFD+GA+CSP solver in Rust, run the Monte Carlo statistical validations, and synthesize the final academic English manuscript under the direction of the author.

11. Data Availability

The Kuberina CLI tool, the synthetic MSC Irina benchmark datasets (`irina_infra.yaml` , `irina_workloads.yaml`), and the verification scripts used in this study are available in the project's open-source repository: <https://github.com/AlexanderSlovov/kuberina>

The source code is released under the GNU Affero General Public License v3.0 (AGPLv3) to ensure modifications and integrations in network-accessible services remain open-source.

Acknowledgments

The author acknowledges the use of Anthropic's Claude and Google's Gemini models as collaborative research assistants for mathematical formalization, Rust software engineering, and language translation during the preparation of this manuscript.

References

- [1] Kuberina project documentation and design notes.
- [2] A. Tanaka et al., "A Benchmark Study of Deep Reinforcement Learning Algorithms for the Container Stowage Planning Problem," arXiv:2510.02589, 2025.
- [3] Z. Wang et al., "Many-Objective Container Stowage Optimization Based on Improved NSGA-III," *Journal of Marine Science and Engineering*, vol. 10, no. 4, p. 517, 2022.
- [4] "Integrating container stowage plan and yard operations," loadmaster.ai.
[Online]. Available: <https://loadmaster.ai/integrating-stowage-and-yard-planning-in-port-operations/>
- [5] P. Kaminsky, "A Multi-stage Decomposition Heuristic for the Container Stowage

Problem," University of California, Berkeley, 2008.

[6] A. Delgado et al., "An accurate model for seaworthy container vessel stowage planning with ballast tanks," DTU Orbit, 2012.

[7] "How Many Containers Fit on a Cargo Ship? (2026 Guide)," Ship4wd. [Online]. Available: <https://ship4wd.com/logistics-shipping/how-many-containers-fit-on-a-cargo-ship>

[8] "What is the largest container ship in the world?" Shipping Containers. [Online]. Available: <https://shipping-containers.com.au/what-is-the-largest-container-ship-in-the-world/>

[9] "The Growing Role of Mega-Ships in International Shipping," IoSCM. [Online]. Available: <https://www.ioscm.com/blog/the-growing-role-of-mega-ships-in-international-shipping/>

[10] "What Is A TEU? Calculating Cargo Ship Capacity (With Examples)," CHS Container Group. [Online]. Available: <https://chs-containergroup.com/us/what-is-a-teu-shipping/>

[11] L. Weiss-Cohen and L. Coelho, "Container Vessel Stowage Planning System Using Genetic Algorithm," Semantic Scholar, 2008.

[12] "An AIMMS-based decision-making model for optimizing the intelligent stowage of export containers in a single bay," *Discrete and Continuous Dynamical Systems - S*, 2019.

[13] D. Pacino, "Fast Generation of Container Vessel Stowage Plans," Ph.D. thesis, IT University of Copenhagen, 2012.

[14] A. Delgado et al., "An Accurate Model for Seaworthy Container Vessel Stowage Planning with Ballast Tanks," Sealytix, 2012.

[15] "Stowage plan for container ships," Grokipedia. [Online]. Available: https://grokipedia.com/page/Stowage_plan_for_container_ships

[16] "Container-Ship Stowage Planning Problem," Encyclopedia.pub. [Online]. Available: <https://encyclopedia.pub/entry/22494>

- [17] D. Pacino et al., "Fast Generation of Container Vessel Stowage Plans using mixed integer programming for optimal master planning and constraint-based slot planning," DTU Orbit, 2012.
- [18] "Matheuristics for Slot Planning of Container Vessel Bays," Sealytix. [Online]. Available: <https://www.sealytix.com/media/eqepwvxj/matheuristicsforslotplanningofcontainervesselbays.pdf>
- [19] "Models and solution algorithms for container terminal operations," DR-NTU. [Online]. Available: <https://dr.ntu.edu.sg/>
- [20] "SAGE — A Tool for Optimal Deployments in Kubernetes Clusters," arXiv:2307.06318, 2023.
- [21] "Google Autopilot cluster: unschedulable pods," Stack Overflow. [Online]. Available: <https://stackoverflow.com/questions/67031113/>
- [22] "Practical Tips for Preventing GPU Fragmentation for Volcano Scheduler," NVIDIA Developer Blog. [Online]. Available: <https://developer.nvidia.com/blog/practical-tips-for-preventing-gpu-fragmentation-for-volcano-scheduler/>
- [23] A. Aleinikov, "GKE Autopilot vs Standard 2026: Which Mode Should You Pick?" [Online]. Available: <https://www.alekseialeinikov.com/en/blog/topics/cloud/gke-autopilot-vs-standard-2026>
- [24] "Autopilot Became the Default Operation Mode for Google Kubernetes Engine," InfoQ, 2023.
- [25] "Auto-scaling Approaches for Cloud-native Applications: A Survey and Taxonomy," arXiv:2507.17128v1, 2025.
- [26] "Kubernetes and OpenStack Orchestration for Multi-Tenant Cloud Environments: Namespace Isolation and GPU Scheduling Strategies," ResearchGate, 2024.
- [27] "Software System for Container Vessel Stowage Planning," GECCO Companion, pp. 1519, 2015.
- [28] "Collaborative Optimization of Vessel Stowage Planning and Yard Pickup in

Automated Container Terminals," *Mathematics*, vol. 12, no. 21, p. 3387, 2024.

[29] "Literature Survey on the Container Stowage Planning Problem," arXiv:2307.07573, 2023.

[30] "Optimising Container Stowage: Minimising Relocations in Maritime Logistics," IE University. [Online]. Available: <https://www.ie.edu/university/>

[31] "Solving integrated problem of stowage planning with crane split by an improved genetic algorithm based on novel encoding mode," ResearchGate, 2022.

[32] "Genetic Algorithm Based Space-Optimised Arrangement of Containers and Stability in Containerships," University of Ibadan, *UIJSLICTR*, 2023.

[33] "A Genetic Algorithm for Solving a Container Storage Problem Using a Residence Time Strategy," *Studies in Informatics and Control*, vol. 26, no. 1, pp. 59-66, 2017.

[34] "Solving the Integrated Multi-Port Stowage Planning and Container Relocation Problems with a Genetic Algorithm and Simulation," *Applied Sciences*, vol. 12, no. 16, p. 8191, 2022.

[35] "Kubernetes Best Practices for Data Teams," [DataExpert.io](https://www.dataexpert.io/blog/kubernetes-best-practices-data-teams). [Online]. Available: <https://www.dataexpert.io/blog/kubernetes-best-practices-data-teams>

[36] "Plugins | Volcano," Volcano v1.8.2 Documentation. [Online]. Available: <https://volcano.sh/docs/v1.8.2/scheduler/plugins/>

[37] "Spark on Kubernetes — Gang Scheduling with YuniKorn," Cloudera Blog. [Online]. Available: <https://www.cloudera.com/blog/technical/spark-on-kubernetes-gang-scheduling-with-yunikorn.html>

[38] "Scheduling Group — Pods," Kubernetes Documentation. [Online]. Available: <https://kubernetes.io/docs/concepts/workloads/pods/scheduling-group/>

[39] "Gang scheduling, Priority scheduling, and Autoscaling for KubeRay CRDs with Kueue," Ray Documentation. [Online]. Available: <https://docs.ray.io/en/latest/cluster/kubernetes/k8s-ecosystem/kueue.html>

[40] "Container Security in 2026: 7 Key Components, Risks & Defenses,"

Checkmarx. [Online]. Available: <https://checkmarx.com/learn/container-security/>