



BRAIN. Broad Research in Artificial Intelligence and Neuroscience

e-ISSN: 2067-3957 | p-ISSN: 2068-0473

Covered in: Web of Science (ESCI); EBSCO; JERIH PLUS (hkdir.no); IndexCopernicus; Google Scholar; SHERPA/RoMEO; ArticleReach Direct; WorldCat; CrossRef; Peeref; Bridge of Knowledge (mostwiedzy.pl); abcdindex.com; Editage; Ingenta Connect Publication; OALib; scite.ai; Scholar9; Scientific and Technical Information Portal; FID Move; ADVANCED SCIENCES INDEX (European Science

Evaluation Centre, neredataltics.org); ivySCI; exaly.com; Journal Selector Tool (letpub.com); Citefactor.org; fatcat!; ZDB catalogue; Catalogue SUDOC (abes.fr); OpenAlex; Wikidata; The ISSN Portal; Socolar; KVK-Volltitel (kit.edu)

2026, Volume 17, Issue 3, pages: 337-348

Submitted: May 24th, 2026 | Accepted for publication: August 23rd, 2026

NeuroCare Grid: A Stigmergic Multi-Agent Architecture Evaluated through the NeuroCare-H Heuristic for Dynamic Surgical Scheduling

Kristijan Cincar *

West University of Timișoara and Tibiscus University of Timișoara, Romania;

Faculty of Computers and Applied Informatics, Tibiscus University of Timișoara, Romania.

kristijan.cincar@e-uvt.ro, k.cincar@tibiscus.ro, <https://orcid.org/0000-0002-5740-6680>

Victoria Iordan

Faculty of Computers and Applied Informatics, Tibiscus University of Timișoara, Romania.

Florentina-Anica Pintea

Faculty of Computers and Applied Informatics, Tibiscus University of Timișoara, Romania;

Center for Multidisciplinary Informatics Research (CCIM), Tibiscus University of Timișoara, Romania. <https://orcid.org/0009-0002-1214-3570>

* Corresponding author

Abstract: *Dynamic surgical scheduling must accommodate uncertain procedure durations, emergency arrivals, resource outages, and competing access and workload objectives. This paper presents NeuroCare Grid, a multi-agent scheduling architecture coordinated through an interpretable, multi-channel stigmergic field, and evaluates NeuroCare-H, its deterministic heuristic implementation. In a disclosed 30-day discrete-event simulation involving 500 synthetic admissions and an emergency share of 15.2%, NeuroCare-H achieved a scheduling rate of $92.47\% \pm 1.92\%$, an emergency waiting time of 0.01 ± 0.01 days, and operating-room utilization of $98.18\% \pm 1.20\%$. The corresponding FCFS results were $87.68\% \pm 1.02\%$, 0.40 ± 0.11 days, and $94.65\% \pm 1.96\%$, respectively. All three comparisons were based on 30 paired seeds and remained significant after Holm correction, with $p < 0.00001$. The scheduling advantage over FCFS persisted when the workload was increased to two and three times the baseline level. However, urgency-based scheduling produced a less equal distribution of waiting times, with a Gini coefficient of 0.63 for NeuroCare-H compared with 0.36 for FCFS. Removing the field slightly increased the scheduling rate from 92.47% to 92.81%, but worsened the Gini coefficient from 0.63 to 0.70. Replacing the multi-channel field with a pooled scalar produced no detectable difference. In a separate bootstrap experiment based on 37 surgical admissions from the open MIMIC-III Demo, the emergency share increased to 78.4%, and the scheduling-rate difference between NeuroCare-H and FCFS was no longer statistically significant ($p = 0.13$). These results establish a reproducible heuristic benchmark within the stated simulator and support the need for external operational validation before clinical use.*

Keywords: *multi-agent systems; stigmergy; operating-room scheduling; heuristic scheduling; discrete-event simulation; MIMIC-III; fairness audit.*

Cincar, K., Iordan, V., & Pintea, F.-A. (2026). NeuroCare Grid: A stigmergic multi-agent architecture evaluated through the NeuroCare-H heuristic for dynamic surgical scheduling. *BRAIN. Broad Research in Artificial Intelligence and Neuroscience*, 17(3), 337-348. <https://doi.org/10.70594/brain/17.3/20>

1. Introduction

Operating rooms (ORs) can be regarded as expensive and tightly integrated resources within hospitals. OR planning affects not only patients' access to treatment but also staff workload, overtime, and downstream bed demand (Cardoen et al., 2010; Macario et al., 1995). Dynamic allocation is challenging because it must accommodate both elective and emergency demand, the lengths of the procedures change unexpectedly, and it is feasible only when operating rooms, surgeons, anaesthesia resources, and beds are available at the same time (van Oostrum et al., 2008; Latorre-Nuez et al., 2016; Rahimi & Gandomi, 2021). Exact optimisation remains a standard approach. However, real-world disruptions usually cause the resorting to rolling horizons, decomposition, or reduced instances (Roshanaei & Naderi, 2021; Miao & Wang, 2021). Multi-agent and reinforcement-learning (RL) approaches provide convenient ways of looking at distributed sequential choices. Earlier works include using MARL for daily scheduling of ORs (Ribino et al., 2022), PPO-based emergency scheduling (Tong et al., 2026), and recently, MARL for allocation of medical resources and physiological monitoring (Hao et al., 2023; Shaik et al., 2025).

In fact, these studies also highlight an important validation requirement: simply being able to run learned policies is not sufficient evidence for a trained policy being efficient. Learning-based policies should have the benchmark model implemented, converged, and tested in the same scenarios for the sake of comparison of their performance (Schulman et al., 2017; Yu et al., 2022; Ning & Xie, 2024).

This paper therefore makes a narrower contribution. It (i) specifies a stable multi-channel stigmergic coordination field; (ii) embeds that field in a resource-explicit surgical-scheduling formulation; (iii) evaluates a deterministic heuristic instantiation rather than presenting it as trained MARL; (iv) compares it with FCFS, priority greedy, and rolling-horizon CP-SAT using common random numbers, ablations, sensitivity tests, and workload stress; and (v) checks the assumed emergency mix against the open MIMIC-III Demo. The conceptual architecture and the executed implementation are explicitly separated throughout the manuscript.

2. Related Work

OR scheduling covers strategic decision-making for resource allocation, intermediate-level planning via block scheduling, and detailed-level activities on job assignment and sequencing (Cardoen et al., 2010; Rahimi & Gandomi, 2021). Cyclic schedules can be used to buffer the possible variance in task durations (van Oostrum et al., 2008); extended formulations involve staff, post-anaesthesia beds, and emergencies (Latorre-Nuez et al., 2016); in addition, integrated planning may benchmark decomposition or exact methods (Roshanaei & Naderi, 2021). RL can be beneficial in situations where frequent decision revisions are required but at a more technical and clinical level, such studies need well-documented baselines and external validation.

PPO is a policy-gradient approach which is frequently used as a benchmark (Schulman et al., 2017). However, cooperative MARL is often implemented through centralised learning and decentralised rollout (Yu et al., 2022). In the domain of healthcare, MARL has been used for OR scheduling (Ribino et al., 2022), resource allocation (Hao et al., 2023), and patient monitoring (Shaik et al., 2025); in fact, MARL overall is very concerned with issues that pertain to the changing environment (non-stationarity), how to divide reward credits (credit assignment), the necessity for communication, and repeatability (Ning & Xie, 2024). Stigmergy is a mechanism where individuals work through the environment as a medium of indirect communication (Grasse, 1959).

Ant-colony optimization typically represents this signal as a scalar trail (Dorigo et al., 2006). NeuroCare Grid instead keeps five interpretable channels on a day–room lattice: demand, urgency, free capacity, conflict pressure, and fairness pressure. Interpretability matters in clinical decision support because recommendations must remain auditable and contestable (Amann et al., 2020). Whether these separate channels outperform a pooled scalar is an empirical question, tested rather than assumed in Section 7.2.

3. Scheduling Formulation

3.1. Sets, Parameters, and Decisions

Let P , D , R , S , and G denote patients, planning days, ORs, surgeons, and specialties. Patient $i \in P$ has release day a_i , latest acceptable day l_i , estimated duration p_i , urgency $u_i \in [0, 1]$, specialty $g_i \in G$, binary downstream-bed demand b_i , and simulated postoperative stay L_i . Room $r \in R$ provides regular capacity C_{rd} on day $d \in D$ and permits at most o'_{rd} overtime. Surgeon $s \in S$ has capacity H_{sd} and compatibility $m_{is} \in \{0, 1\}$. Downstream bed capacity on day d is B_d .

The binary assignment decision is

$$x_{ir ds} \in \{0, 1\},$$

where $x_{ir ds} = 1$ means patient i is assigned to room r , day d , and surgeon s ; zero means no such assignment. The non-negative variable o_{rd} is overtime and w_i is realised waiting time. Within-day sequencing is handled by the simulator and rolling-horizon repair.

3.2. Constraints

Each patient is scheduled at most once:

$$\sum_{r \in R} \sum_{d \in D} \sum_{s \in S} x_{ir ds} \leq 1, i \in P.$$

Assignments must fall inside the release–deadline window:

$$x_{ir ds} = 0, d < a_i \text{ or } d > l_i.$$

Room and surgeon capacity constraints are given by

$$\sum_{i,s} p_i x_{ir ds} \leq C_{rd} + o_{rd}, r, d, 0 \leq o_{rd} \leq o'_{rd}, r, d, \sum_{i,r} p_i x_{ir ds} \leq H_{sd}, s, d.$$

Specialty compatibility requires

$$x_{ir ds} \leq m_{is}, i, r, d, s.$$

The downstream-bed constraint is

$$\sum_{i,r,s} \sum_{\tau=\max(a_i, d-L_i+1)}^d b_i x_{ir \tau s} \leq B_d, d.$$

3.3. Conceptual Objective

The architecture's conceptual objective combines normalised terms:

$$\min J = \alpha W \sim + \beta E \sim + \gamma O \sim + \delta D \sim + \eta F \sim - \zeta U \sim,$$

where $\tilde{W}$, $\tilde{E}$, $\tilde{O}$, $\tilde{D}$, $\tilde{F}$, and $\tilde{U} \in [0, 1]$ denote mean wait, emergency lateness, overtime, deferrals, waiting-time inequality, and utilisation; α , β , γ , δ , η , and $\zeta \geq 0$ are conceptual weights. This expression belongs to the architectural specification. In the reported evaluation, CP-SAT maximises urgency-weighted scheduled cases, whereas NeuroCare-H applies the fixed heuristic coefficients reported in Table 2. The experiments therefore do not optimise J directly.

4. Data and Simulation Environment

MIMIC-III contains information on admission, service, procedure, and ICU stays, but does not provide a complete OR schedule, surgeon calendar, or reliable OR start and finish times (Johnson et al., 2016a; Johnson et al., 2016b). The main 500-admission cohort is therefore synthetic and only calibrated to selected admission-level descriptors. All OR capacities, calendars, durations, arrivals, and disruptions are simulated (Figure 1).

Data Provenance and Simulation Construction

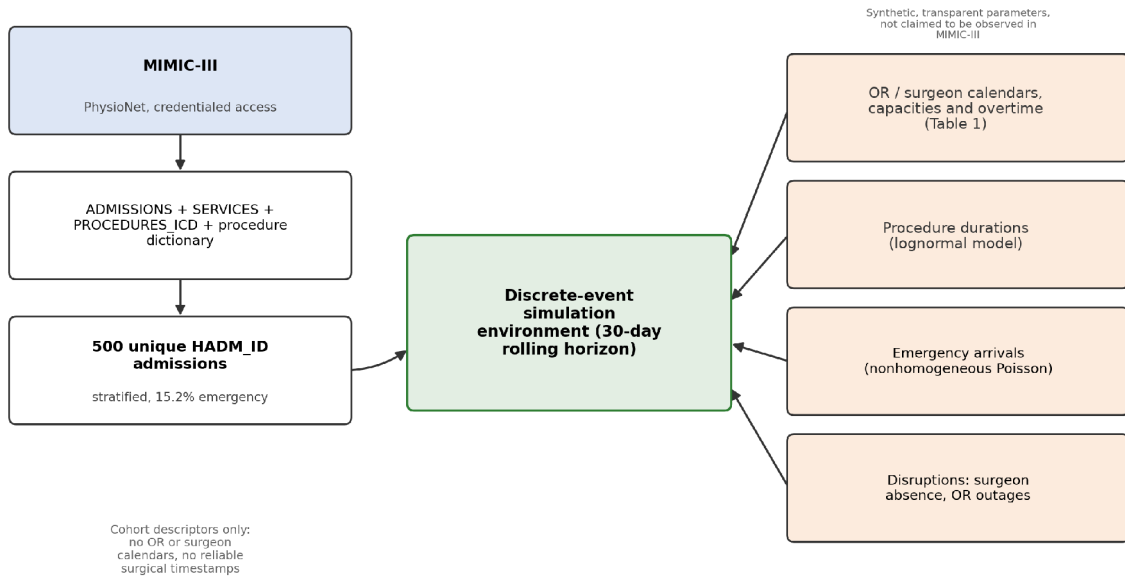


Figure 1. Data provenance and simulation construction. MIMIC-III supplies admission-level descriptors only; operational variables are synthetic.

For specialty g , the lognormal duration parameters are

$$\sigma_g^2 = \ln(1 + v_g^2), \mu_g = \ln(M_g), p_i \sim \text{LogNormal}(\mu_g, \sigma_g^2),$$

where M_g and v_g are the assumed specialty-specific median and coefficient of variation. These are scenario parameters, not quantities estimated from MIMIC-III. At decision time t , priority is

$$q_i(t) = 0.55u_i + 0.25\min\left(1, \frac{t-a_i}{l_i-a_i+1}\right) + 0.20c_i,$$

where $c_i \in [0, 1)$ is a prespecified complexity score. Only information available at scheduling time is used.

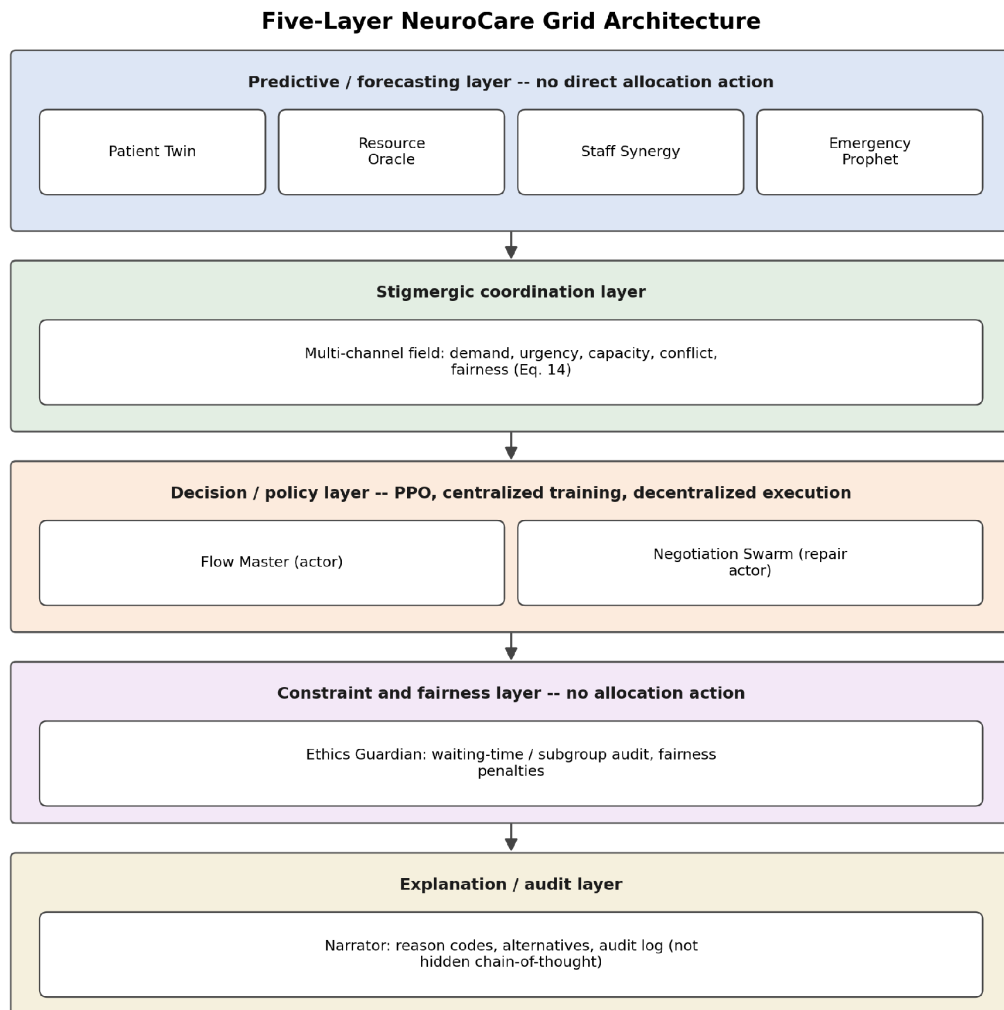
Table 1. Cohort and simulation parameters

Component	Setting	Generation rule or interpretation
Horizon	30 days	Rolling daily decisions; repair after each disruption.
Cohort	500 admissions	Synthetic stand-in; 15.2% emergency/urgent and 84.8% elective.
Rooms and overtime	5 ORs	10 regular hours/day plus at most 2 overtime hours/day per room.
Surgeons	10	Two per specialty; 8 regular hours/day; days off sampled before each run.
Specialties	Five	Cardiac, neurological, orthopaedic, general, thoracic.
Duration medians	240, 210, 150, 120, 180 min	Specialty-specific lognormal medians; CVs 0.35, 0.40, 0.35, 0.45, 0.40.
Turnover	20–35 min	Triangular distribution between consecutive cases in the same room.
Arrivals	Nonhomogeneous Poisson	Intensity calibrated to target emergency share; time sampled within day.

Component	Setting	Generation rule or interpretation
Deadlines	1-day emergency target	Elective deadline 7–30 days by procedure class.
Disruptions	Absence/outage	Surgeon absence 0.02/day; OR outage 0.01/day; common random numbers.
Stress	1×, 2×, 3×	Arrival-load multipliers; no improvement under stress is assumed.
Replications	30 seeds	Identical patient and disruption streams across methods within seed.

5. Conceptual Multi-Agent Architecture

Eight functional agent types form the five-layer architecture in Figure 2. Flow Master and Negotiation Swarm are the only components intended to host learned allocation and repair policies. The other agents encode state, forecast resources or arrivals, audit constraints and equity, or generate structured explanations. In the present study, none is trained; NeuroCare-H implements the coordination logic with fixed scores.



Clinician oversight: the scheduler proposes a slot plus binding constraints and alternatives; clinicians may lock or reject the proposal, and every decision is logged.

Figure 2. Five-layer conceptual NeuroCare Grid architecture. The current evaluation uses the heuristic NeuroCare-H controller.

5.1. Multi-Channel Stigmergic Field

Let $\Phi_{dr}^{(n)} \in R^K$ denote the K -channel field at day d , room r , and coordination iteration n . Its update is

$$\Phi_{dr}^{(n+1)} = (1 - \lambda)\Phi_{dr}^{(n)} + \rho \sum_{(d',r') \in N_{dr}} \left(\Phi_{d'r'}^{(n)} - \Phi_{dr}^{(n)} \right) + Z_{dr}^{(n)},$$

where $\lambda \in (0, 1)$ is decay, $\rho \geq 0$ is diffusion, N_{dr} is the four-neighbour lattice, and $Z_{dr}^{(n)} \in R^K$ is a clipped deposit vector. A sufficient convex-combination stability condition is

$$0 \leq \rho \leq \frac{1-\lambda}{\max_{d,r} |N_{dr}|}.$$

Componentwise soft thresholding removes low-amplitude numerical noise:

$$T_{\tau}(z) = \text{sign}(z) \max(|z| - \tau, 0),$$

where $\tau \geq 0$ is the threshold. Patient i scores assignment (d, r) as

$$A_i(d, r) = v_i^{\top} \Phi_{dr} + \theta^{\top} f_{i,dr},$$

where $v_i \in R^K$ contains channel weights, $f_{i,dr}$ is the local feature vector (urgency, expected wait, compatibility, downstream pressure, remaining capacity), and θ contains its fixed heuristic coefficients. NeuroCare-H evaluates feasible assignments with A_i and greedily selects the highest-scoring open slot. Figure 3 shows the coordination cycle executed by NeuroCare-H. At each iteration, agents observe the current local state, greedily score feasible open slots, update and diffuse the shared field, and use the revised field when evaluating subsequent assignments.

One Decentralized Coordination Cycle

Agents repeat this cycle each coordination iteration $n \rightarrow n+1$

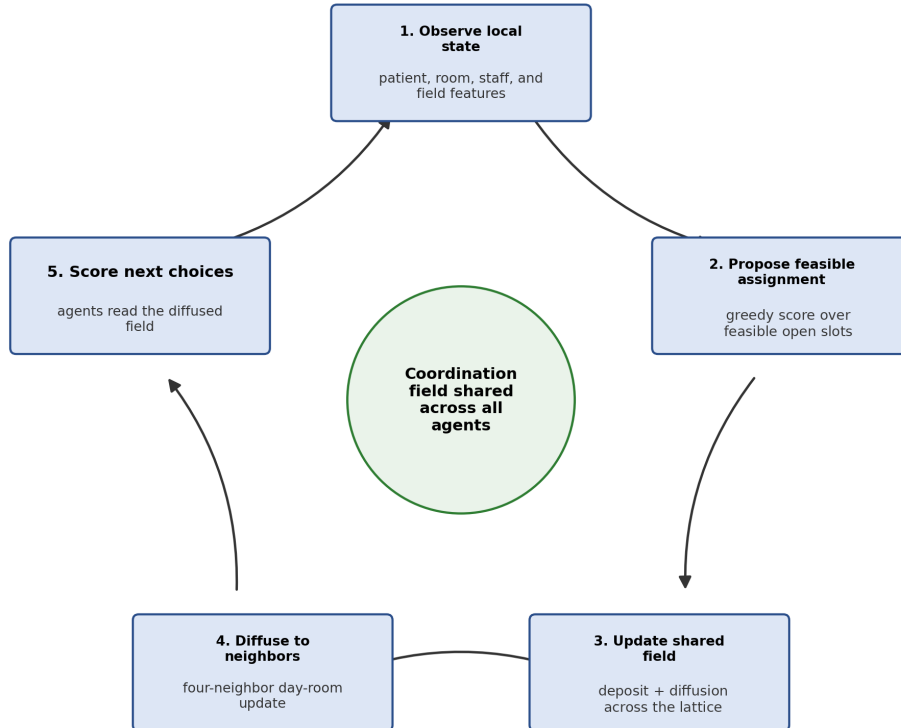


Figure 3. Coordination cycle executed by NeuroCare-H. Agents observe local state, greedily score feasible open slots, update and diffuse the shared field, and use the revised field in the next iteration.

5.2. Symbol Audit

All symbols are defined at first use. Indices i , d , r , s , g , and n identify the patient, day, room, surgeon, specialty, and coordination iteration. The symbol τ denotes a day summation index in the downstream-bed constraint and a threshold in the soft-thresholding operator; the meaning is explicit from context. Variables x , o , and w are decisions or outcomes; a , ℓ , p , u , g , b , L , c , and q are patient parameters; C , $\bar{o}$, H , m , and B are capacities or compatibility indicators. Tilded variables are normalised objective terms, Greek objective coefficients are non-negative weights, and Φ , λ , ρ , $\mathcal{N}$, Z , K , v , f , θ , A , and A define the stigmergic field and action score.

6. Executed Heuristic Evaluation

6.1. Compared Methods

The executed methods are random feasible, FCFS, priority greedy, rolling-horizon CP-SAT, NeuroCare-H, and a pooled-field proxy. Priority greedy is the field-ablated urgency scheduler used again in Section 7.2; the same method therefore serves as both a comparison baseline and the no-field ablation. The pooled and multi-channel controllers use identical fixed coefficients, which isolates the effect of field representation. PPO and trained MARL variants are outside the scope of the reported experiments. The parameter values used by the executed controllers are reported in Table 2. These values were held constant across paired runs so that differences between scheduling methods could be attributed to their decision rules rather than to different parameter settings.

The evaluation procedure is summarised in Algorithm 1. Every method receives the same patient cohort, procedure durations, staff absences, and operating-room outages within each seed, enabling paired statistical comparisons.

Table 2. Executed configuration

Parameter	Value
Field decay λ	0.30
Field diffusion ρ	0.15 (bound 0.175)
Threshold τ	0.01
Channel weights v	[0.6, 1.0, 0.5, -0.7, 0.8]
Field term weight	0.5
Urgency/emergency/fairness/complexity	0.5/0.9/0.6/0.2
CP-SAT budget	0.5 s/day; 1 worker
Anaesthesia teams	6 shared; 45 min/case
Seeds (main/sensitivity/stress)	30/10/15

Scheduling rate, emergency and elective wait, OR utilisation, throughput, overtime, Gini coefficient, and Jain index are recorded. Gini and Jain quantify the distribution of waiting time across the simulated cohort; they are not demographic-fairness measures. Methods share identical arrival, duration, absence, and outage streams within each seed. Paired Wilcoxon tests, paired bootstrap intervals, Holm multiplicity correction, and Hedges' g are used.

Algorithm 1. Paired reproducible evaluation

For seeds $k = 1, \dots, 30$, generate the cohort, durations, absences, and outages.

For every method m , reset the simulator to the identical seed- k scenario.

Run the 30-day horizon and store patient- and day-level outcomes.

Apply paired inference and publish configuration and run-level data.

7. Results

7.1. Main Synthetic-Cohort Results

Table 3 reports the 30-seed outcomes. Relative to FCFS, NeuroCare-H improved scheduling rate by 4.79 percentage points (95% paired-bootstrap CI [4.15, 5.44], Hedges' $g = 3.07$, Holm-corrected $p < 10^{-5}$), while emergency wait fell from 0.40 to 0.01 days. However, FCFS was more equitable in aggregate waiting time: Gini rose from 0.36 to 0.63 and Jain fell from 0.72 to 0.40 under NeuroCare-H (both $p < 10^{-8}$). Priority greedy and CP-SAT displayed still larger aggregate inequity. Priority greedy is the field-ablated urgency scheduler examined again in Section 7.2.

Table 3. Executed rerun: mean $\pm$ SD over 30 paired seeds

Metric	FCFS	Priority greedy	Rolling CP-SAT	NeuroCare-H	Pooled-field proxy
Scheduling rate (%)	87.68 ± 1.02	92.81 ± 1.61	93.95 ± 1.07	92.47 ± 1.92	92.43 ± 1.92
Emergency wait (days)	0.40 ± 0.11	0.02 ± 0.02	0.03 ± 0.02	0.01 ± 0.01	0.01 ± 0.01
Elective wait (days)	3.56 ± 0.39	4.32 ± 0.41	2.85 ± 0.26	4.37 ± 0.38	4.38 ± 0.40
OR utilization (%)	94.65 ± 1.96	99.33 ± 0.58	97.75 ± 1.20	98.18 ± 1.20	98.12 ± 1.03
Throughput (patients/day)	14.61 ± 0.17	15.47 ± 0.27	15.66 ± 0.18	15.41 ± 0.32	15.40 ± 0.32
Gini coefficient	0.36 ± 0.03	0.70 ± 0.02	0.79 ± 0.01	0.63 ± 0.02	0.63 ± 0.02
Jain index	0.72 ± 0.04	0.33 ± 0.02	0.23 ± 0.02	0.40 ± 0.02	0.40 ± 0.02
Overtime (h/day)	6.25 ± 0.15	6.50 ± 0.25	7.40 ± 0.24	6.11 ± 0.52	6.08 ± 0.46

7.2. Agent and Field Ablations

Removing Negotiation Swarm reduced scheduling rate from 92.47% to 77.64%, largely because that ablation also removed access to overtime; this is a coupled implementation effect, not evidence that an agent identity alone caused the difference. Removing Ethics Guardian reduced scheduling rate by 2.85 points and worsened Gini from 0.63 to 0.82. Removing Emergency Prophet left scheduling rate almost unchanged (92.53%) but increased emergency wait from 0.01 to 0.29 days.

The field-specific result is more nuanced. Priority greedy, which is the same field-ablated urgency scheduler introduced in Section 6.1, scheduled 92.81% of cases, compared with 92.47% for NeuroCare-H. The field therefore did not improve throughput in this setting. It was associated

instead with a more equal waiting-time distribution: Gini improved from 0.70 to 0.63 and Jain from 0.33 to 0.40. The difference between the pooled-field proxy and the five-channel field was so small that one could barely tell them apart (92.43% vs 92.47%, and the other metrics were almost the same too). One possible reason could be that it was possible to replace one of the rooms with another, and a fixed score for each room was given regardless of the room, which made the separate channels not have much structure to work with. Testing that explanation requires room-specific constraints and matched controllers evaluated with pooled, multi-channel, and no-field inputs under the same computational budget.

7.3. Sensitivity and Workload Stress

Across four, five, and six ORs and duration scales of 0.8 and 1.2, FCFS was worst in four of five settings, but all three heuristic methods converged within 0.5 points when durations were inflated by 20%. Thus the FCFS disadvantage is not universal. Under $2\times$ load, FCFS dropped 10.88 points while NeuroCare-H dropped 0.08; under $3\times$ load, the drops were 14.19 and 1.37 points. This supports lower degradation under the simulated stress test, not improvement under stress or clinical resilience.

7.4. Open MIMIC-III Demo Check

Among 129 admissions in MIMIC-III Demo v1.4, 37 had a surgical service: 19 general, 7 cardiac, 5 neurological, 5 thoracic, and 1 orthopaedic (Johnson et al., 2019). Their emergency/urgent share was 78.4%, not 15.2%. A 500-case bootstrap resample from their joint specialty, emergency-status, and length-of-stay distribution still required synthetic OR times, priorities, and deadlines.

On this real-derived emergency mix, scheduling rates were $86.46\% \pm 1.80\%$ (FCFS), $88.26\% \pm 1.99\%$ (priority), $92.12\% \pm 1.22\%$ (CP-SAT), and $86.85\% \pm 1.78\%$ (NeuroCare-H). The NeuroCare-H-FCFS difference shrank to 0.39 points and was not significant ($p = 0.13$); CP-SAT and priority greedy both outperformed NeuroCare-H ($p < 0.001$). This reversal is given equal prominence to the main synthetic result. Figure 4 summarises the external grounding experiment based on the open MIMIC-III Demo. It presents the observed surgical-specialty distribution, contrasts the assumed emergency share with the real-derived share, and reports the scheduling results obtained after bootstrap resampling.

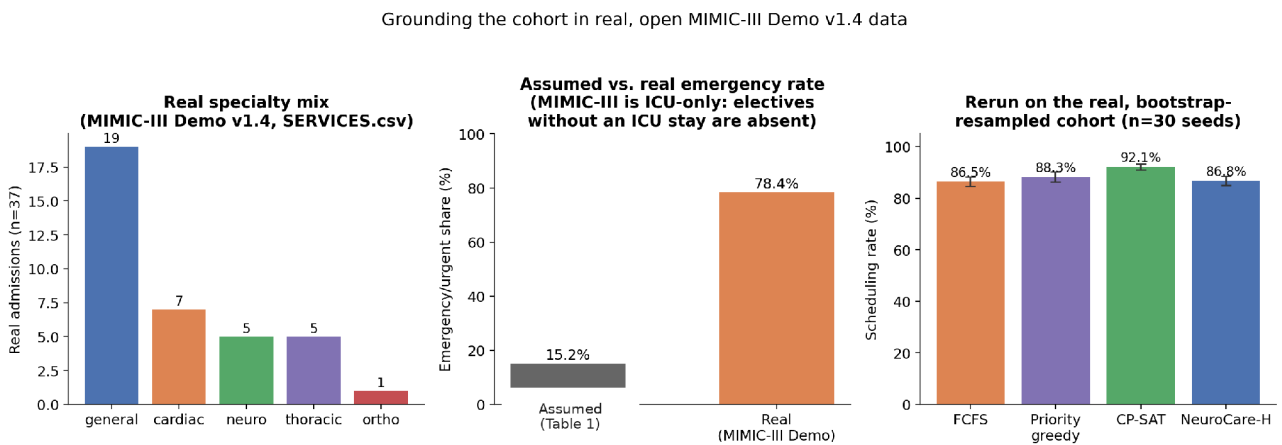


Figure 4. Open MIMIC-III Demo check: surgical-service mix, assumed versus observed emergency share, and the rerun under the real-derived mix.

7.5. Fairness and Subgroup-Analysis Boundary

Gini and Jain describe how waiting time is distributed across the simulated cohort; in this paper, we refer to that property as aggregate waiting-time equity. They do not measure demographic fairness. The main cohort is synthetic, and the open Demo contains only 37 surgical admissions,

including one orthopaedic case, so neither source supports stable estimates by ethnicity, sex, insurance status, or their intersections. A later research ought to identify and preregister potential subgroups, and their minimum cell sizes. The results of analyses stratified on variables like ethnicity, gender, and the type of practice should be given with their level of uncertainty. Such a reporting should be done at the minimum, for a sufficiently large, multi-hospital cohort. The current role of the Ethics Guardian is only that of a checker and enforcer but not a proof that demographic bias has been eradicated.

8. Discussion

The experiments provide limited support to the conclusions. NeuroCare-H is a reproducible demonstration of a heuristic system that gives results which are a lot higher than FCFS in the authors' main synthetic test and is much more stable under workload stress. Also, these findings also indicate that CP-SAT is the most effective system scheduler, multi-channel field doesn't surpass pooled scalar, urgency prioritisation harms waiting times fairness compared to FCFS, and the NeuroCare-H, FCFS efficiency gap disappears when the MIMIC-III Demo-derived emergency mix is considered. In order for the controller to work, it uses fixed scores instead of learned policies. If the extension is to be learning-based, matched PPO and MARL baselines, convergence evidence, learning curves, computational budgets and evaluation on held-out seeds are necessary. It is equally important to have the external operational data to calibrate OR capacity, surgeon rosters, procedure durations, outages, and emergency arrivals as MIMIC-III lacks those scheduling parameters.

Field ablation study only gets meaningful attention when it's negative. So far, in the room-interchanger simulator, no five-channel representation has outperformed a scalar proxy. A future factorial experiment should vary room specialisation and communication structure and compare identical trained policies with multi-channel, pooled, and no-field observations. That design would isolate the field from agent identity and heuristic score construction. For prospective evaluation, a stepped-wedge cluster design across OR units could compare scheduling rate and emergency wait while monitoring utilisation, overtime, override rate, cognitive load, and subgroup equity. Institutional review, clinician override, data governance, and prespecified stopping rules for emergency-delay or equity harm are mandatory.

9. Conclusion

NeuroCare Grid defines the stigmergic multi-agent architecture, while NeuroCare-H is the fixed heuristic evaluated in this study. In the synthetic scenario with a 15.2% emergency share, NeuroCare-H improved scheduling rate and emergency wait relative to FCFS, but produced a less equal waiting-time distribution and did not match rolling CP-SAT on scheduling rate. The multi-channel field improved aggregate equity relative to the no-field urgency scheduler, although it offered no detectable advantage over a pooled scalar. When the experiment was repeated with the MIMIC-III Demo-derived emergency mix of 78.4%, the scheduling-rate difference between NeuroCare-H and FCFS disappeared. These results provide a transparent benchmark for further development and show where external data and stronger comparative evaluation are most needed.

Data and Code Availability

MIMIC-III is available through PhysioNet under credentialed access (MIT Laboratory for Computational Physiology 2016); MIMIC-III Demo v1.4 is open access (Johnson et al., 2019). The main cohort was generated synthetically from the disclosed configuration. The generators, scheduler implementations, run-level CSV files, and analysis scripts should accompany the submission so that the complete evaluation can be reproduced.

Conflict of Interest

The authors declare no conflict of interest.

References

- Amann, J., Blasimme, A., Vayena, E., Frey, D., & Madai, V. I. (2020). Explainability for artificial intelligence in healthcare: A multidisciplinary perspective. *BMC Medical Informatics and Decision Making*, 20, 310. <https://doi.org/10.1186/s12911-020-01332-6>
- Cardoen, B., Demeulemeester, E., & Beliën, J. (2010). Operating room planning and scheduling: A literature review. *European Journal of Operational Research*, 201(3), 921–932. <https://doi.org/10.1016/j.ejor.2009.04.011>
- Dorigo, M., Birattari, M., & Stützle, T. (2006). Ant colony optimization. *IEEE Computational Intelligence Magazine*, 1(4), 28–39. <https://doi.org/10.1109/MCI.2006.329691>
- Grassé, P.-P. (1959). La reconstruction du nid et les coordinations interindividuelles chez *Bellicositermes natalensis* et *Cubitermes* sp. La théorie de la stigmergie: Essai d'interprétation du comportement des termites constructeurs. *Insectes Sociaux*, 6(1), 41–80. <https://doi.org/10.1007/BF02223791>
- Hao, Q., Xu, F., Chen, L., Hui, P., & Li, Y. (2023). Hierarchical multi-agent model for reinforced medical resource allocation with imperfect information. *ACM Transactions on Intelligent Systems and Technology*, 14(1), 1–27. <https://doi.org/10.1145/3552436>
- Johnson, A. E. W., Pollard, T. J., Shen, L., Lehman, L.-W. H., Feng, M., Ghassemi, M., Moody, B., Szolovits, P., Celi, L. A., & Mark, R. G. (2016a). MIMIC-III, a freely accessible critical care database. *Scientific Data*, 3, 160035. <https://doi.org/10.1038/sdata.2016.35>
- Johnson, A., Pollard, T., & Mark, R. (2019). MIMIC-III clinical database demo (Version 1.4). PhysioNet. <https://doi.org/10.13026/C2HM2Q>
- Latorre-Núñez, G., Lüer-Villagra, A., Marianov, V., Obreque, C., Ramis, F., & Neriz, L. (2016). Scheduling operating rooms with consideration of all resources, post anesthesia beds, and emergency surgeries. *Computers & Industrial Engineering*, 97, 248–257. <https://doi.org/10.1016/j.cie.2016.05.016>
- Macario, A., Vitez, T. M., Dunn, B., & McDonald, T. (1995). Where are the costs in perioperative care? Analysis of hospital costs and charges for inpatient surgical care. *Anesthesiology*, 83(6), 1138–1144. <https://doi.org/10.1097/0000542-199512000-00002>
- Miao, H., & Wang, J.-J. (2021). Scheduling elective and emergency surgeries at shared operating rooms with emergency uncertainty and waiting time limit. *Computers & Industrial Engineering*, 160, 107551. <https://doi.org/10.1016/j.cie.2021.107551>
- Johnson, A., Pollard, T., & Mark, R. (2016b). MIMIC-III clinical database (Version 1.4). PhysioNet. <https://doi.org/10.13026/C2XW26>
- Ning, Z., & Xie, L. (2024). A survey on multi-agent reinforcement learning and its application. *Journal of Automation and Intelligence*, 3(2), 73–91. <https://doi.org/10.1016/j.jai.2024.02.003>
- van Oostrum, J. M., van Houdenhoven, M., Hurink, J. L., Hans, E. W., Wullink, G., & Kazemier, G. (2008). A master surgical scheduling approach for cyclic scheduling in operating room departments. *OR Spectrum*, 30(2), 355–374. <https://doi.org/10.1007/s00291-006-0068-x>
- Rahimi, I., & Gandomi, A. H. (2021). A comprehensive review and analysis of operating room and surgery scheduling. *Archives of Computational Methods in Engineering*, 28(3), 1667–1688. <https://doi.org/10.1007/s11831-020-09432-2>
- Ribino, P., Di Napoli, C., & Serino, L. (2022). A multi-agent RL algorithm for single-day operating room scheduling. In H. H. Alvarez Valera & M. Luštrek (Eds.), *Workshops at 18th International Conference on Intelligent Environments (IE2022)* (Vol. 31, pp. 277–287). IOS Press. <https://doi.org/10.3233/AISE220053>
- Roshanaei, V., & Naderi, B. (2021). Solving integrated operating room planning and scheduling: Logic-based Benders decomposition versus Branch-Price-and-Cut. *European Journal of Operational Research*, 293(1), 65–78. <https://doi.org/10.1016/j.ejor.2020.12.004>
- Schulman, J., Wolski, F., Dhariwal, P., Radford, A., & Klimov, O. (2017). Proximal policy optimization algorithms. arXiv. <https://doi.org/10.48550/arXiv.1707.06347>

- Shaik, T., Tao, X., Li, L., Xie, H., Dai, H.-N., Zhao, F., & Yong, J. (2025). AI-driven multi-agent reinforcement learning framework for real-time monitoring of physiological signals in stress and depression contexts. *Brain Informatics*, 12(1), Article 14. <https://doi.org/10.1186/s40708-025-00262-1>
- Tong, Y., Li, Z., Guo, T., Huang, X., & Li, H. (2026). Improving efficiency of general surgery department by dynamic surgical scheduling with emergency uncertainty. *Complex & Intelligent Systems*, 12(3), Article 105. <https://doi.org/10.1007/s40747-025-02216-w>
- Yu, C., Velu, A., Vinitzky, E., Gao, J., Wang, Y., Bayen, A., & Wu, Y. (2022). The surprising effectiveness of PPO in cooperative multi-agent games. *Advances in Neural Information Processing Systems*, 35, 24611–24624.