

Lifting in Multi-agent Systems under Uncertainty

Tanya Braun¹, Marcel Gehrke², Florian Lau², Ralf Möller²

¹Computer Science Department, University of Münster,

²Institute of Information Systems, University of Lübeck, ³Institute of Telematics, University of Lübeck

Multi-agent Systems

- Joint state and reward for transition, sensor, reward functions
- Individual actions and observations

Complexity exponential in number of agents n



Problems with large n not computable, e.g., nano-scale systems

Under Symmetries

- Agent types: same actions and observations available
- Additional symmetries in functions

Best case symmetry: worst-case dependency down to $\log n$ only



Reuse existing algorithms for types; new task: optimise w.r.t. n

Lifting in Multi-agent Systems under Uncertainty

Tanya Braun¹, Marcel Gehrke², Florian Lau², Ralf Möller²

¹University of Münster, Münster, Germany

²University of Lübeck, Lübeck, Germany



Link to preliminary paper (QR code):
https://www.ifis.uni-luebeck.de/uploads/tx_wapublications/braun_143_public.pdf
 Please cite full version when published by PMLR.

DecPOMDP

- Decentralised Partially Observable Markov Decision Problem
 - Set of agents working towards a joint reward
 - Environment a Markov-1 stochastic process
- DecPOMDP model M

$$(I, S, \{A_i\}_{i=1}^N, T, R, \{O_i\}_{i=1}^N, \Omega)$$
 - I a set of agents, $|I| = N$,
 - S a random variable for the state space,
 - A_i a decision variable; $A = \times_{i=1}^N \text{ran}(A_i)$ set of joint actions
 - $T(S', S, A) = P(S' | S, A)$ a transition model,
 - $R(S)$ a reward function,
 - O_i a random variable; $O = \times_{i=1}^N \text{ran}(O_i)$ set of joint observations,
 - $\Omega(O, S) = P(O | S)$ a sensor model.
 - Given horizon τ , each agent i has a local policy π_i :
 $\text{ran}(O_{i,(0:t)}) \mapsto \text{ran}(A_i); t < \tau; \pi = (\pi_i)_{i=1}^N$ a joint policy
- DecPOMDP Semantics: all possible joint policies Π_M
- DecPOMDP Problem: find joint policy π^* that maximises the expected utility with discount factor $\gamma \in [0, 1]$

$$U_M^\pi(s_t, \mathbf{o}_{0:t}) = R(s_t, \pi(\mathbf{o}_{0:t})) + \gamma \sum_{s_{t+1} \in \text{ran}(S)} T(s_{t+1}, s_t, \pi(\mathbf{o}_{0:t})) \cdot \sum_{\mathbf{o}_{t+1} \in \text{ran}(O)} \Omega(\mathbf{o}_{t+1}, s_{t+1}) U_M^\pi(s_{t+1}, \mathbf{o}_{0:t+1})$$

Problem

- Complexities** T, R, Ω of T, R, Ω , evaluation cost $\mathbb{C}$ of a joint policy and size of policy space $\mathbb{P}$ **exponential** in N

$$T \in O(s^2 a^N) \quad R \in O(s a^N) \quad \Omega \in O(s o^N)$$

$$\mathbb{C} \in O(s o^{N\tau}) \quad \mathbb{P} \in O(a^{\frac{N(o^\tau - 1)}{o - 1}})$$
 - $s = |\text{ran}(S)|, a = \max_i |\text{ran}(A_i)|, o = \max_i |\text{ran}(O_i)|$
- Especially a problem with $N \gg 10,000$

Symmetric & Partitioned DecPOMDP

- Symmetric DecPOMDP: $K \ll N$ partitions in agent set
 - In each partition: the same set of actions and observations
 - Formally, $I = \bigcup_{k=1}^K \mathfrak{S}_k, \mathfrak{S}_k \neq \emptyset, \mathfrak{S}_k \cap \mathfrak{S}_{k'} = \emptyset$; for each $\mathfrak{S}_k$:
 $\forall i, j \in \mathfrak{S}_k : \text{ran}(A_i) = \text{ran}(A_j) \wedge \text{ran}(O_i) = \text{ran}(O_j)$
- Partitioned model

$$(\bar{I}, S, \{\bar{A}_k\}_{k=1}^K, \bar{T}, \bar{R}, \{\bar{O}_k\}_{k=1}^K, \bar{\Omega})$$
 - $\bar{I}$ a partitioning $\{\mathfrak{S}_k\}_{k=1}^K$ of an agent set, $n_k = |\mathfrak{S}_k|, |\bar{I}| = N$,
 - $\bar{A}_k, \bar{O}_k$ variables per partition, with joint actions and observations defined over the K partitions
 - $\bar{T}, \bar{R}, \bar{\Omega}$ like T, R, Ω but defined over partitioned inputs

Counting DecPOMDP

- Indistinguishability of agents in a partition yields invariance towards which particular partition agents perform an action or observe an event: it only matters how many agents do or observe something
 - E.g., 2 actions performed by 5 agents each: 10 over 5 and 5 = 255 different permutations of 10 agents to do the actions with the same outcome
 - Count occurrences $[n_1, \dots, n_l], l = |\text{ran}(A_k)|$
- Counting DecPOMDP: a partitioned DecPOMDP with
 - $\bar{A}_k = \#_{X_k}[A_k(X_k)]$ a counting random variable
 - $\bar{O}_k = \#_{X_k}[O_k(X_k)]$ a counting random variable
 - $\bar{T}, \bar{R}, \bar{\Omega}$ defined over the counting random variables

$$(\bar{I}, S, \{\#_{X_k}[A_k(X_k)]\}_{k=1}^K, \bar{T}, \bar{R}, \{\#_{X_k}[O_k(X_k)]\}_{k=1}^K, \bar{\Omega})$$
- Theorem:**
 Counting model M_c has an equivalent ground model $\text{gr}(M_c)$.
- Theorem:**
 Optimal policy in M_c also optimal in $\text{gr}(M_c)$.
- Complexities** $T_c, R_c, \Omega_c, \mathbb{C}_c$ **polynomial**, $\mathbb{P}_c$ **exponential** in $n < N$

$$T_c \in O(s^2 n^{Ka}) \quad R_c \in O(s n^{Ka}) \quad \Omega_c \in O(s n^{Ko})$$

$$\mathbb{C}_c \in O(s n^{K\tau o}) \quad \mathbb{P}_c \in O(n^a \frac{K(n^{\tau o} - 1)}{n^o - 1})$$
 - $n = \max_k n_k, |\text{ran}(\bar{A}_k)| \leq n^a, |\text{ran}(\bar{O}_k)| \leq n^o$

Isomorphic DecPOMDP

- Independence assumption between agents of a partition
 - $\bar{T}, \bar{R}, \bar{\Omega}$ factorise identically for each agent within each partition
 - Higher efficiency at the expense of lower expressiveness
- Isomorphic DecPOMDP: a partitioned DecPOMDP with
 - $\bar{A}_k = A_k(X_k)$ a parameterised random variable
 - $\bar{O}_k = O_k(X_k)$ a parameterised random variable
 - $\bar{T}, \bar{R}, \bar{\Omega}$ defined over the parameterised random variables

$$(\bar{I}, S, \{A_k(X_k)\}_{k=1}^K, \bar{T}, \bar{R}, \{O_k(X_k)\}_{k=1}^K, \bar{\Omega})$$
- Lemma:**
 Isomorphic model M_i has an equivalent counting model M_c .
- Lemma:**
 Enough to define policies in M_i over $\text{ran}(A_k)$ and $\text{ran}(O_k)$.
- Corollary:**
 Partition sizes only influence U_M^π , not π^* itself.
- Complexities** $T_i, R_i, \Omega_i, \mathbb{C}_i, \mathbb{P}_i$ **logarithmic** in $n < N$

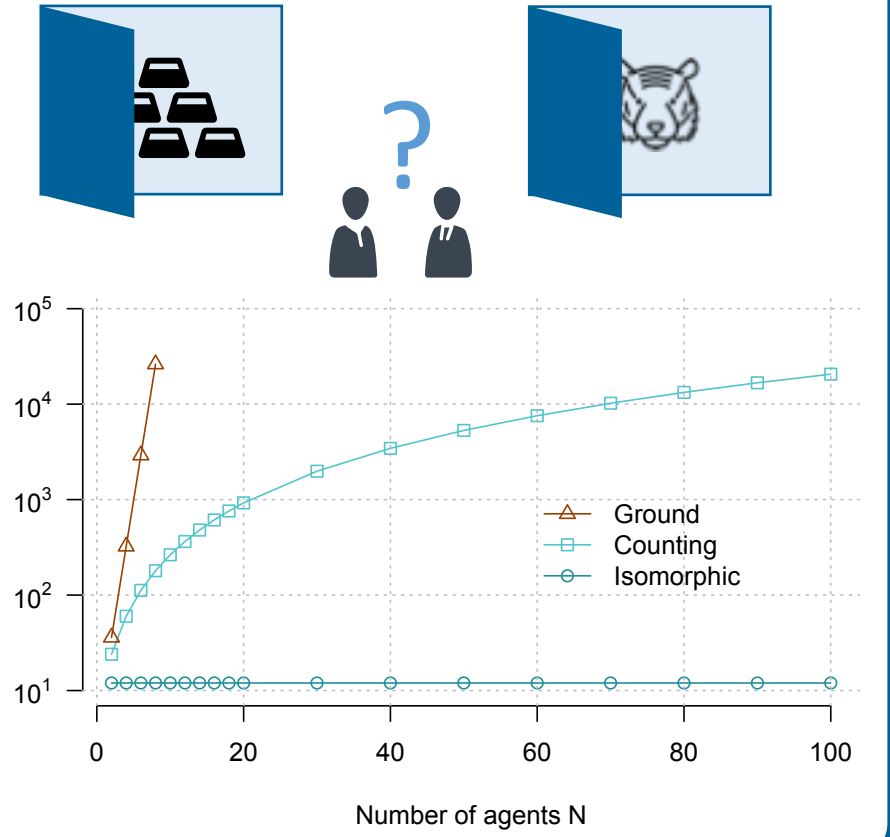
$$T_i \in O(s^2 a^K) \quad R_i \in O(s a^K) \quad \Omega_i \in O(s o^K)$$

$$\mathbb{C}_i \in O(\log_2(n) s o^{K\tau}) \quad \mathbb{P}_i \in O(a^{\frac{K(o^\tau - 1)}{o - 1}})$$

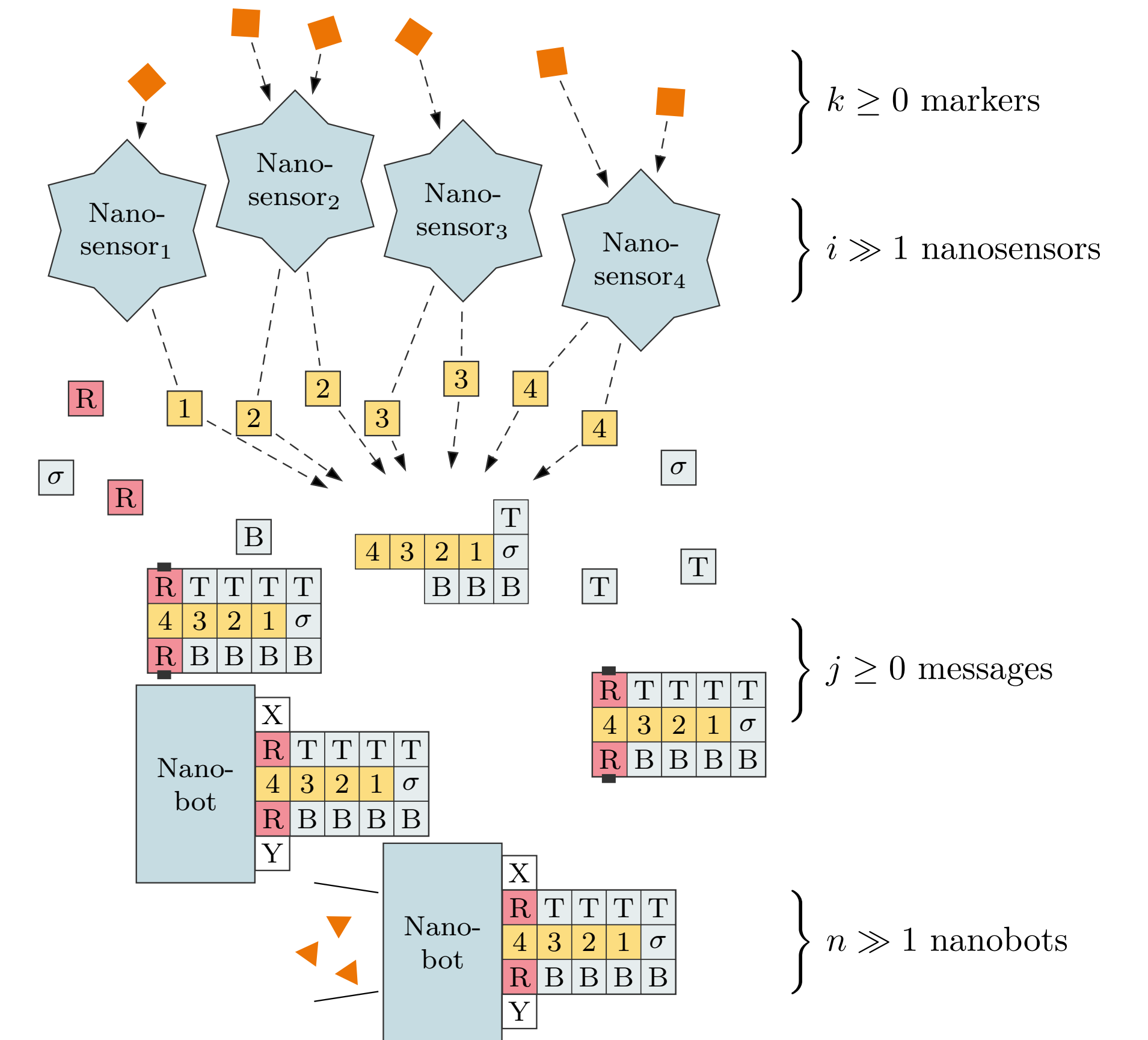
Possible to reuse existing solution approaches to solve DecPOMDP for K partitions; adapt result w.r.t. sizes n_k to reach a given U

Example: DecTiger

- $I = \{\text{agent}_1, \text{agent}_2\}, \mathfrak{S}_{K=1} = I$
- $\text{ran}(S) = \{\text{tigerleft}, \text{tigerright}\}$
- $A_K(X_K) = \{\text{listen}, \text{openleft}, \text{openright}\}$
- $O_K(X_K) = \{\text{hearleft}, \text{hearright}\}$
- As a counting model:
 T, R, Ω equivalently encoded in $\bar{T}, \bar{R}, \bar{\Omega}$
- As an isomorphic model:
 - Reset in T not possible to encode in $\bar{T}$
 - Agreeing on an action yielding a lower punishment not possible to encode in $\bar{R}$
 - Ω equivalently encoded in $\bar{\Omega}$



Application: Nanoscale Medical Systems



- Agent set: $\kappa = 4$ different types of markers/sensors, $\iota = 1$ different types of messages/nanobots; $\kappa + \iota = K, n_k \sim 64,000$ (preliminary experiments)
- Each $\mathfrak{S}_k$: 2 actions (output, no act.), 2 observations (sense/receive, no obs.)