



A convex approach for Markov chain estimation from aggregate data via inverse optimal transport

Downloaded from: <https://research.chalmers.se>, 2026-09-14 17:59 UTC

Citation for the original published paper (version of record):

Mascherpa, M., Ringh, A., Taghvaei, A. et al (2026). A convex approach for Markov chain estimation from aggregate data via inverse optimal transport. European Journal of Control. <http://dx.doi.org/10.1016/j.ejcon.2026.101592>

N.B. When citing this work, cite the original published paper.



Contents lists available at ScienceDirect

European Journal of Control

journal homepage: www.sciencedirect.com/journal/european-journal-of-control

A convex approach for Markov chain estimation from aggregate data via inverse optimal transport

Michele Mascherpa ^{a,*}, Axel Ringh ^b, Amirhossein Taghvaei ^c, Johan Karlsson ^a

^a Department of Mathematics, KTH Royal Institute of Technology, Stockholm, Sweden

^b Department of Mathematical Sciences, Chalmers University of Technology and University of Gothenburg, Gothenburg, Sweden

^c William E. Boeing Department of Aeronautics and Astronautics, University of Washington, Seattle, WA, USA

ARTICLE INFO

Recommended by Prof. T Parisini

Keywords:

Markov chain estimation
Inverse optimal transport
Schrödinger bridge problem
Convex optimization
Entropic proximal method

ABSTRACT

We address the problem of identifying the dynamical law governing the evolution of a population of indistinguishable particles, when only aggregate distributions at successive times are observed. Assuming a Markovian evolution on a discrete state space, the task reduces to estimating the underlying transition probability matrix from distributional data. We formulate this inverse problem within the framework of entropic optimal transport, as a joint optimization over the transition matrix and the transport plans connecting successive distributions. This formulation results in a convex optimization problem, and we propose an efficient iterative algorithm based on the entropic proximal method. We illustrate the accuracy and convergence of the method in two numerical setups, considering estimation from independent snapshots and estimation from a time series of aggregate observations, respectively.

1. Introduction

Consider a scenario where a large number of indistinguishable particles move according to an unknown dynamical law. Although the individual trajectories of the particles are not accessible, one can observe their aggregate configurations at successive time instances. The fundamental challenge is to identify the underlying—potentially stochastic—law that governs the evolution of these particles using only such aggregate data. This inverse problem arises naturally in diverse domains. In fluid mechanics, for example, Particle Tracking Velocimetry (PTV) techniques rely on the movement of tracer particles illuminated by lasers to infer local flow dynamics (Dabiri & Pecora, 2019). Beyond fluid flow estimation, similar challenges appear in many other applications. For example, in orbital object tracking, where space debris or satellites follow uncertain dynamical laws (Anz-Meador, 2020; Berger et al., 2020), and in swarm robotics, where collective motion patterns of large groups of simple agents must be analyzed and predicted (Rubenstein et al., 2012). Further applications include understanding pedestrian movement from anonymized cell phone data (Willenborg, 2025) or inferring individual-level animal behavior from aggregate ecology data (King, 2013).

Given the indistinguishable nature of the observed particles, it is natural to represent each snapshot as a distribution over the state space. When this space is discrete, the aggregate evolution can be modeled as a sequence of distributions generated by a Markov chain. In this setting,

the underlying stochastic law corresponds to the transition probability matrix that governs how mass is transported between states over time. The key scientific problem, therefore, is to estimate this transition matrix from a sequence of noisy or incomplete distributional measurements.

The estimation of transition probabilities from Markov models is a classical topic (Kemeny et al., 1969; Miller, 1952), and in particular the estimation from aggregate discrete-time observations has been considered in the monograph by Lee et al. (1970), which derives least-squares and maximum-likelihood estimators formulated as quadratic programming problems. For a continuous-time Markov process, the same type of estimators have been studied by Kalbfleisch et al. (1983), whose results have been consolidated by Lawless and McLeish (1984), analyzing the amount of information contained in aggregate data and showing the efficiency loss that arises when individual transitions are unobserved. In Bernstein and Sheldon (2016), the scenario in which data are corrupted by noise is considered, and an estimator based on the method of moments is proposed.

While previous methods are mainly based on least squares or covariance/moment based estimates, we here propose a new likelihood-based method using concepts from optimal transport. The optimal transport problem, which is a classical problem (Villani, 2021), has recently been used to address problems in estimation and control (Chen et al., 2021; Emerick et al., 2024; Haasler et al., 2021a; Mascherpa et al., 2023; Terpin et al., 2024a). In Stuart and Wolfram (2020) the concept of

* Corresponding author.

E-mail addresses: micmas@kth.se (M. Mascherpa), axelri@chalmers.se (A. Ringh), amirtag@uw.edu (A. Taghvaei), johan.karlsson@math.kth.se (J. Karlsson).

<https://doi.org/10.1016/j.ejcon.2026.101592>

Received 18 June 2026; Accepted 22 July 2026

Available online 24 July 2026

0947-3580/© 2026 The Author(s). Published by Elsevier Ltd on behalf of European Control Association. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

inverse optimal transport was introduced, in which, given the marginals, the nonconvex problem of simultaneously finding both the cost and the transport matrices is considered. This concept has been explored in various settings, often motivated by applications in biology. In [Lamoline et al. \(2025\)](#) and [Elvander and Haasler \(2025\)](#), inverse optimal transport problems are used to identify the dynamics given sampling data, and nonconvex problems of jointly estimating dynamics and transport plans are considered. Other related works include ([Swaminathan & Murray, 2014](#)), where Markov-chain identification from distributional data is posed as maximum a posteriori estimation with a multinomial likelihood under dynamic constraints, yielding a non-convex program. Related work ([Yang & Zha, 2023](#)) also considers inverse problems from aggregate marginal observations, but focuses on recovering time-varying optimal transport cost matrices. In contrast, we estimate a single Markov transition matrix shared across all observations, which enables us to analyze identifiability and uniqueness of the recovered transition matrix. Other related approaches use optimal transport or Wasserstein-gradient-flow models to infer state dynamics from population snapshots ([Lavenant et al., 2024](#); [Schiebinger et al., 2019](#); [Terpin et al., 2024b](#); [Tronstad et al., 2025](#)). While related in their use of marginal snapshot data, their aim is to recover couplings, path laws, or energy functionals, rather than a finite-state transition matrix.

Our approach can be seen as a regularized version of inverse optimal transport, where transport plans and transition probabilities are estimated. The resulting optimization problem is jointly convex and admits an efficient proximal algorithm that updates the transport matrices and normalizes their aggregate to recover the transition probabilities. We derive the corresponding dual problem, which we use to provide uniqueness conditions as well as convergence of the method.

The article is organized as follows. In [Section 2](#), we provide background on the Schrödinger bridge problem and its connection to entropic optimal transport. In [Section 3](#), we formulate the joint optimization problem over transport matrices and the transition probability matrix, and analyze existence, uniqueness, and the corresponding dual formulation. In [Section 4](#), we introduce an iterative algorithm for efficiently solving the problem, while [Section 5](#) discusses identifiability issues and presents numerical results under varying data dependencies and noise levels, including a comparison with a least-squares estimator. The paper concludes with a summary of findings in [Section 6](#).

2. Background

2.1. Notation

We denote by $\mathbf{1}_n$ a vector of ones, omitting the dimension n if clear from context. With $\odot$ we indicate entrywise matrix multiplication, and we use $\langle \cdot, \cdot \rangle$ to denote the standard (vector/Frobenius) inner product. The nonnegative orthant of $\mathbb{R}^n$ is denoted with $\mathbb{R}_+^n$, and $[m] := \{1, \dots, m\}$.

2.2. Log-likelihoods for collections of particles in a Markov chain and connection to optimal transport

Consider a collection of indistinguishable particles, where each particle evolves over a finite set of states $X = \{X_1, X_2, \dots, X_n\}$ ([Haasler et al., 2019](#); [Pavon & Ticozzi, 2010](#)). Assume that the underlying probability law is specified by a state transition matrix $A = (a_{ij})_{i,j=1}^n$, where $a_{ij} = P(q_{t+1} = X_j | q_t = X_i)$, and q_t denotes the state of a particle at time t . Now consider the case where we observe the particle distributions $\mu_0, \mu_1 \in \mathbb{R}^n$ before and after a transition, where the i th component $\mu_i(i)$ denotes the number of particles in state X_i at time $t \in \{0, 1\}$. The particle transitions between states can be described by a matrix $M = (m_{ij})_{i,j=1}^n$, where m_{ij} denotes the number of particles that move from state X_i to state X_j . By construction, the mass transition matrix satisfies $\sum_i m_{ij} = \mu_1(j)$ for all $j \in [n]$ and $\sum_j m_{ij} = \mu_0(i)$ for all $i \in [n]$.

Given the initial distribution μ_0 and the transition probability matrix A , the probability that the particles evolve according to a given transi-

tion matrix M is

$$P_{\mu_0, A}(M) = \prod_{i=1}^n \left(\binom{\mu_0(i)}{m_{i1}, m_{i2}, \dots, m_{in}} \prod_{j=1}^n a_{ij}^{m_{ij}} \right), \quad (1)$$

where $\binom{\cdot}{\cdot, \dots, \cdot}$ denotes a multinomial coefficient. When the number of particles grows, the probability of a given transition can be approximated by the relative entropy ([Haasler et al., 2019](#))

$$P_{\mu_0, A}(M) \approx e^{-D(M | \text{diag}(\mu_0)A)},$$

where D is the Kullback-Leibler (KL) divergence.

Definition 1. Let p and q be two nonnegative vectors or matrices of the same dimension. The Kullback-Leibler (KL) functional of p from q is defined as

$$D(p|q) := \sum_i p_i \log \frac{p_i}{q_i}, \quad (2)$$

where $0 \log 0$ is defined to be 0. When p and q have the same total mass, i.e., $\sum_i p_i = \sum_i q_i$, this functional is a divergence and is nonnegative. In this work, we also use D with arguments that may have different total mass, in which case it can be negative.

This formulation allows us to both approximate the negative log-likelihood of (1) in terms of the KL divergence and to consider particles as continuous quantities, that is, $M \in \mathbb{R}_+^{n \times n}$; as the number of particles becomes large, the discrete particle distributions can be approximated by continuous densities. A feasible evolution matrix between the two distributions $\mu_0 \in \mathbb{R}_+^n$ and $\mu_1 \in \mathbb{R}_+^n$ must satisfy the marginal constraints $M\mathbf{1} = \mu_0$ and $M^T\mathbf{1} = \mu_1$. Hence, the most likely evolution matrix M between the two distributions can be approximated by the solution to

$$\min_{M \in \mathbb{R}_+^{n \times n}} D(M | \text{diag}(\mu_0)A) \quad (3)$$

subject to $M\mathbf{1} = \mu_0, \quad M^T\mathbf{1} = \mu_1$.

This problem can naturally be extended to a Markov chain of length T , and can be referred to as a time- and space-discrete form of the Schrödinger bridge problem ([Haasler et al., 2019, 2021b](#); [Mascherpa et al., 2023](#); [Pavon & Ticozzi, 2010](#)).

It should be noted that the formulation (3) is closely connected to the (entropy regularized) optimal transport problem. In fact, if A and μ_0 are strictly positive, then (3) is equivalent to the problem

$$\min_{M \in \mathbb{R}_+^{n \times n}} \langle C, M \rangle + D(M | \mathbf{1}_{n \times n}) \quad (4)$$

subject to $M\mathbf{1} = \mu_0, \quad M^T\mathbf{1} = \mu_1$,

where the cost matrix is $C = -\log(\text{diag}(\mu_0)A)$. A key difference between the two formulations is that $D(M | \text{diag}(\mu_0)A)$ is jointly convex in M and A , whereas $\langle C, M \rangle$ is not jointly convex in M and C . This distinction is consequential in the present setting because A (or equivalently C) is unknown, and our aim is to identify the transition matrix and the transport matrix simultaneously. This is the main reason for adopting the former formulation here.

3. Problem formulation, duality, existence and uniqueness

For a sequence of observed marginal pairs (μ_t, ν_t) , with $t = 1, \dots, \mathcal{T}$, satisfying $\mu_t^T \mathbf{1} = \nu_t^T \mathbf{1}$ for all t , we want to simultaneously find the $\mathcal{T}$ transport matrices $M_t \in \mathbb{R}_+^{n \times n}$ connecting each initial distribution $\mu_t \in \mathbb{R}_+^n$ to its corresponding final distribution $\nu_t \in \mathbb{R}_+^n$, and their common prior $A \in \mathbb{R}_+^{n \times n}$, which is the most likely discrete Markov chain, in the Kullback-Leibler divergence sense, to have generated the marginal observations. Formulated as an optimization problem, this means that we want to solve

$$\min_{\substack{A, M_t \in \mathbb{R}_+^{n \times n} \\ t \in [\mathcal{T}]}} \sum_{t=1}^{\mathcal{T}} D(M_t | \text{diag}(\mu_t)A) \quad (5a)$$

$$\text{subject to } A\mathbf{1} = \mathbf{1} \quad (5b)$$

$$M_t \mathbf{1} = \mu_t \quad \text{for } t \in [\mathcal{T}] \quad (5c)$$

$$M_t^T \mathbf{1} = \nu_t \quad \text{for } t \in [\mathcal{T}]. \quad (5d)$$

Constraint (5b) ensures that the matrix A is row-stochastic, and thus its entries can be interpreted as transition probabilities. Conditions (5c) and (5d) are marginal constraints on the mass transport matrices M_t .

Note that, under the above constraints, we have

$$\begin{aligned} D(M_t | \text{diag}(\mu_t)A) &= \sum_{i,j} M_t(i,j) \log \left(\frac{M_t(i,j)}{\mu_t(i)A(i,j)} \right) \\ &= \sum_{i,j} M_t(i,j) \log \left(\frac{M_t(i,j)}{A(i,j)} \right) - \sum_i \mu_t(i) \log \mu_t(i). \end{aligned}$$

The last term is constant and thus does not affect the optimization. The problem can thus be formulated as

$$\begin{aligned} \min_{\substack{A, M_t \in \mathbb{R}_+^{n \times n} \\ t \in [\mathcal{T}]}} \sum_{t=1}^{\mathcal{T}} D(M_t | A) \\ \text{subject to } A\mathbf{1} = \mathbf{1} \\ M_t \mathbf{1} = \mu_t \quad \text{for } t \in [\mathcal{T}] \\ M_t^T \mathbf{1} = \nu_t \quad \text{for } t \in [\mathcal{T}]. \end{aligned} \quad (6)$$

We first consider the minimization over A , with the transport matrices M_t fixed. Then, the part of the objective function of (6) depending on A is $D(\bar{M} | A)$, where $\bar{M}(i,j) := \sum_{t=1}^{\mathcal{T}} M_t(i,j)$. Since the constraint $A\mathbf{1} = \mathbf{1}$ acts row-wise, the minimization separates over the rows of A . Using standard properties of the Kullback-Leibler divergence, A is the solution where each row is proportional to the corresponding row of $\bar{M}$, i.e.,

$$A(i,j) = \frac{\bar{M}(i,j)}{\sum_{\ell} \bar{M}(i,\ell)}. \quad (7)$$

Plugging (7) into (6), and removing constant terms in the objective function, we obtain the equivalent formulation

$$\begin{aligned} \min_{\substack{\bar{M}, M_t \in \mathbb{R}_+^{n \times n} \\ t \in [\mathcal{T}]}} \sum_{t=1}^{\mathcal{T}} D(M_t | \bar{M}) \\ \text{subject to } M_t \mathbf{1} = \mu_t \quad \text{for } t \in [\mathcal{T}] \\ M_t^T \mathbf{1} = \nu_t \quad \text{for } t \in [\mathcal{T}] \\ \bar{M} = \sum_{t=1}^{\mathcal{T}} M_t, \end{aligned} \quad (8)$$

where A is removed from the problem and can be recovered from $\bar{M}$ via (7). Since (5) and (8) are equivalent, and in particular yield the same optimal transport plans $(M_t)_{t=1}^{\mathcal{T}}$, we will study the latter formulation.

3.1. Existence of an optimal solution

We make the following assumption on the observed marginals μ_t, ν_t .

Assumption 1. For all index pairs $(i,j) \in [n] \times [n]$, there exists $t \in [\mathcal{T}]$ such that $\mu_t(i) > 0$ and $\nu_t(j) > 0$.

Remark 1. If $\mu_t(i) = 0$ or $\nu_t(j) = 0$, then the only feasible value is $M_t(i,j) = 0$. In particular, if Assumption 1 does not hold, i.e., if there is an index pair $(i,j) \in [n] \times [n]$ such that for every $t \in [\mathcal{T}]$ either $\mu_t(i) = 0$ or $\nu_t(j) = 0$, then we have that $M_t(i,j) = 0$ for all t , implying $\bar{M}(i,j) = 0$, and the problem can be restricted to the remaining variables.

Lemma 1. There exists a feasible solution to (8). Moreover, under Assumption 1, there exists a feasible solution such that $\bar{M} > 0$.

Proof. Let $s_t := \mu_t^T \mathbf{1} = \nu_t^T \mathbf{1}$. Consider $M_t := \frac{1}{s_t} \mu_t \nu_t^T$, for all $t \in [\mathcal{T}]$. Then $M_t \geq 0$, $M_t \mathbf{1} = \mu_t$ and $M_t^T \mathbf{1} = \nu_t$ for all t , hence the problem is feasible. Furthermore, under Assumption 1, for every $(i,j) \in [n] \times [n]$ there exists $t \in [\mathcal{T}]$ such that $\mu_t(i) > 0$ and $\nu_t(j) > 0$. Therefore we also have $\bar{M} = \sum_{t=1}^{\mathcal{T}} M_t = \sum_{t=1}^{\mathcal{T}} \frac{1}{s_t} \mu_t \nu_t^T > 0$. $\square$

The next result is that the set of optimal solutions of (8) is non-empty.

Proposition 1. There exists an optimal solution to (8).

Proof. The feasible set of (8) is non-empty by Lemma 1. Furthermore, the feasible set is bounded, as $0 \leq M_t(i,j) \leq \mu_t(i)$ for all $t \in [\mathcal{T}]$ and $i,j \in [n]$, and closed, since it is defined by linear equality and inequality constraints. The objective function is finite and continuous, as it is smooth on the interior, and on the boundary we adopt the usual continuous extension $0 \log \frac{0}{0} := 0$. Note that $M_t(i,j) \log(M_t(i,j)/\bar{M}(i,j)) \rightarrow 0$ as $M_t(i,j), \bar{M}(i,j) \rightarrow 0$ since $M_t(i,j) \leq \bar{M}(i,j)$. By Weierstrass theorem we can then conclude existence of a minimizer. $\square$

3.2. Duality

In order to derive the corresponding dual problem the following lemma is useful.

Lemma 2. Let $\lambda \in \mathbb{R}^{\mathcal{T}}$. The problem

$$\inf_{x \in \mathbb{R}_+^{\mathcal{T}}} \sum_{t=1}^{\mathcal{T}} \left(x_t \log \frac{x_t}{\bar{x}} \right) - \langle x, \lambda \rangle,$$

where $\bar{x} = \mathbf{1}^T x$, has an optimal solution if and only if $\sum_t \exp(\lambda_t) \leq 1$. In this case, the minimum value is 0 and the set of optimal solutions is given by

$$\begin{aligned} \{0\} & \quad \text{if } \sum_{t=1}^{\mathcal{T}} \exp(\lambda_t) < 1, \\ \{x = \alpha \exp(\lambda) \mid \alpha \geq 0\} & \quad \text{if } \sum_{t=1}^{\mathcal{T}} \exp(\lambda_t) = 1. \end{aligned}$$

If $\sum_t \exp(\lambda_t) > 1$, then the objective value tends to $-\infty$ for $x^{(\alpha)} = \alpha \exp(\lambda)$ as $\alpha \rightarrow \infty$.

Proof. Let $\beta := \sum_t \exp(\lambda_t)$, and note that $\beta > 0$. Define $F(x)$

$$\begin{aligned} F(x) &:= \sum_{t=1}^{\mathcal{T}} \left(x_t \log \frac{x_t}{\bar{x}} \right) - \langle x, \lambda \rangle = \sum_{t=1}^{\mathcal{T}} x_t \log \frac{x_t}{\bar{x} \exp(\lambda_t)} \\ &= \sum_{t=1}^{\mathcal{T}} x_t \log \frac{x_t}{\bar{x} e^{\lambda_t} / \beta} - \bar{x} \log \beta. \end{aligned}$$

Next we will use the inequality $z \log(z/y) - z + y \geq 0$, where equality holds if and only if $z = y$, to show that the first term is non-negative. Thus we obtain

$$\sum_{t=1}^{\mathcal{T}} x_t \log \frac{x_t}{\bar{x} e^{\lambda_t} / \beta} \geq \sum_{t=1}^{\mathcal{T}} (x_t - \bar{x} e^{\lambda_t} / \beta) = \bar{x} - \frac{\bar{x}}{\beta} \sum_{t=1}^{\mathcal{T}} e^{\lambda_t} = 0,$$

from which it follows that $F(x) \geq -\bar{x} \log \beta$.

- If $\beta < 1$, then $F(x) > 0$ for all $\bar{x} > 0$, and $F(0) = 0$. Hence $x = 0$ is the unique minimizer.
- If $\beta = 1$, then $F(x) \geq 0$, with equality if and only if $x_t = \bar{x} e^{\lambda_t}$ for all t . Therefore all minimizers are in the form $x = \alpha \exp(\lambda)$, for $\alpha \geq 0$.
- If $\beta > 1$, then $-\bar{x} \log \beta < 0$, and along $x^{(\alpha)} = \alpha \exp(\lambda)$ we have $F(x^{(\alpha)}) = -\alpha \beta \log \beta \rightarrow -\infty$ as $\alpha \rightarrow \infty$, so the problem is unbounded from below.

This completes the proof. $\square$

We use the result of Lemma 2 to derive the Lagrange dual of (8). We begin by expressing the corresponding Lagrangian

$$\begin{aligned} \mathcal{L}(M, \lambda, \rho) &= \sum_{i=1}^T \mathcal{D}(M_i | \bar{M}) + \sum_{i=1}^T \lambda_i^T (\mu_i - M_i \mathbf{1}) \\ &+ \sum_{i=1}^T \rho_i^T (v_i - M_i^T \mathbf{1}) = \sum_{i=1}^T (\lambda_i^T \mu_i + \rho_i^T v_i) \\ &+ \sum_{i,j} \left(M_{i,j} \log \frac{M_{i,j}}{\bar{M}(i,j)} - M_{i,j} (\lambda_i(i) + \rho_i(j)) \right), \end{aligned}$$

where $\lambda_i, \rho_i \in \mathbb{R}^n$, for $t = 1, \dots, T$, are the Lagrange multipliers. Note that for a given element (i, j) , the problem of finding the minimizing $M_{i,j}$, for $t = 1, \dots, T$, is of the form considered in Lemma 2. Therefore, the Lagrangian is bounded from below only if $\sum_i \exp(\lambda_i(i) + \rho_i(j)) \leq 1$, and the corresponding minimizers $M_{i,j}$ are given by

$$M_{i,j}^* = \bar{M}^*(i, j) \exp(\lambda_i(i) + \rho_i(j)). \quad (9)$$

In addition $M_{i,j}^* = \bar{M}^*(i, j) = 0$ for all t if $\sum_i \exp(\lambda_i(i) + \rho_i(j)) < 1$. The dual problem then becomes

$$\sup_{\lambda_i, \rho_i \in \mathbb{R}^n, i \in [T]} \sum_{i=1}^T (\lambda_i^T \mu_i + \rho_i^T v_i) \quad (10a)$$

$$\text{subject to } \sum_{i=1}^T \exp(\lambda_i(i) + \rho_i(j)) \leq 1, \quad \forall i, j \in [n]. \quad (10b)$$

If $\mu_i > 0$ and $v_i > 0$ for $t \in [T]$, then $M_t = \mu_i v_i^T / (\mu_i^T \mathbf{1})$ yield a feasible point with $M_t > 0$ for all t , and hence $\bar{M} > 0$. Therefore, by Slater's condition, strong duality holds. If some entry of μ_i or v_i is zero, then the corresponding row or column of M_t is necessarily zero and can be removed without changing the problem. In that case, the dual optimum in (10) is not attained, but the duality gap is zero, and it is approached in the limit as the dual variables associated with those entries tend to $-\infty$. For the uniqueness results in the following section we will assume that $\mu_i > 0$ and $v_i > 0$ for all $t \in [T]$.

3.3. Uniqueness

Problem (8) is convex, but not necessarily strictly convex, so its optimal solution may not be unique. We now characterize the structure of the optimal solution set of (8) and derive a necessary and sufficient condition for uniqueness.

Proposition 2 (Characterization of uniqueness). *Let $((M_i^*)_{i=1}^T, \bar{M}^*)$ be an optimal solution of (8) and assume that $\mu_i > 0$ and $v_i > 0$ for all $t \in [T]$. Let $u_i := \exp(\lambda_i)$ and $v_i := \exp(\rho_i)$ be the associated dual scalings, where $(\lambda_i, \rho_i)_{i=1}^T$ are optimal dual variables, and denote by $\mathcal{I} = \{(i, j) \mid \sum_i u_i(i) v_i(j) < 1\}$ the index pairs for which the dual constraints are inactive. Then, the optimal solution is the unique solution if and only if the only solution $X \in \mathbb{R}^{n \times n}$ to the system of equations*

$$(X \odot (u_i v_i^T)) \mathbf{1} = 0, \quad t \in [T], \quad (11a)$$

$$(X \odot (u_i v_i^T))^T \mathbf{1} = 0, \quad t \in [T], \quad (11b)$$

$$X(i, j) = 0 \text{ for } (i, j) \in \mathcal{I}, \quad (11c)$$

$$X(i, j) \geq 0 \text{ for all } (i, j) \text{ where } \bar{M}^*(i, j) = 0, \quad (11d)$$

is the trivial solution $X = 0$.

Proof. Exploiting the convexity of the optimal solution set, we consider optimal solutions of the form $M_i^* + \delta M_i$ to give conditions for uniqueness. Let (λ, ρ) be optimal dual variables. By strong duality, any primal optimal solution of (8) minimizes the Lagrangian with (λ, ρ) fixed. Fix an index pair (i, j) , and denote $m_i := M_{i,j}$, $\bar{m} = \sum_i m_i$. The (i, j) th component of the Lagrangian of (8) reads

$$\mathcal{L}_{ij}(m) = \sum_{i=1}^T [m_i (\log(m_i / \bar{m}))] - \sum_{i=1}^T (\lambda_i(i) + \rho_i(j)) m_i.$$

Minimizing $\mathcal{L}_{ij}$ with respect to m yields the same problem as in Lemma 2, with λ_i replaced by $\lambda_i(i) + \rho_i(j)$. Hence, for the entries $(i, j) \in \mathcal{I}$ corresponding to an inactive dual constraint, the set of minimizers is $\{0\}$. Therefore, $M_{i,j}^* = 0$ is the unique solution for $(i, j) \in \mathcal{I}$.

Consider now the entries $(i, j) \notin \mathcal{I}$, corresponding to active dual constraints, i.e., (i, j) for which $\sum_i u_i(i) v_i(j) = 1$. Then, the entire ray

$$\{\alpha (u_i(i) v_i(j))_{i=1}^T : \alpha \geq 0\}$$

gives cost-equivalent primal values. Collecting all entries together, an infinitesimal perturbation that preserves optimality in every active entry can then be written as

$$\delta M_t = X \odot (u_t v_t^T), \quad t \in [T],$$

for a matrix $X \in \mathbb{R}^{n \times n}$ with $X_{ij} = 0$ if $(i, j) \in \mathcal{I}$. Feasibility with respect to the marginal constraints $M_t \mathbf{1} = \mu_t$ and $M_t^T \mathbf{1} = v_t$ imposes

$$(\delta M_t) \mathbf{1} = 0, \quad (\delta M_t)^T \mathbf{1} = 0, \quad t \in [T].$$

Thus, we obtain the linear feasibility system

$$(X \odot (u_t v_t^T)) \mathbf{1} = 0, \quad t \in [T],$$

$$(X \odot (u_t v_t^T))^T \mathbf{1} = 0, \quad t \in [T].$$

Finally, the positivity constraint $M \geq 0$ imposes that, if $\bar{M}^*(i, j) = 0$, it must also hold $X(i, j) \geq 0$. $\square$

The condition that $X = 0$ is the only solution to (11) can be expressed directly in terms of the dual scaling vectors.

Corollary 1 (Dual characterization). *Let $(u_i, v_i)_{i=1}^T$ be the dual scaling vectors at the optimum, with $u_i = \exp(\lambda_i)$ and $v_i = \exp(\rho_i)$. Then a sufficient condition for a solution M^* of problem (8) to be unique is that*

$$\text{span}\{v_i\}_{i=1}^T = \mathbb{R}^n \quad \text{or} \quad \text{span}\{u_i\}_{i=1}^T = \mathbb{R}^n.$$

Furthermore, if $\bar{M}^* > 0$, the condition is also necessary.

Proof. To prove sufficiency, note that if either $\{v_i\}_i$ or $\{u_i\}_i$ spans $\mathbb{R}^n$, then from (11a) or (11b) it follows that $X(i, :) = 0$ or $X(:, j) = 0$ for all i, j , hence $X = 0$, proving uniqueness. To prove necessity, assume that $\bar{M}^*(i, j) > 0$ for all (i, j) . Now, assume that both families have deficient span, i.e. there exist nonzero vectors

$$a \in \text{span}\{u_i\}_{i=1}^T \perp, \quad b \in \text{span}\{v_i\}_{i=1}^T \perp.$$

Define $X \in \mathbb{R}^{n \times n}$ by $X(i, j) = a_i b_j$. For each t ,

$$[(X \odot (u_t v_t^T)) \mathbf{1}]_i = u_t(i) a_i \langle b, v_t \rangle = 0,$$

$$[(X \odot (u_t v_t^T))^T \mathbf{1}]_j = v_t(j) b_j \langle a, u_t \rangle = 0,$$

since $a \perp u_t$ and $b \perp v_t$ for all t . Thus $X \neq 0$ satisfies (11a) and (11b). Under the positivity assumption $\bar{M}^* > 0$, also (11c) and (11d) are satisfied. Therefore, we have constructed a nonzero solution of (11). By Proposition 2, the solution is then not unique. Hence, if the solution is unique, at least one of the families $\{u_i\}_{i=1}^T$ or $\{v_i\}_{i=1}^T$ must span $\mathbb{R}^n$, proving necessity. $\square$

The dual characterization provides a compact and practically verifiable condition, with uniqueness holding if the dual scaling vectors u_i or v_i span $\mathbb{R}^n$. We note that for this condition to be satisfied it must hold $T \geq n$. The following result presents a primal interpretation of Corollary 1.

Corollary 2 (Primal characterization). *Let M^* be an optimal solution to problem (8), and assume that $\bar{M}^* > 0$. Then M^* is the unique optimal solution if and only if, for every i (equivalently, for every j),*

$$\text{span}\{M_i^*(i, :)\}_{i=1}^T = \mathbb{R}^n \quad \text{or} \quad \text{span}\{M_i^*(:, j)\}_{i=1}^T = \mathbb{R}^n,$$

i.e., each of the rows (or the columns) of $(M_i^*)_{i=1}^T$ spans the whole space.

Proof. By optimality condition (9),

$$M_i^*(i, j) = \bar{M}^*(i, j) u_i(i) v_i(j).$$

Fix i . For each t ,

$$M_t^*(i, :) = u_t(i) (\bar{M}^*(i, :) \odot v_t^T),$$

so all the row vectors $\{M_t^*(i, :)\}_t$ are obtained from $\{v_t\}_t$ by diagonal scaling with the fixed positive vector $\bar{M}^*(i, :)$. Since scaling by a fixed diagonal matrix with strictly positive entries preserves span, we get

$$\text{span}\{M_t^*(i, :)\}_{t=1}^T = \text{span}\{v_t\}_{t=1}^T.$$

A symmetric argument with the columns gives

$$\text{span}\{M_t^*(:, j)\}_{t=1}^T = \text{span}\{u_t\}_{t=1}^T,$$

and the uniqueness result follows from [Corollary 1](#). $\square$

Remark 2. Linear independence of the dual scalings (u_t, v_t) reflects the marginals (μ_t, ν_t) . If two observations are proportional, $(\mu_{t_1}, \nu_{t_1}) = c(\mu_{t_2}, \nu_{t_2})$, then $u_{t_1} \propto u_{t_2}$ and $v_{t_1} \propto v_{t_2}$, so the second observation adds no information. The precise link between the rank of (u_t, v_t) and that of (μ_t, ν_t) remains an open question.

4. Numerical method and convergence analysis

The main idea for solving the optimization problem (8) is to iteratively update the transport plans and transition probability matrix. Updating the transport plans consists of computing the solution of $\mathcal{T}$ entropy regularized optimal transport problems, which can be solved efficiently (and in parallel) with Sinkhorn iterations ([Cuturi, 2013](#); [Peyré & Cuturi, 2019](#)). The prior estimate is then updated as the sum of the optimal mass transport matrices, and the procedure iterated. This is detailed in [Algorithm 1](#).

Algorithm 1 Proximal iterative scheme for (8).

```

 $M_t \leftarrow \mu_t \nu_t^T / \mu_t^T \mathbf{1}, X \leftarrow \sum_{t=1}^T M_t$ 
while change in  $X$  above tolerance do
   $M_t \leftarrow \arg \min_{M_t \in \mathbb{R}_+^{n \times n}} D(M_t | X)$ 
  subject to  $M_t \mathbf{1} = \mu_t, M_t^T \mathbf{1} = \nu_t, t \in [\mathcal{T}]$ 
   $X \leftarrow \sum_{t=1}^T M_t$ 
end while
 $A \leftarrow \text{diag}(X\mathbf{1})^{-1} X$ 

```

A way to view this algorithm, and to show its convergence, is to consider it as an instance of the entropic proximal method ([Teboulle, 1992](#)), with Kullback-Leibler divergence (2) as the penalization. The method allows to minimize a closed, convex, and proper function $f : \mathbb{R}^m \rightarrow (-\infty, +\infty]$, by generating the sequence $\{x^k\}$, starting from an initial strictly positive point $x^0 > 0$, according to the update rule

$$x^{k+1} = \arg \min_{x \in \mathbb{R}_+^m} \{f(x) + \varepsilon_k D(x | x^k)\}, \quad (12)$$

for a sequence of positive parameters $\{\varepsilon_k\}$. The next proposition shows convergence of [Algorithm 1](#), and its connection to the proximal iterations (12).

Proposition 3. Under [Assumption 1](#), [Algorithm 1](#) converges to an optimal solution $\bar{M}^*$ of (8).

Proof. Consider the function

$$\begin{aligned} f(X) &= \inf_{\substack{M_t \in \mathbb{R}_+^{n \times n} \\ t \in [\mathcal{T}]}} \sum_{t=1}^T D(M_t | X) \\ \text{subject to } M_t \mathbf{1} &= \mu_t & \text{for } t \in [\mathcal{T}] \\ M_t^T \mathbf{1} &= \nu_t & \text{for } t \in [\mathcal{T}] \\ X &= \sum_{t=1}^T M_t. \end{aligned} \quad (13)$$

The function f corresponds to problem (8), with the prior X being considered as the variable. First, observe that f is closed, convex and

proper. Under [Assumption 1](#), we have seen ([Lemma 1](#)) that there exists a strictly positive feasible point, so that $\text{dom} f \cap \mathbb{R}_{>0}^{n \times n} \neq \emptyset$. In addition, [Proposition 1](#) guarantees that the optimal solution set is non-empty, thus $\{X : f(X) = \inf_{X \in \mathbb{R}_+^{n \times n}} f\} \neq \emptyset$. With the parameter choice $\varepsilon_k \equiv 1$, the entropic proximal point algorithm converges according to [Teboulle \(1997, Theorem 4.3\)](#), with the objective gap decaying as $O(1/K)$, denoting with K the number of iterations. The iterations of (12), with f defined as (13), become

$$\begin{aligned} \arg \min_{X \in \mathbb{R}_+^{n \times n}} \inf_{\substack{M_t \in \mathbb{R}_+^{n \times n} \\ t \in [\mathcal{T}]}} \sum_{t=1}^T D(M_t | X) + D(X | X^k) \\ \text{subject to } M_t \mathbf{1} = \mu_t, \quad M_t^T \mathbf{1} = \nu_t, \quad \text{for } t \in [\mathcal{T}] \\ X = \sum_{t=1}^T M_t. \end{aligned}$$

Using the constraint $X = \sum_{t=1}^T M_t$, the objective can be equivalently expressed as $\sum_{t=1}^T D(M_t | X^k)$. Thus, the variable X is not part of the objective function, and will be determined by the optimal mass transport matrices M_t via the constraint $X = \sum_t M_t$. The updates become

$$\begin{aligned} M^* &\leftarrow \begin{cases} \arg \min_{M_t \in \mathbb{R}_+^{n \times n}} \sum_{t=1}^T D(M_t | X^k), \\ \text{s.t. } M_t \mathbf{1} = \mu_t, \quad M_t^T \mathbf{1} = \nu_t, \quad t \in [\mathcal{T}], \end{cases} \\ X^* &\leftarrow \sum_{t=1}^T M_t^*, \end{aligned}$$

which correspond to [Algorithm 1](#). $\square$

Remark 3. As it is stated, [Algorithm 1](#) requires to solve, at each iteration, $\mathcal{T}$ bi-marginal entropy regularized OT problems. This computational cost can be greatly alleviated by performing only a few sweeps of the Sinkhorn iterations in the inner loop, often just one. In particular, if each minimization in (12) is performed with an accuracy η_k , and $\sum_{k=1}^\infty \varepsilon_k \eta_k < \infty$, then the inexact iterations converge to an optimal solution ([Teboulle, 1997, Theorem 4.3](#)). Recent works have specialized this result for optimal transport and some of its related formulations ([Chen et al., 2024](#); [Xie et al., 2020](#); [Yang & Toh, 2022](#)), giving sufficient descent conditions for convergence.

5. Applications

A typical application of this problem is when the marginal distributions are derived from independent Markov chain trajectories. In particular, if $\{X_t^i\}_{i=1}^N$ are N independent trajectories of a Markov chain with n states $\{1, \dots, n\}$ and Markov transition matrix A , define

$$\mu_t(k) = \frac{1}{N} \sum_{i=1}^N \mathbf{1}_{X_t^i=k}, \quad \text{for } k \in [n], \quad (14)$$

and $\nu_t = \mu_{t+1}$. For finite but large N , we have

$$\nu_t = A^T \mu_t + \text{"noise"}, \quad (15)$$

where the "noise" part becomes zero when $N \rightarrow \infty$. In this section, we numerically investigate the estimation of the Markov transition matrix A from such measurements. First, in [Section 5.1](#), we consider independent measurements, where multiple Markov transitions of length 2 are considered. Then, in [Section 5.2](#), we consider sequential observations where each pair is of the form (μ_t, μ_{t+1}) for $t \in [\mathcal{T}]$. In [Section 5.3](#) we discuss the identifiability issue in connection with the excitation of the states. Finally, in [Section 5.4](#) we compare the proposed methodology with a constrained least-squares estimator. In all experiments, the outer iterations stopping tolerance has been set to 10^{-6} , while the inner Sinkhorn updates were performed inexactly, with one step per outer iteration.

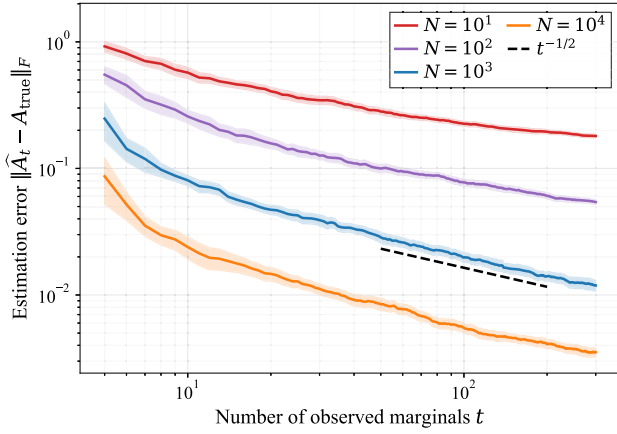


Fig. 1. Numerical results for the application of Algorithm 1 on the independent observation case of Section 5.1, for A given by (16). The plot shows estimation error as a function of observed marginals (in log-log scale), for four values of the particle number N .

5.1. Independent observations

To illustrate Algorithm 1, we consider the matrix A , arising from the estimation from flow cytometry data in Swaminathan and Murray (2014),

$$A = \begin{bmatrix} 0.48 & 0.50 & 0 & 0.02 & 0 \\ 0.33 & 0.27 & 0 & 0.40 & 0 \\ 0 & 0 & 0 & 0.54 & 0.46 \\ 0.26 & 0 & 0.45 & 0.29 & 0 \\ 0 & 0 & 0.51 & 0 & 0.49 \end{bmatrix}. \quad (16)$$

We simulate a scenario where the Markov chain A produces $\mathcal{T}$ marginal pairs (μ_t, ν_t) , corresponding to two consecutive trajectory observations as in (14). In particular, for each t , we sample μ_t from a random uniform distribution, and then ν_t is constructed according to (14). Thus, each pair represents distinct realizations of the same process with independent initial condition, ensuring excitation of all the states of the system. Increasing finite values of N are used to model varying noise levels: as $N \rightarrow \infty$ we reach the transition $\nu = A^T \mu$.

We test the recovery of the true transition matrix A as $\mathcal{T}$ increases, using the Frobenius norm between A and its estimate (7); see Fig. 1. The system has 5 states, and an increasing number of observation pairs $\mathcal{T} \in \{5, \dots, 300\}$ is considered. For each value of N , a total of $K = 30$ simulations with different observation pairs have been averaged. The shaded area is the 95% confidence interval of the mean error across the repeats.

After an initial transient due to early noise, we observe that, when N is large enough for the system to approach the regime (15), the error converges at a rate close to $t^{-1/2}$, typical of consistent estimators with increasing i.i.d. data (Van der Vaart, 2000, Ch. 5).

5.2. Sequential observation

We now consider a scenario in which the number of particles N is small, and the marginal distributions correspond to consecutive observations of particle trajectories as in (14), so that each observation pair is in the form (μ_t, μ_{t+1}) , for $t \in [\mathcal{T}]$.

The results are shown in Fig. 2. We consider a number of marginals in the range $\{40, \dots, 2000\}$. The best performance is achieved with $N = 2$ particles: the reduced number of particles introduces higher stochasticity, making the transitions closer to independent. On the other hand, with few observations, a small number of particles such as $N = 2$ may not explore all states, whereas larger values of N yield a solution even

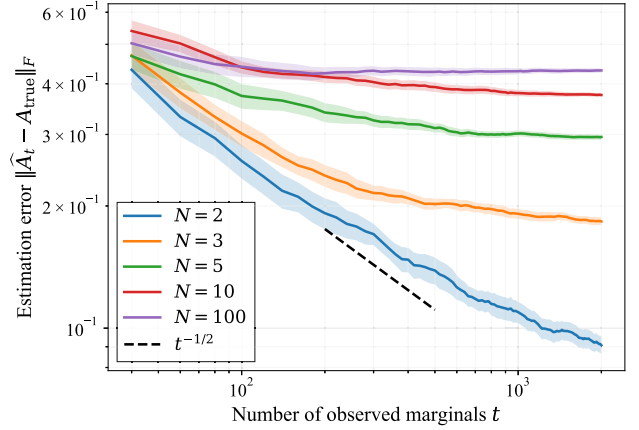


Fig. 2. Numerical results for the application of Algorithm 1 on the sequential observation case of Section 5.2, for A given by (16). The plot shows estimation error as a function of observed marginals (in log-log scale), for five values of the particle number N .

for smaller $\mathcal{T}$, but they do not provide enough excitation to the system, and their estimation error reaches a nearly constant level. This phenomenon will be further discussed in the next subsection.

5.3. Discussion on identifiability

If the matrix A is irreducible and aperiodic, we observe exponential convergence of the marginal observations (14) to the unique stationary distribution π of A (Levin & Peres, 2017),

$$d(t) := \sup_{\mu_0} \|(A^t)^T \mu_0 - \pi\|_{TV} \leq C \alpha^t,$$

for $C > 0$, while the coefficient $\alpha \in (0, 1)$ is related to the second largest (in magnitude) eigenvalue of A , and $\|\mu - \nu\|_{TV} := \frac{1}{2} \sum_i |\mu(i) - \nu(i)|$ is the total variation norm. It is then possible to define the mixing time

$$t_{\text{mix}}(\epsilon) := \min\{t : d(t) \leq \epsilon\},$$

at which the distance to stationarity is ϵ -small. As a consequence, approaching the mixing time, any two Markov chains A and B that share the same stationary distribution π will lead to very similar trajectories, $(A^t)^T \mu_0 \approx (B^t)^T \mu_0 \approx \pi$. Thus after a few time steps, the observations reduce to (noisy) draws from the stationary distribution π , and it will not be possible to distinguish among the elements of the set

$$\mathcal{A}(\pi) := \{A \in \mathbb{R}_+^{n \times n} : A\mathbf{1} = \mathbf{1}, A^T \pi = \pi\}.$$

The convex optimization problem (6) will then pick a (possibly unique, if the conditions in Proposition 2 are satisfied) optimal solution in $\mathcal{A}(\pi)$, corresponding to the best fit to near-stationary, noisy marginals. This behavior is observed in Fig. 2: as N increases, the system receives insufficient excitation, and the estimation error stabilizes to an almost constant level. In particular, for $N = 100$, the estimate of A at $\mathcal{T} = 2000$, averaged across the repeats, is

$$\hat{A} = \begin{bmatrix} 0.39 & 0.37 & 0.07 & 0.11 & 0.06 \\ 0.30 & 0.23 & 0.08 & 0.32 & 0.07 \\ 0.08 & 0.05 & 0.10 & 0.41 & 0.35 \\ 0.24 & 0.06 & 0.34 & 0.27 & 0.09 \\ 0.08 & 0.04 & 0.39 & 0.12 & 0.37 \end{bmatrix}. \quad (17)$$

Denoting with $\pi, \hat{\pi}$, the stationary distributions of $A, \hat{A}$, respectively, we observe that $\|\pi - \hat{\pi}\|_{TV} \approx 0.001$. Therefore, the method selected an element $\hat{A} \in \mathcal{A}(\pi)$ (up to numerical approximation), whose entries are all positive and well separated from zero. This discrepancy likely originates from the zero entries of (16), which are penalized by the maximum-entropy formulation that favors a more uniform distribution of mass,

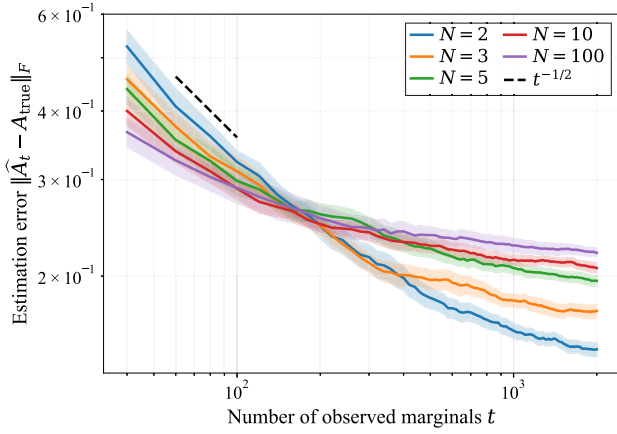


Fig. 3. Numerical results for the application of Algorithm 1 on the sequential observation case of Section 5.2, with random positive A . The plot shows estimation error as a function of observed marginals (in log-log scale), for five values of the particle number N .

as in (17). We test this hypothesis by repeating the experiment with a randomly generated, positive A . The results are reported in Fig. 3. We observe better convergence properties than in Fig. 2, although slower than the $-\frac{1}{2}$ rate observed with independent data. If a sparsity pattern is known a priori, it can be incorporated directly into (5) by restricting the support of A and of the transport plans, i.e., by imposing $A_{ij} = 0$ and $M_t(i, j) = 0$ on the forbidden entries.

To correctly identify the underlying transition probability matrix A , it is then necessary that the observations contain sufficient excitation. As $\|\mu_t - \pi\|$ decays exponentially with time (in the noise-free case $N \rightarrow \infty$), an effective way of improving identifiability is to collect several independent episodes with different initial data, analogous to the setup in Section 5.1. Furthermore, when the number of particles N is finite, it also offers a trade-off between noise and exploration of A . Low values of N lead to higher stochasticity, which enhances exploration and makes the marginal distributions more informative, at the cost of a less accurate measurement of A .

5.4. Comparison with a least-squares baseline

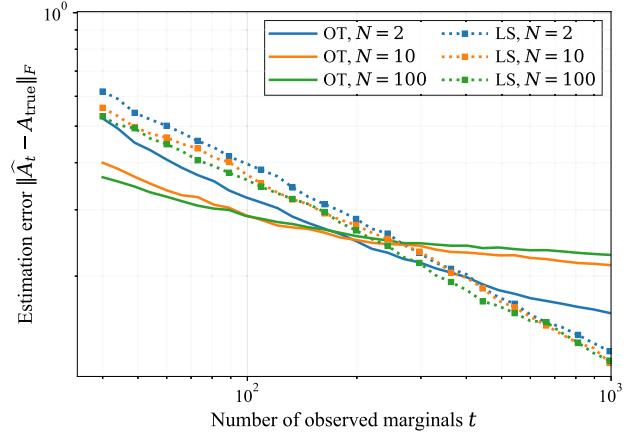
We complement the previous experiments with a comparison against a constrained least-squares estimator, proposed in Lee et al. (1970). Given $\mathcal{T}$ aggregate observations (μ_t, v_t) , we consider the following optimization problem

$$\min_{A \in \mathbb{R}^{n \times n}} \sum_{t=1}^{\mathcal{T}} \|\mu_t - A^T \mu_t\|_2^2 \quad (18)$$

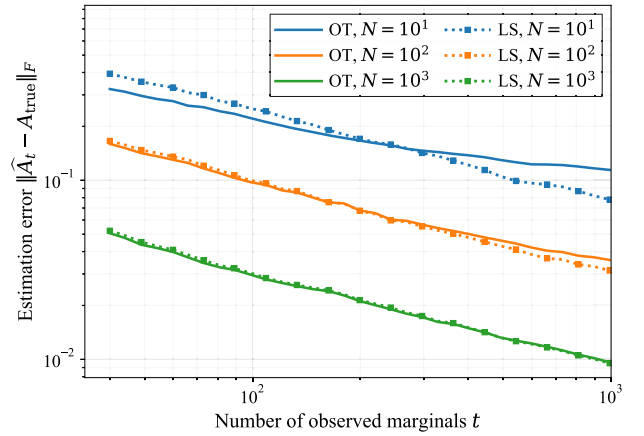
subject to $A\mathbf{1} = \mathbf{1}, \quad A \geq 0$.

Problem (18) is a convex quadratic program and is solved using CVXPY (Diamond & Boyd, 2016) with an off-the-shelf QP solver. In this comparison, we use the same random dense positive transition matrix A as was used in Section 5.3 for which the error estimates were presented in Fig. 3. We compare the two estimators with sequential observations (i.e., $v_t = \mu_{t+1}$) and with independently generated pairs of marginals, reporting the average over 30 observation scenarios.

The results are reported in Fig. 4. In both settings, the proposed methodology gives smaller reconstruction error for small values of $\mathcal{T}$. As the number of observed marginal pairs increases, the least-squares estimator eventually attains a lower reconstruction error. The difference is particularly prominent in the *sequential observation* scenario, while with independent observations, as N increases, the two estimators have a closer performance.



(a) Sequential observations.



(b) Independent observations.

Fig. 4. Comparison between the proposed estimator and the constrained least-squares estimator for a randomly generated dense transition matrix. The plots show estimation error as a function of observed marginal pairs (in log-log scale), for three values of the particle number N .

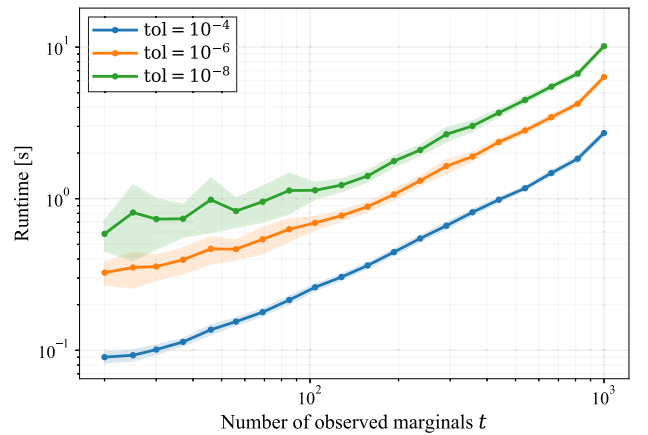


Fig. 5. Runtime of the OT estimator in the independent-observation setting with $N = 100$, averaged over 30 scenarios.

This comparison illustrates the different behavior induced by the two formulations. While the least-squares baseline directly fits the empirical transition relations $v_i \approx A^T \mu_i$, which are obtained from the particle dynamics up to sampling fluctuations, the proposed formulation instead estimates A through an entropy-based inverse optimal transport problem, and simultaneously reconstructs also the optimal mass transportation matrices M_i .

We also evaluate the computational scaling of the proposed OT estimator, by reporting in Fig. 5 the runtime for the independent observation setting with $N = 100$ particles per marginal and an increasing number $\mathcal{T}$ of marginals, for different stopping tolerances. The same dense transition matrix as in the previous experiments has been used. The reported runtimes are averaged over 30 observation scenarios. The simulation was performed on a laptop with an Intel Core i7-1165G7 CPU and 16 GB of RAM, using Python 3.9. Fig. 5 shows that the problem can be solved within seconds on a standard laptop computer with high precision even when the number of marginals is large. As expected, the runtime increases with the number of observed marginals $\mathcal{T}$, reflecting the increasing size of the optimization problem. Stricter stopping tolerances lead to longer runtimes, although the dependence on $\mathcal{T}$ appears similar across tolerances.

6. Conclusions

In this work, we addressed the problem of estimating a Markov chain from aggregate observations by formulating it as an entropy-regularized optimal transport problem, where both the cost function and the transport matrices are treated as optimization variables under marginal constraints given by the data. We proposed an entropic proximal algorithm that alternately updates the transport plans and the transition matrix, showing with numerical experiments that the method successfully retrieves the true transition probabilities, provided sufficient excitation of the system.

Future research directions include a theoretical investigation of identifiability, as well as providing a proof of consistency and establishing the convergence rate observed empirically under independent observations. Another line of work concerns deriving conditions for uniqueness of the optimal solution based on the properties of the observed data rather than on the solution itself. Possible generalizations include partial model identification, where either prior information on the transition structure is incorporated or certain entries of the transition matrix are constrained or fixed, as well as investigating extensions to hidden Markov models, cf. (Singh et al., 2022).

CRedit authorship contribution statement

Michele Mascherpa: Writing – original draft, Writing – review & editing, Conceptualization, Methodology; **Axel Ringh:** Writing – original draft, Writing – review & editing, Conceptualization, Methodology; **Amirhossein Taghvaei:** Writing – original draft, Writing – review & editing, Conceptualization, Methodology; **Johan Karlsson:** Conceptualization, Methodology, Writing – original draft, Writing – review & editing.

Declaration of generative AI and AI-assisted technologies in the manuscript preparation process

During the preparation of this work the authors used ChatGPT to assist with text refinement. After using this tool/service, the authors reviewed and edited the content as needed and take full responsibility for the content of the published article.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgment

This work was supported by KTH Digital Futures; the Wallenberg AI, Autonomous Systems and Software Program (WASP) funded by the Knut and Alice Wallenberg Foundation, Sweden; the Swedish Research Council (VR) [grant number 2020–03454]; and the National Science Foundation (NSF) [grant numbers EPCN-2318977, EPCN-2347358].

References

- Anz-Meador, P. D. (2020). Orbital debris quarterly news. *Orbital Debris Quarterly News (ODQN)*, 24(1), 1–16. NASA report JSC-E-DAA-TN77633.
- Berger, T. E., Holzinger, M. J., Sutton, E. K., & Thayer, J. P. (2020). Flying through uncertainty. *Space Weather*, 18(1), e2019SW002373.
- Bernstein, G., & Sheldon, D. (2016). Consistently estimating Markov chains with noisy aggregate data. In *Artificial intelligence and statistics* (pp. 1142–1150). PMLR.
- Chen, X., Wang, F., Liu, J., & Cui, L. (2024). An inexact Bregman proximal point method and its acceleration version for unbalanced optimal transport. arXiv:2402.16978
- Chen, Y., Georgiou, T. T., & Pavon, M. (2021). Optimal transport in systems and control. *Annual Review of Control, Robotics, and Autonomous Systems*, 4(1), 89–113.
- Cuturi, M. (2013). Sinkhorn distances: Lightspeed computation of optimal transport. *Advances in Neural Information Processing Systems*, 26, 2292–2300.
- Dabiri, D., & Pecora, C. (2019). Particle tracking velocimetry. IOP Publishing.
- Diamond, S., & Boyd, S. (2016). CVXPY: A Python-embedded modeling language for convex optimization. *Journal of Machine Learning Research*, 17(83), 1–5.
- Elvander, F., & Haasler, I. (2025). Mixtures of ensembles: System separation and identification via optimal transport. *IEEE Control Systems Letters*, 9, 1646–1651.
- Emerick, M., Jonas, J., & Bamieh, B. (2024). Causal tracking of distributions in Wasserstein space: A model predictive control scheme. In *2024 IEEE 63rd conference on decision and control (CDC)* (pp. 7606–7611). IEEE.
- Haasler, I., Karlsson, J., & Ringh, A. (2021a). Control and estimation of ensembles via structured optimal transport. *IEEE Control Systems Magazine*, 41(4), 50–69.
- Haasler, I., Ringh, A., Chen, Y., & Karlsson, J. (2019). Estimating ensemble flows on a hidden Markov chain. In *2019 IEEE 58th conference on decision and control (CDC)* (pp. 1331–1338). IEEE.
- Haasler, I., Ringh, A., Chen, Y., & Karlsson, J. (2021b). Multimarginal optimal transport with a tree-structured cost and the Schrödinger bridge problem. *SIAM Journal on Control and Optimization*, 59(4), 2428–2453.
- Kalbfleisch, J. D., Lawless, J. F., & Vollmer, W. M. (1983). Estimation in Markov models from aggregate data. *Biometrics*, 39(4), 907–919.
- Kemeny, J. G., Snell, J. L. et al. (1969). Finite Markov chains (vol. 26). van Nostrand, Princeton, NJ.
- King, G. (2013). A solution to the ecological inference problem: Reconstructing individual behavior from aggregate data. Princeton University Press.
- Lamoline, F., Haasler, I., Karlsson, J., Gonçalves, J., & Aalto, A. (2025). Dynamic gene regulatory network inference from single-cell data using optimal transport. *Bioinformatics*, 41(8), btaf394.
- Lavenant, H., Zhang, S., Kim, Y.-H., Schiebinger, G. et al. (2024). Toward a mathematical theory of trajectory inference. *The Annals of Applied Probability*, 34(1A), 428–500.
- Lawless, J. F., & McLeish, D. L. (1984). The information in aggregate data from Markov chains. *Biometrika*, 71(3), 419–430.
- Lee, T.-C., Judge, G. G., & Zellner, A. (1970). Estimating the parameters of the Markov probability model from aggregate time series data (vol. 65). Contributions to Economic Analysis. Amsterdam: North-Holland Publishing Company.
- Levin, D. A., & Peres, Y. (2017). *Markov chains and mixing times* (vol. 107). American Mathematical Soc.
- Mascherpa, M., Haasler, I., Ahlgren, B., & Karlsson, J. (2023). Estimating pollution spread in water networks as a Schrödinger bridge problem with partial information. *European Journal of Control*, 74, 100846.
- Miller, G. A. (1952). Finite Markov processes in psychology. *Psychometrika*, 17(2), 149–167.
- Pavon, M., & Ticozzi, F. (2010). Discrete-time classical and quantum Markovian evolutions: Maximum entropy problems on path space. *Journal of Mathematical Physics*, 51(4), 042104.
- Peyré, G., & Cuturi, M. (2019). Computational optimal transport: With applications to data science. *Foundations and Trends® in Machine Learning*, 11(5–6), 355–607.
- Rubenstein, M., Ahler, C., & Nagpal, R. (2012). Kilobot: A low cost scalable robot system for collective behaviors. In *2012 IEEE International conference on robotics and automation* (pp. 3293–3298). IEEE.
- Schiebinger, G., Shu, J., Tabaka, M., Cleary, B., Subramanian, V., Solomon, A., Gould, J., Liu, S., Lin, S., Berube, P. et al. (2019). Optimal-transport analysis of single-cell gene expression identifies developmental trajectories in reprogramming. *Cell*, 176(4), 928–943.
- Singh, R., Zhang, Q., & Chen, Y. (2022). Learning hidden Markov models from aggregate observations. *Automatica*, 137, 110100.
- Stuart, A. M., & Wolfram, M.-T. (2020). Inverse optimal transport. *SIAM Journal on Applied Mathematics*, 80(1), 599–619.
- Swaminathan, A., & Murray, R. M. (2014). Identification of Markov chains from distributional measurements and applications to systems biology. *IFAC Proceedings Volumes*, 47(3), 4400–4405.

- Teboulle, M. (1992). Entropic proximal mappings with applications to nonlinear programming. *Mathematics of Operations Research*, 17(3), 670–690.
- Teboulle, M. (1997). Convergence of proximal-like algorithms. *SIAM Journal on Optimization*, 7(4), 1069–1083.
- Terpin, A., Lanzetti, N., & Dörfler, F. (2024a). Dynamic programming in probability spaces via optimal transport. *SIAM Journal on Control and Optimization*, 62(2), 1183–1206.
- Terpin, A., Lanzetti, N., Gadea, M., & Dörfler, F. (2024b). Learning diffusion at lightspeed. *Advances in Neural Information Processing Systems*, 37, 6797–6832.
- Tronstad, M., Karlsson, J., & Dahlin, J. S. (2025). MultistageOT: Multistage optimal transport infers trajectories from a snapshot of single-cell data. *Proceedings of the National Academy of Sciences*, 122(50), e2516046122.
- Van der Vaart, A. W. (2000). Asymptotic statistics. Cambridge University Press.
- Villani, C. (2021). Topics in optimal transportation (vol. 58). American Mathematical Soc.
- Willenborg, L. (2025). Using cell phones to compute dynamic population densities safely: A theoretical exploration. Discussion Paper. Statistics Netherlands (CBS), The Hague.
- Xie, Y., Wang, X., Wang, R., & Zha, H. (2020). A fast proximal point method for computing exact Wasserstein distance. In *Uncertainty in artificial intelligence* (pp. 433–453). PMLR.
- Yang, L., & Toh, K.-C. (2022). Bregman proximal point algorithm revisited: A new inexact version and its inertial variant. *SIAM Journal on Optimization*, 32(3), 1523–1554.
- Yang, S., & Zha, H. (2023). Estimating latent population flows from aggregated data via inverting multi-marginal optimal transport. In *Proceedings of the 2023 SIAM international conference on data mining (SDM)* (pp. 181–189). SIAM.