

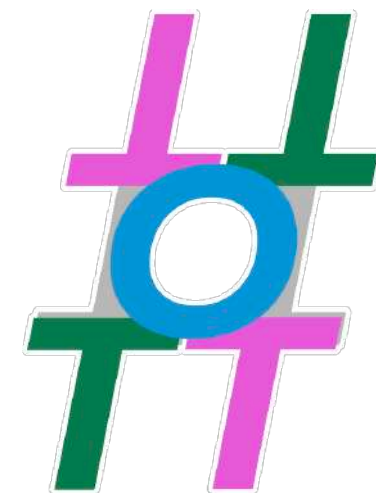
Optimal Transport

Marco Cuturi

MLSS NTU 2021



Google AI
Brain Team



***Optimal
Transport
Tools***

First, a big thanks !

To the MLSS NTU 2021 organisers

To all my collaborators on these topics, including here a recent subset



M. Seigal



L. Petrini



L. Papaxanthos



O. Teboul



JP Vert



J. Niles-Weed



C. Bunne



T. Lin



A. Genevay



M. Scetbon



A. Gramfort



J. Josse



H. Janati



F.P. Paty



G. Peyré



B. Muzellec



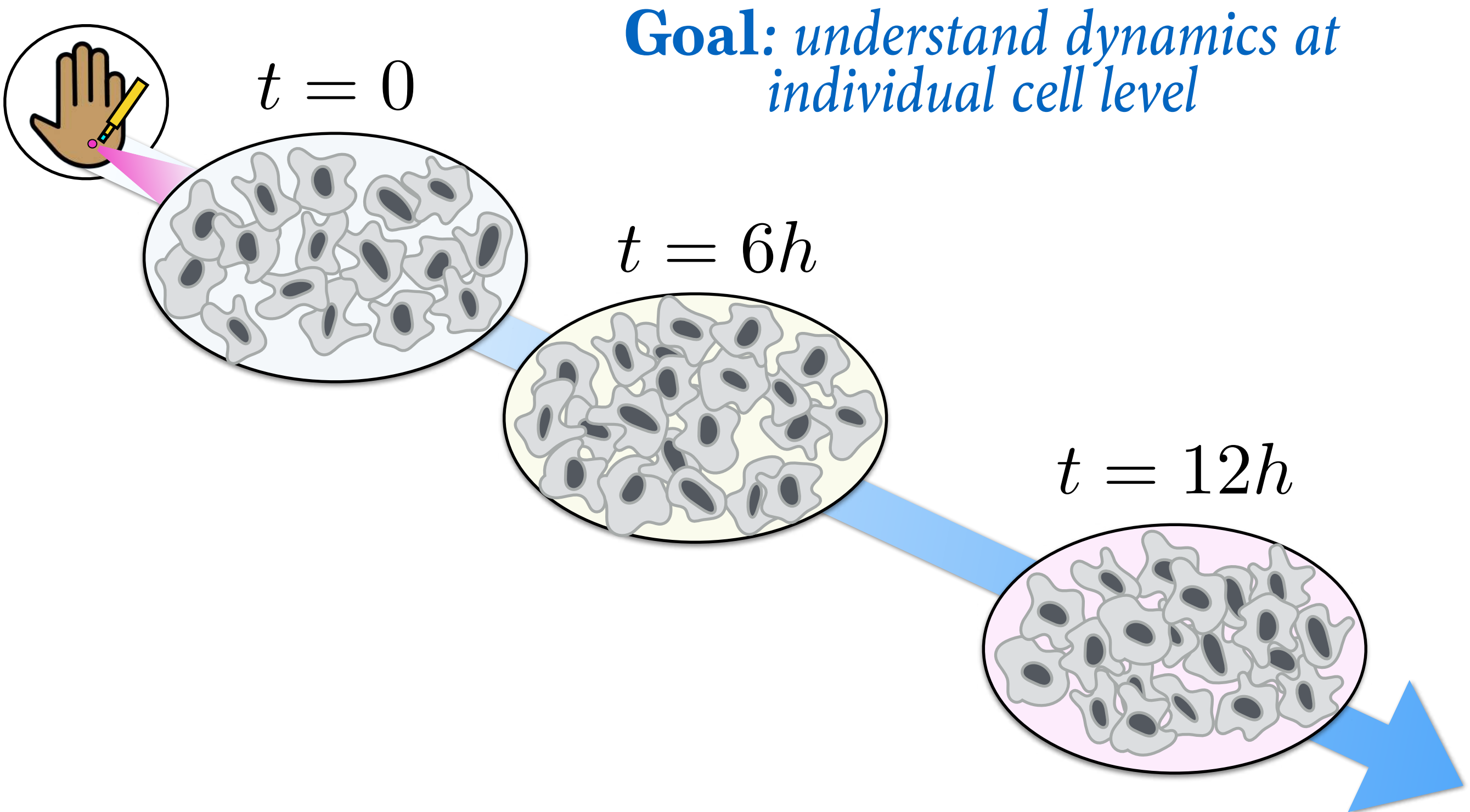
C. Boyer



Google AI
Brain Team

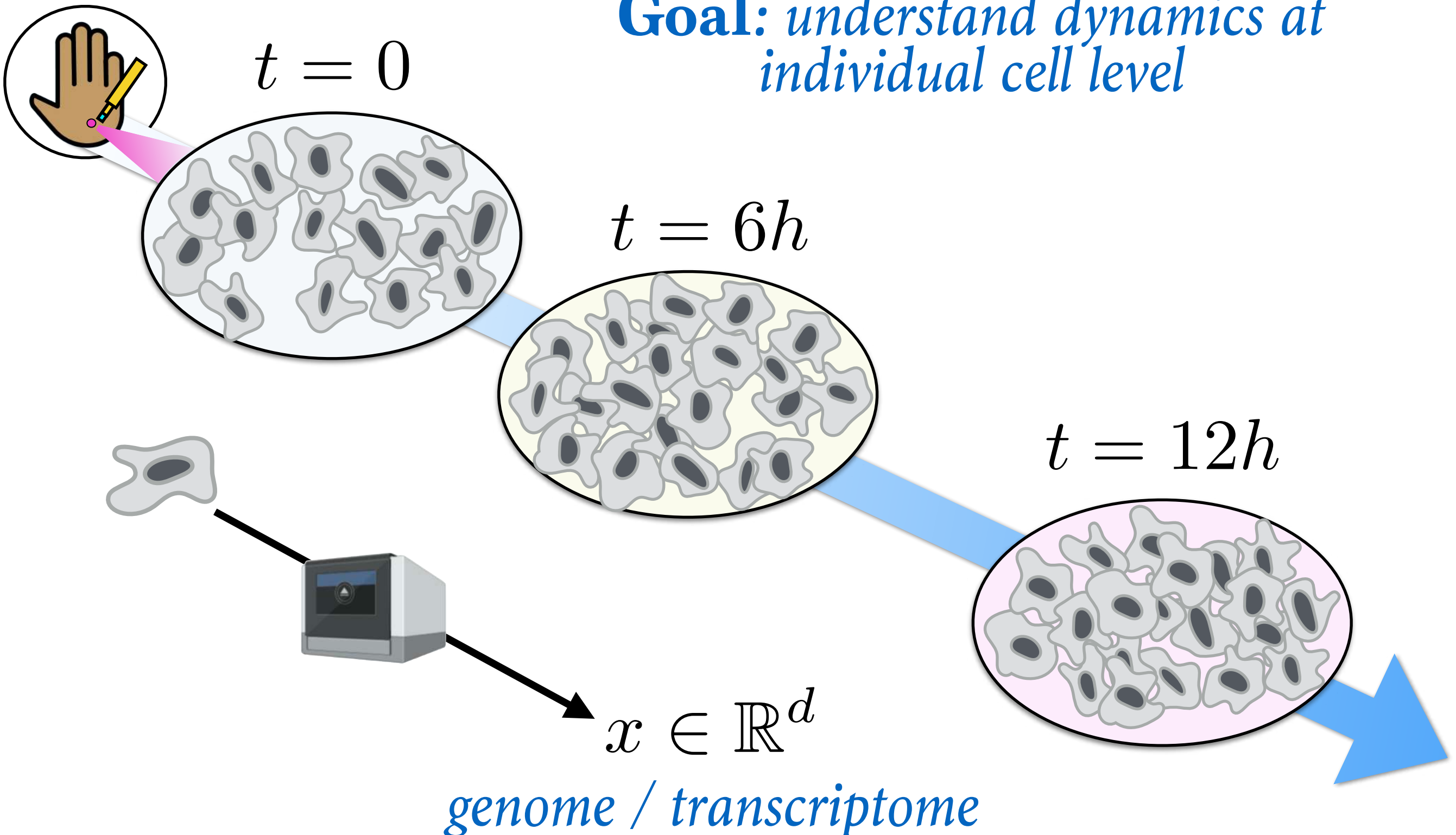


An Example from Single-Cell Genomics

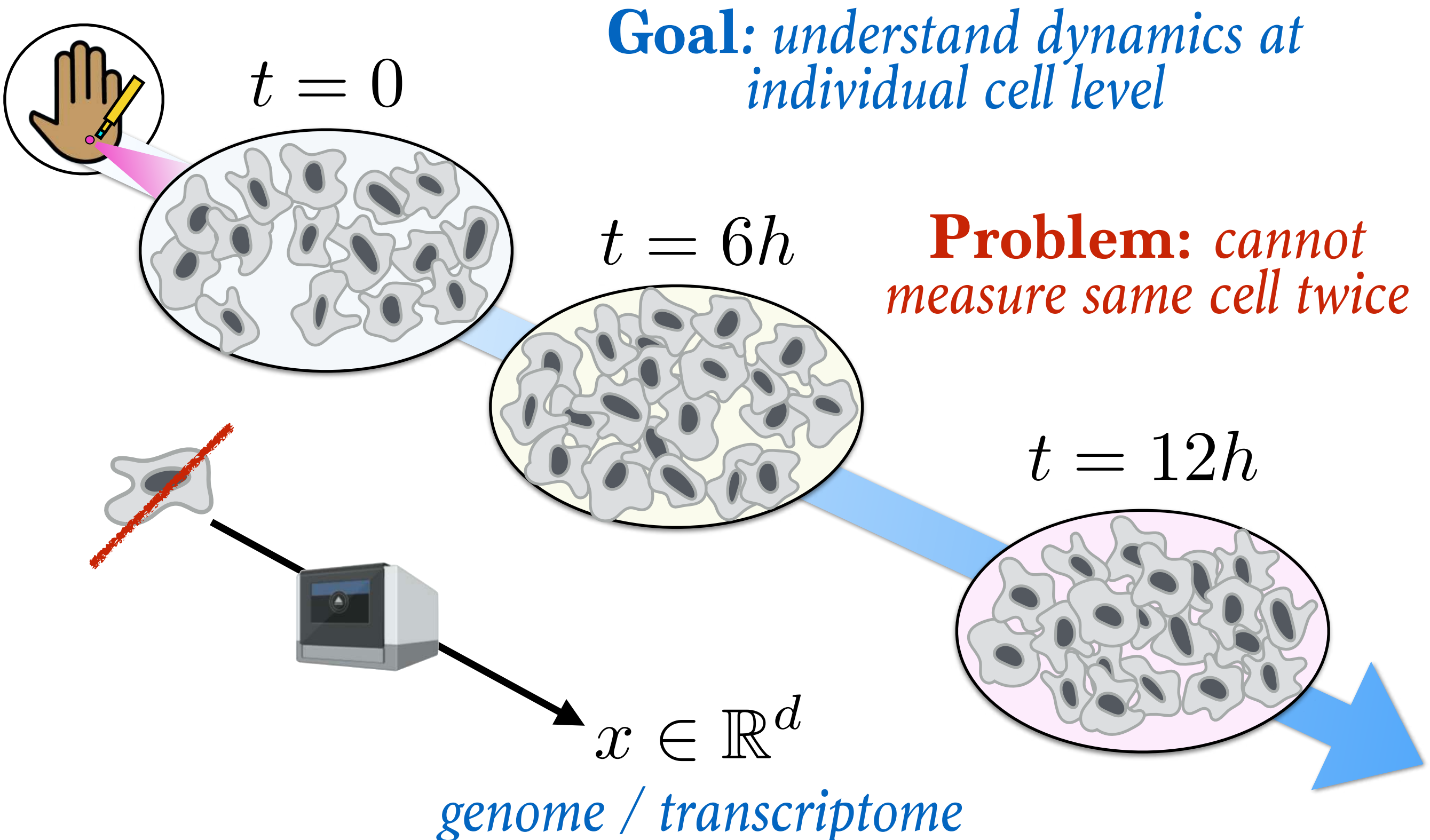


An Example from Single-Cell Genomics

Goal: *understand dynamics at individual cell level*

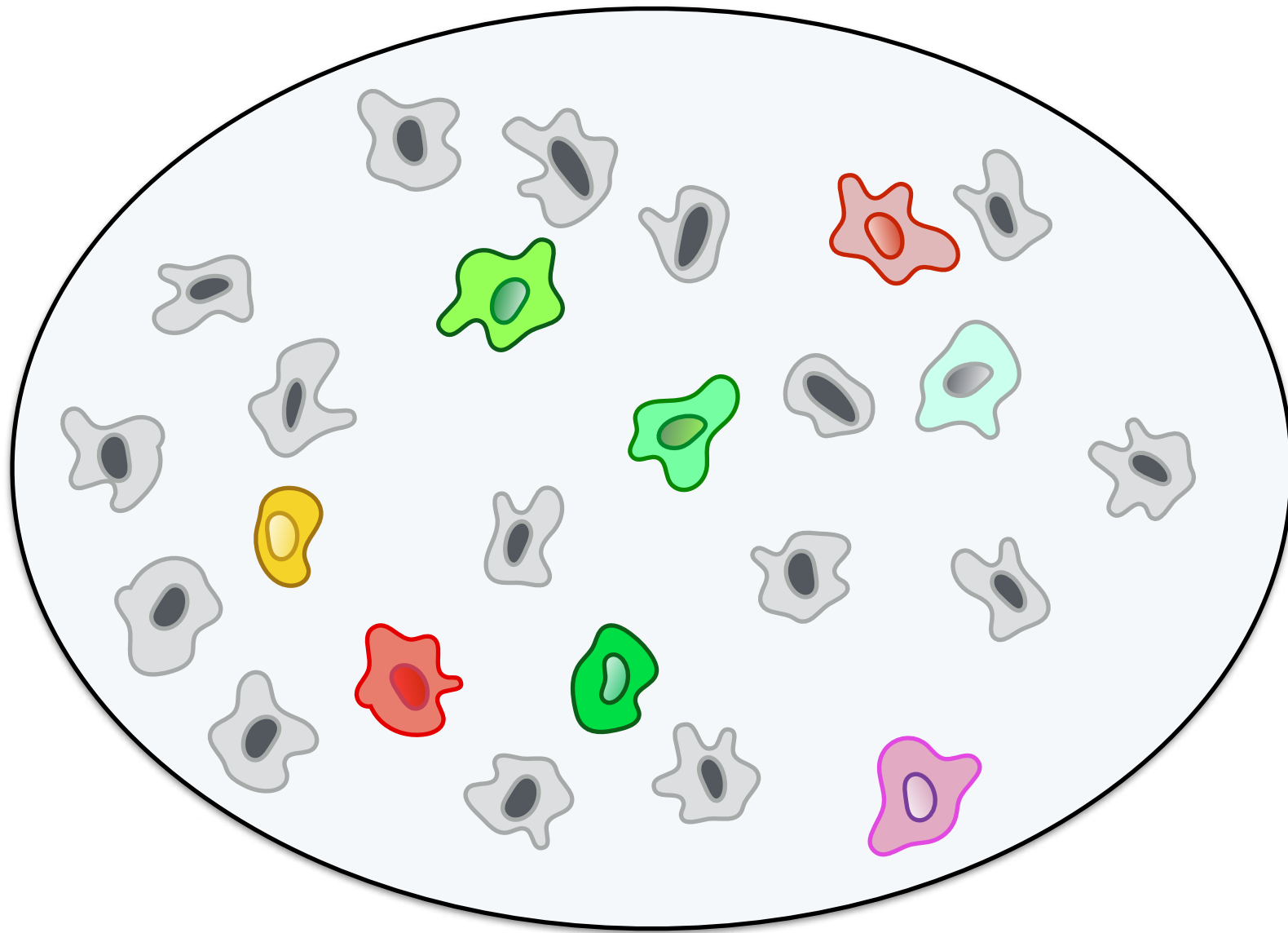


An Example from Single-Cell Genomics



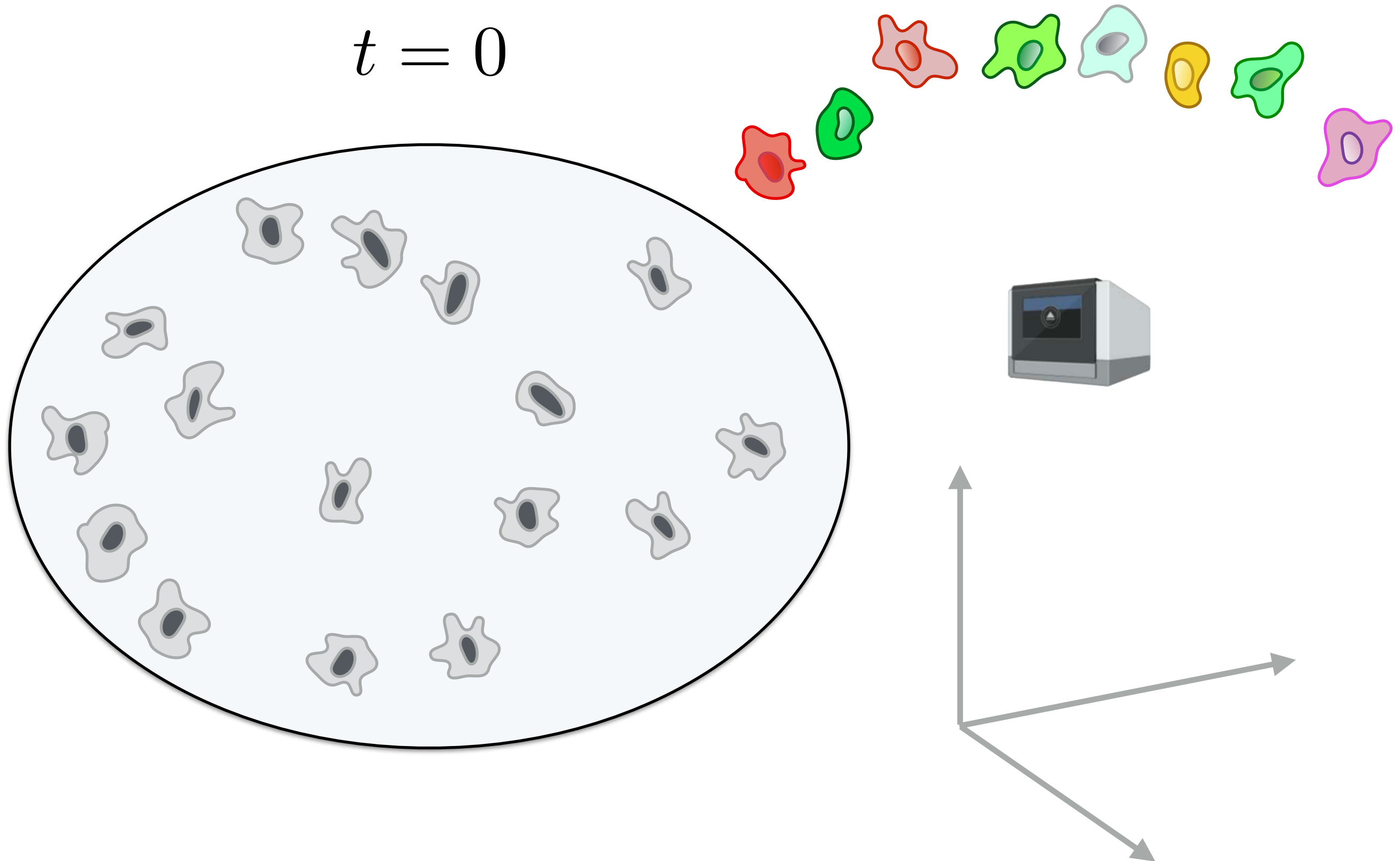
An Example from Single-Cell Genomics

$t = 0$



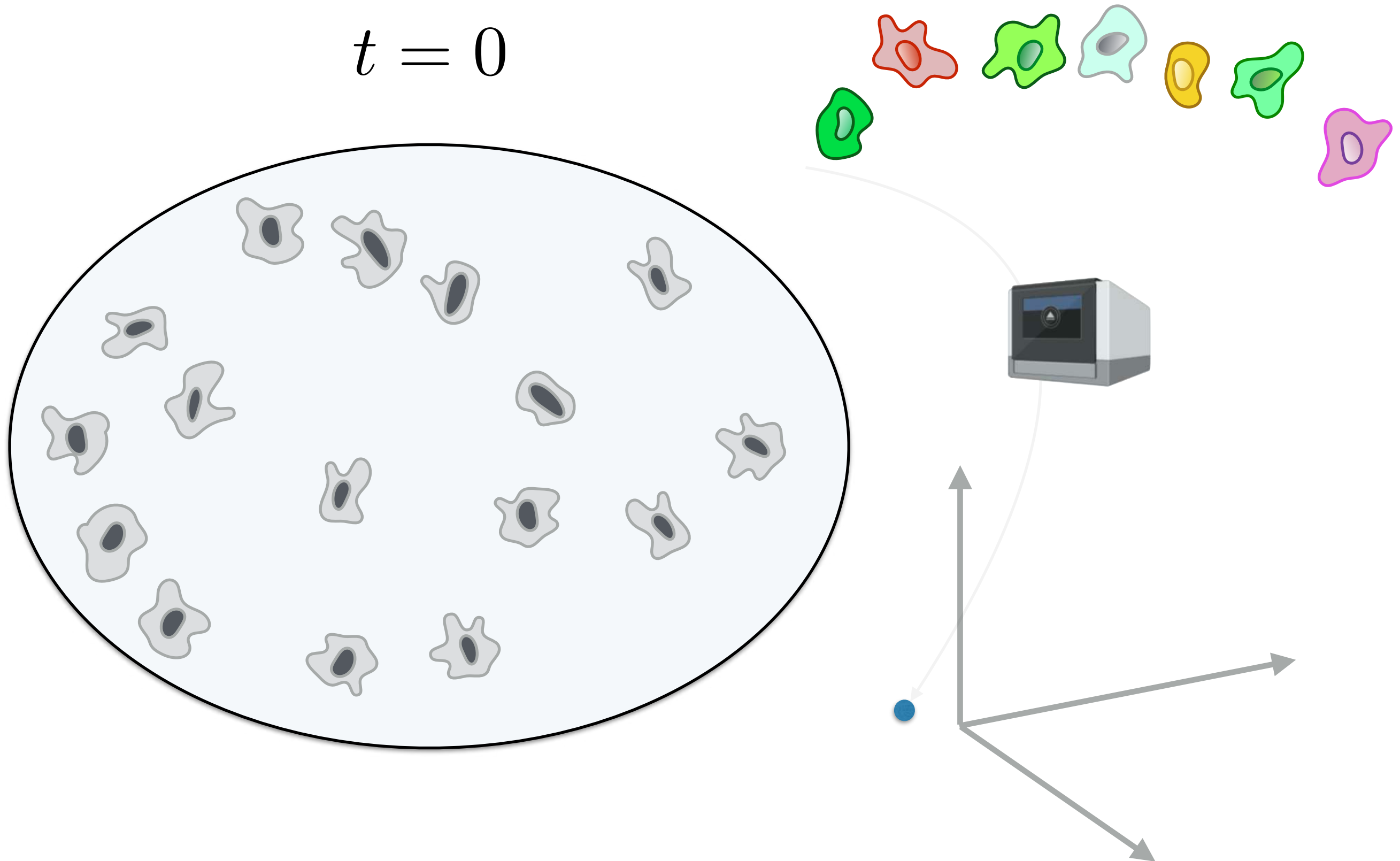
An Example from Single-Cell Genomics

$t = 0$



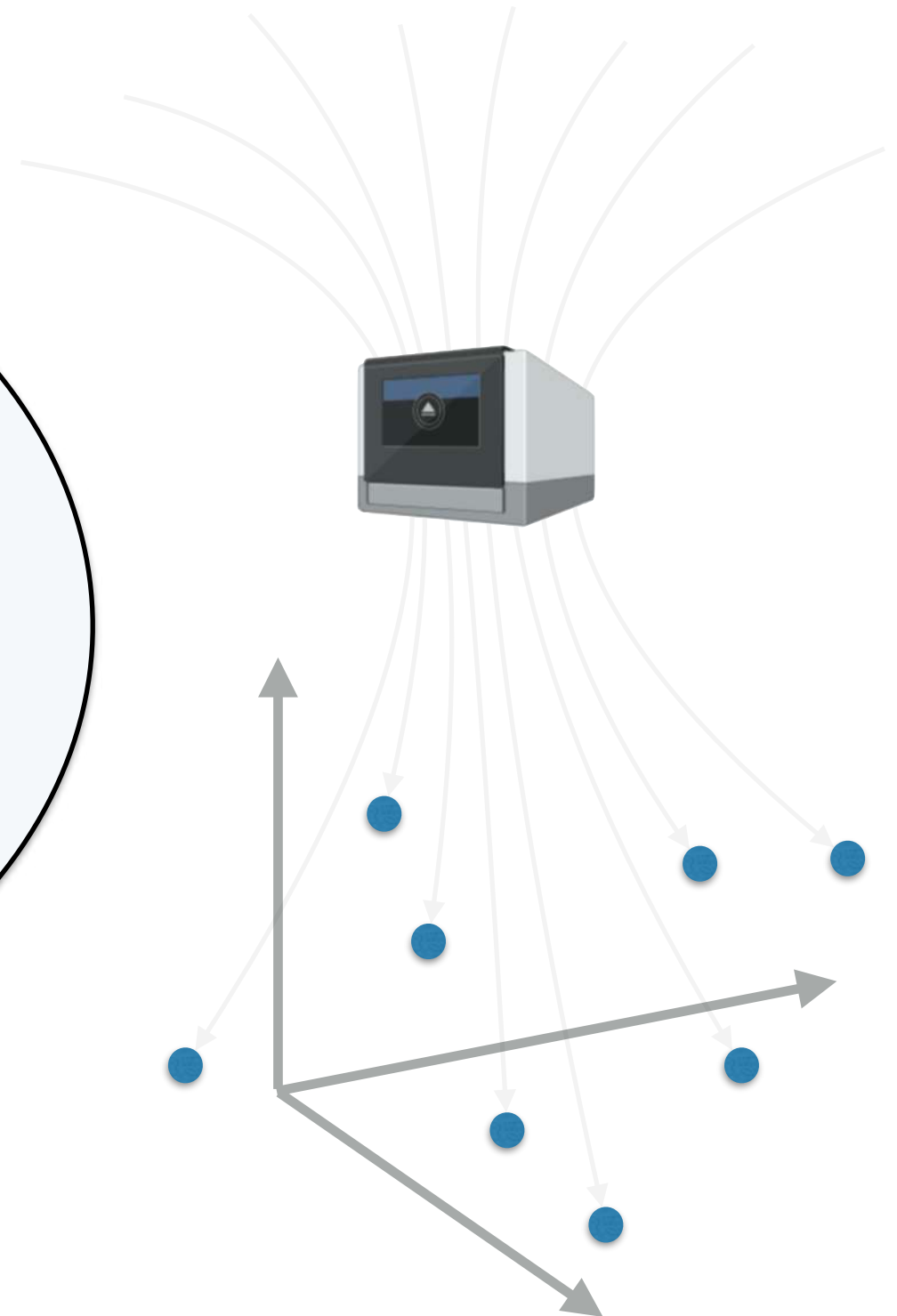
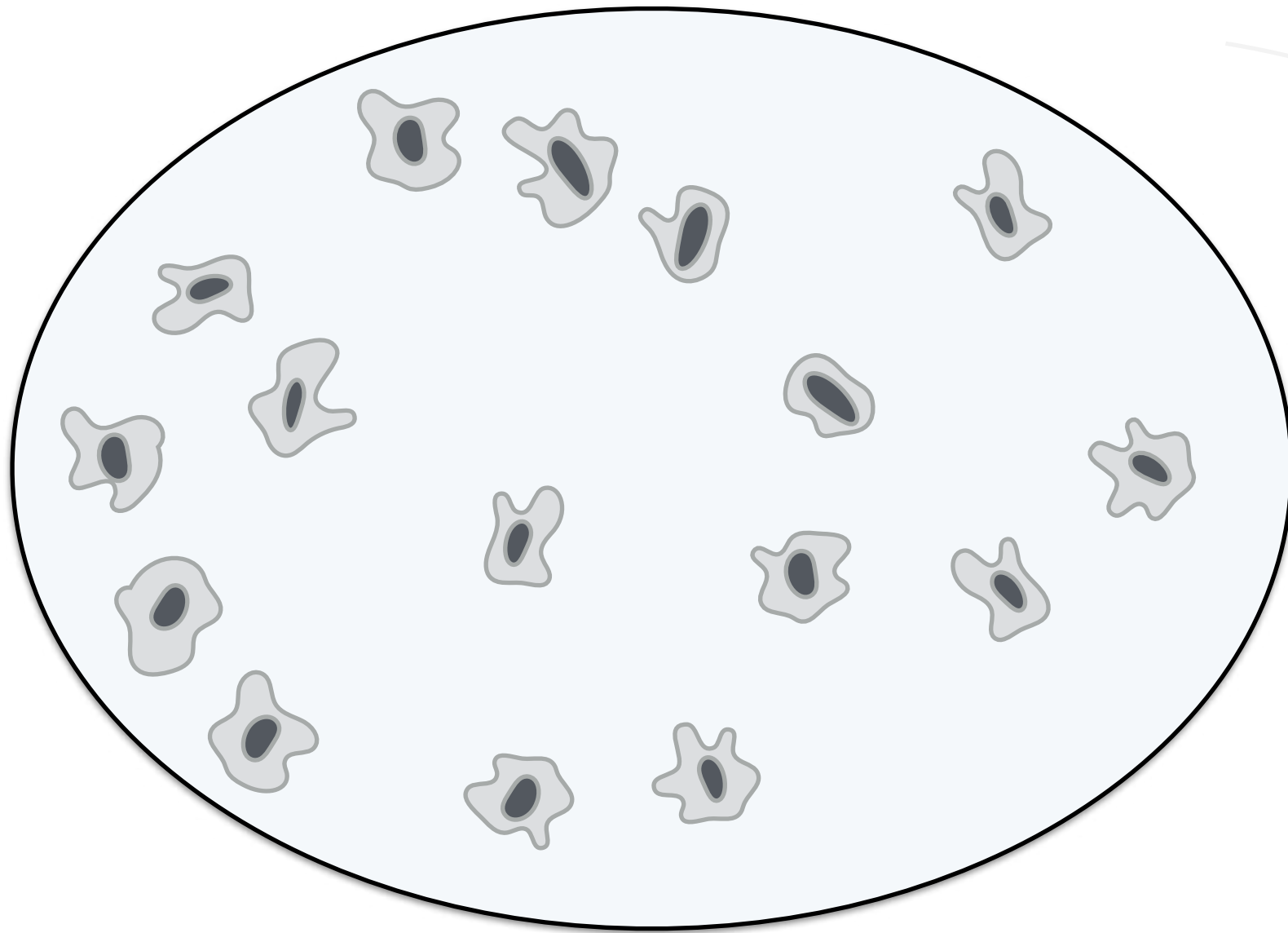
An Example from Single-Cell Genomics

$t = 0$



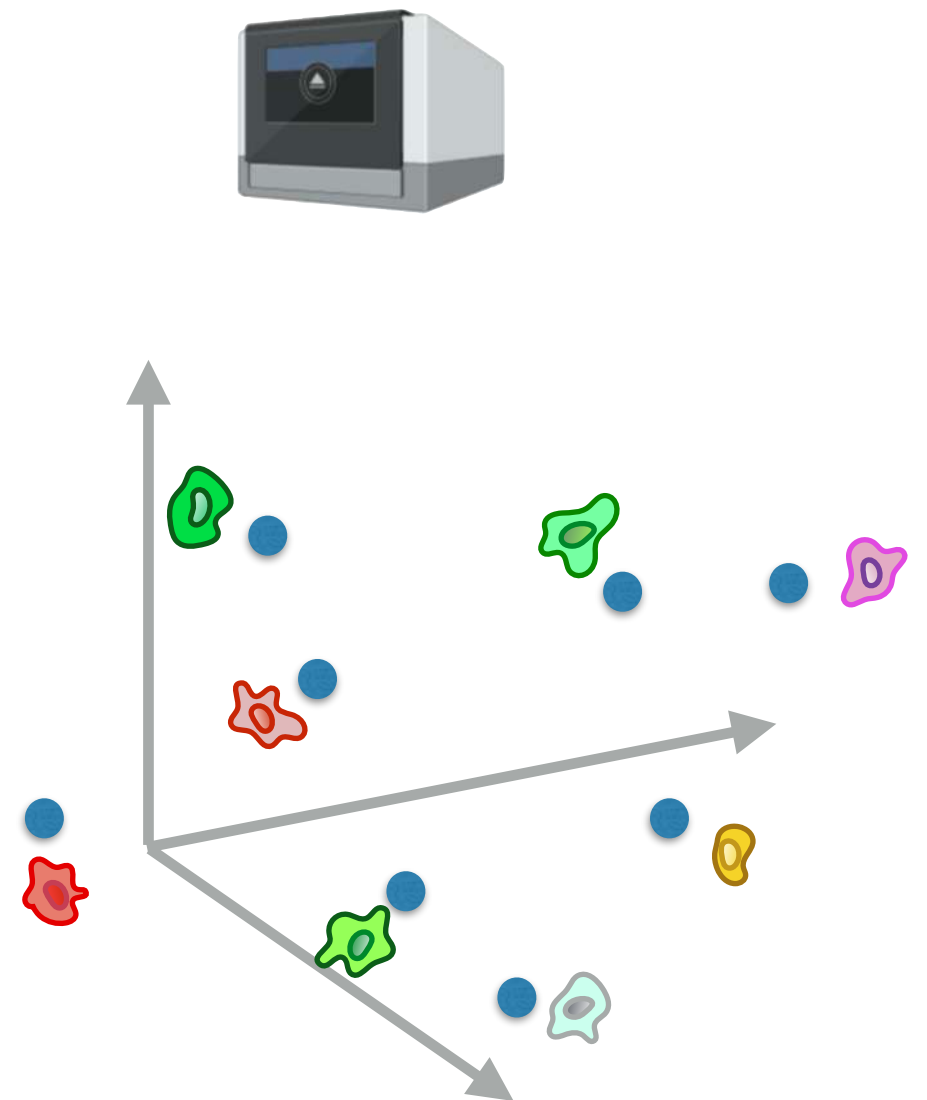
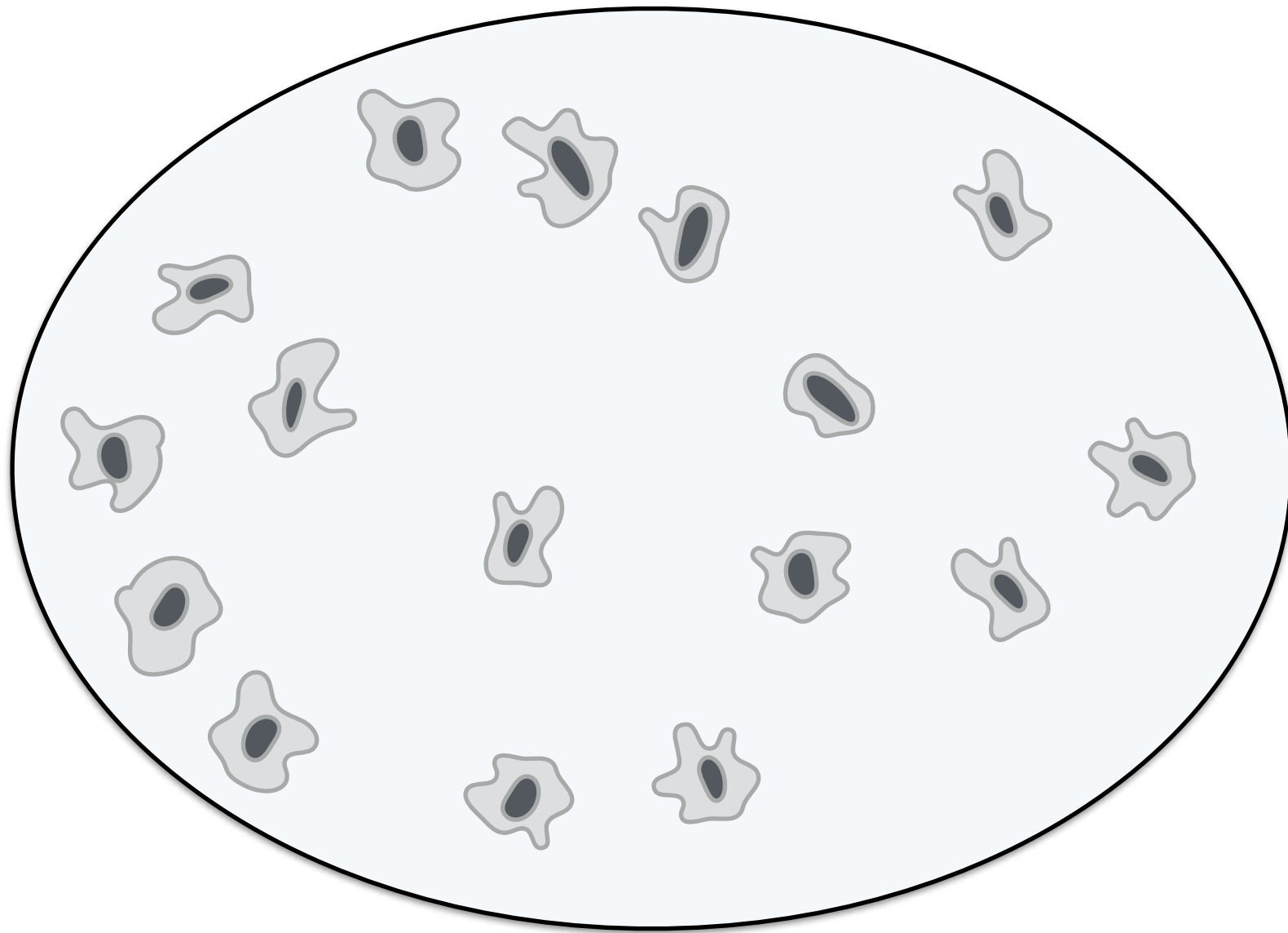
An Example from Single-Cell Genomics

$t = 0$



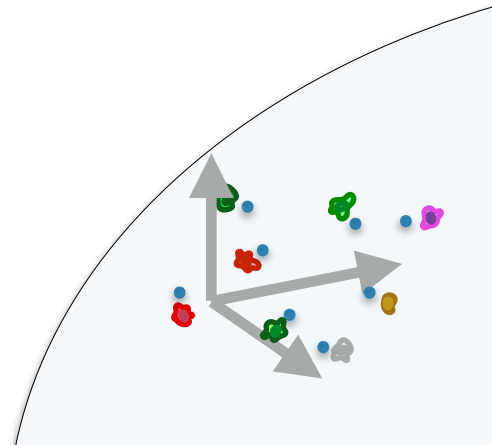
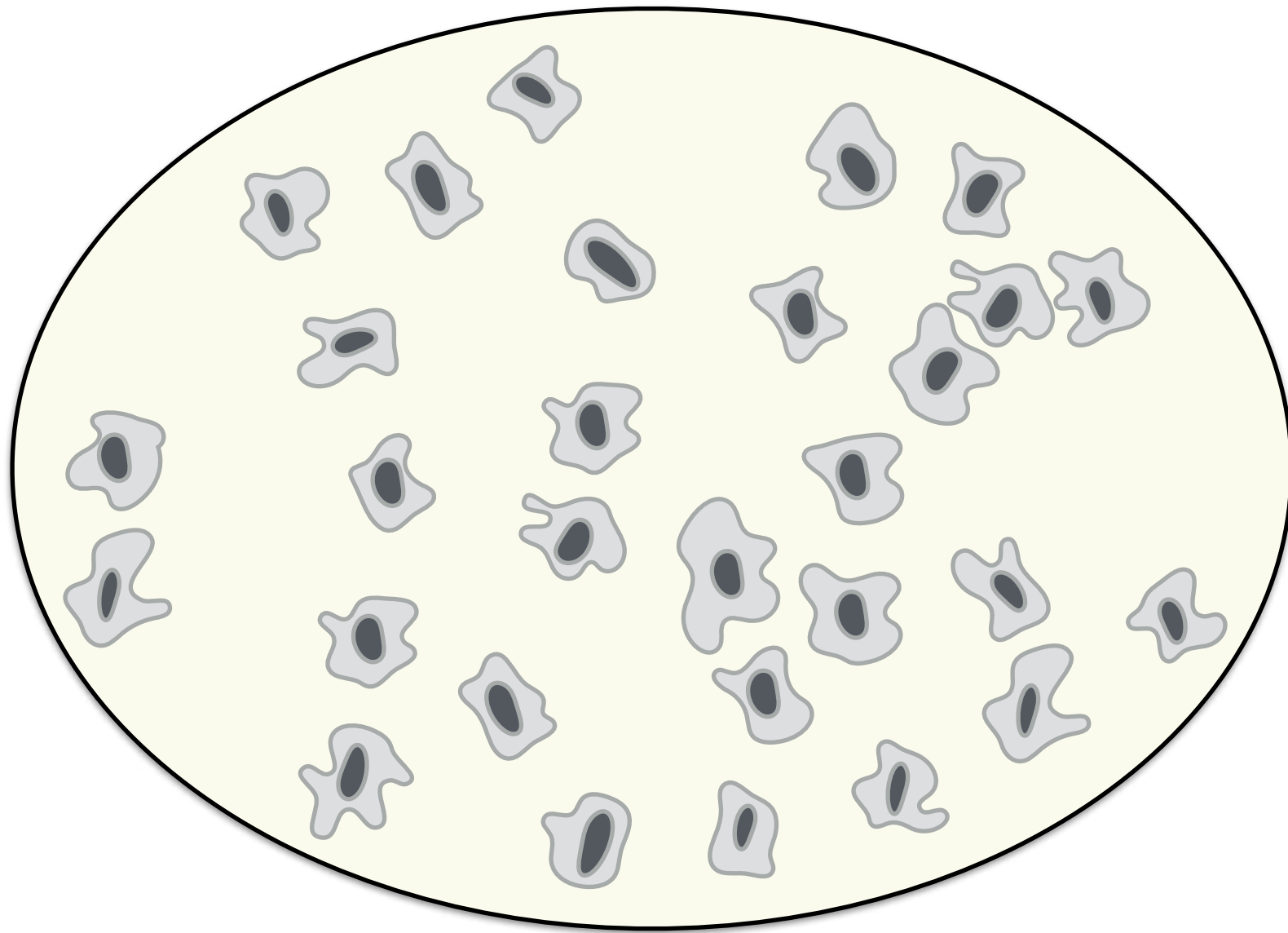
An Example from Single-Cell Genomics

$t = 0$



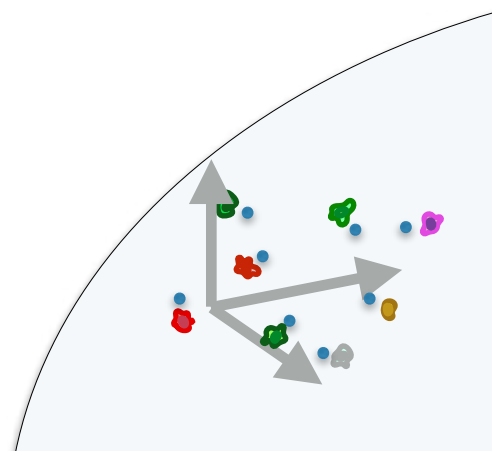
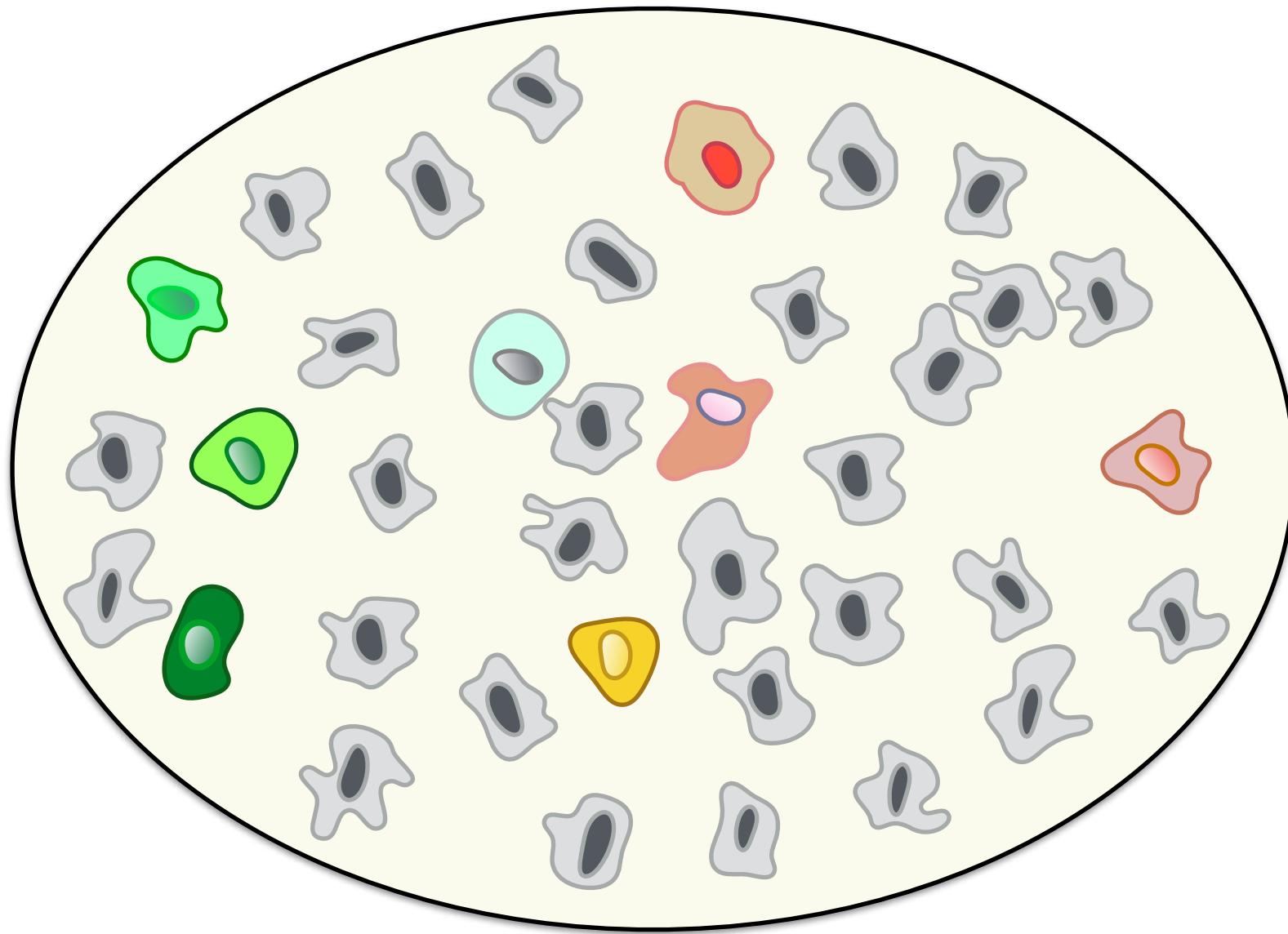
An Example from Single-Cell Genomics

$t = 6h$



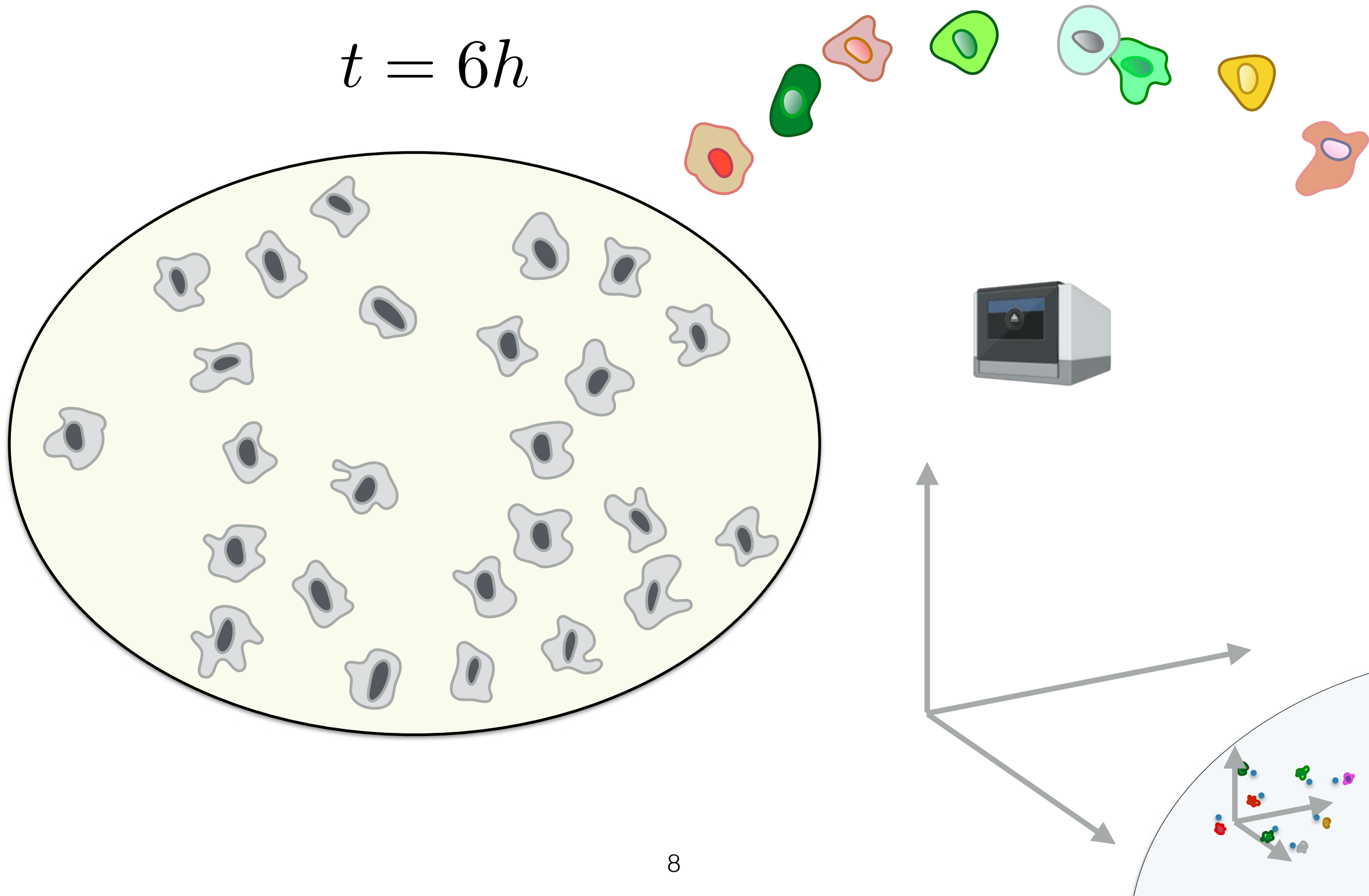
An Example from Single-Cell Genomics

$t = 6h$



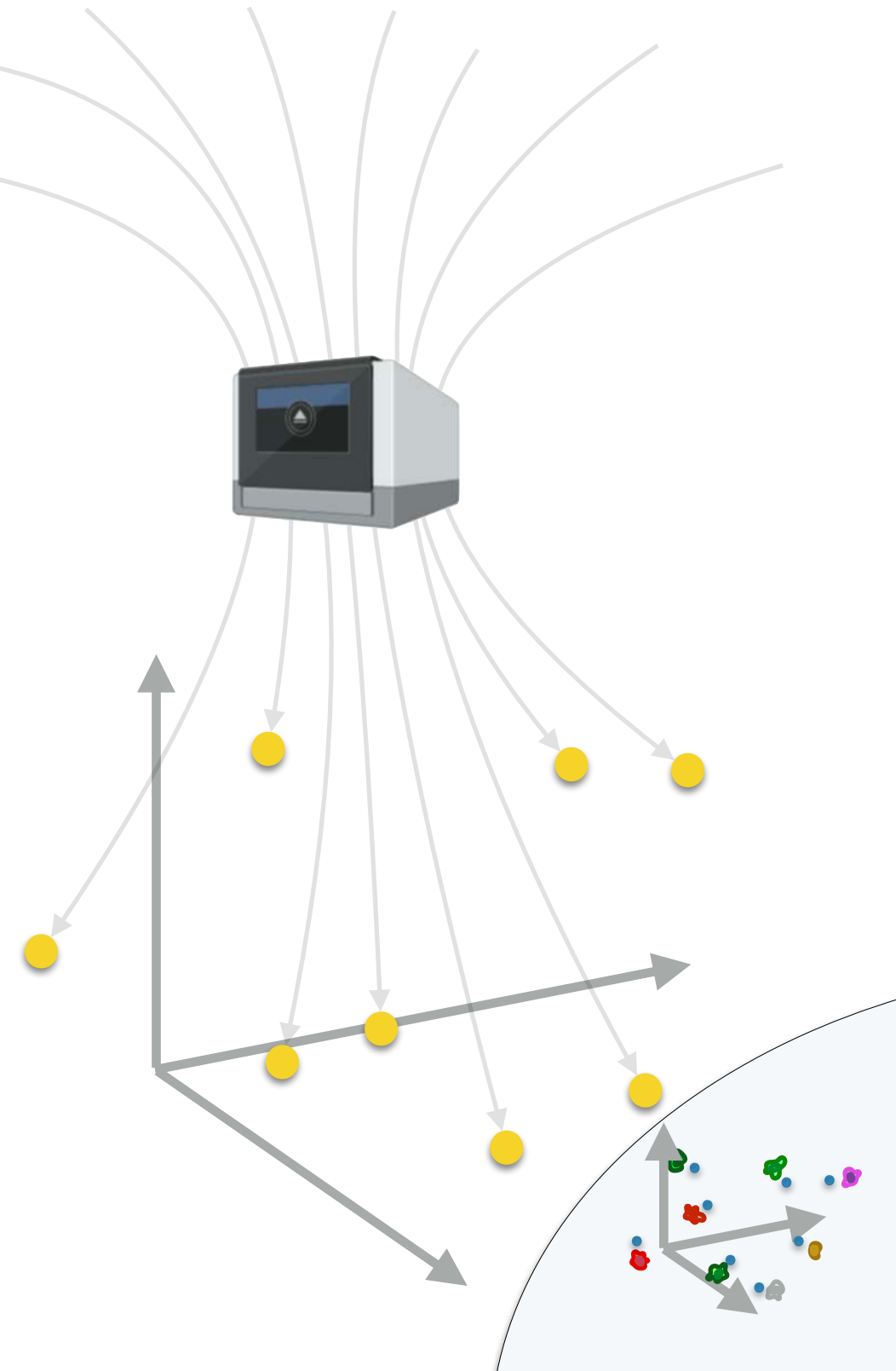
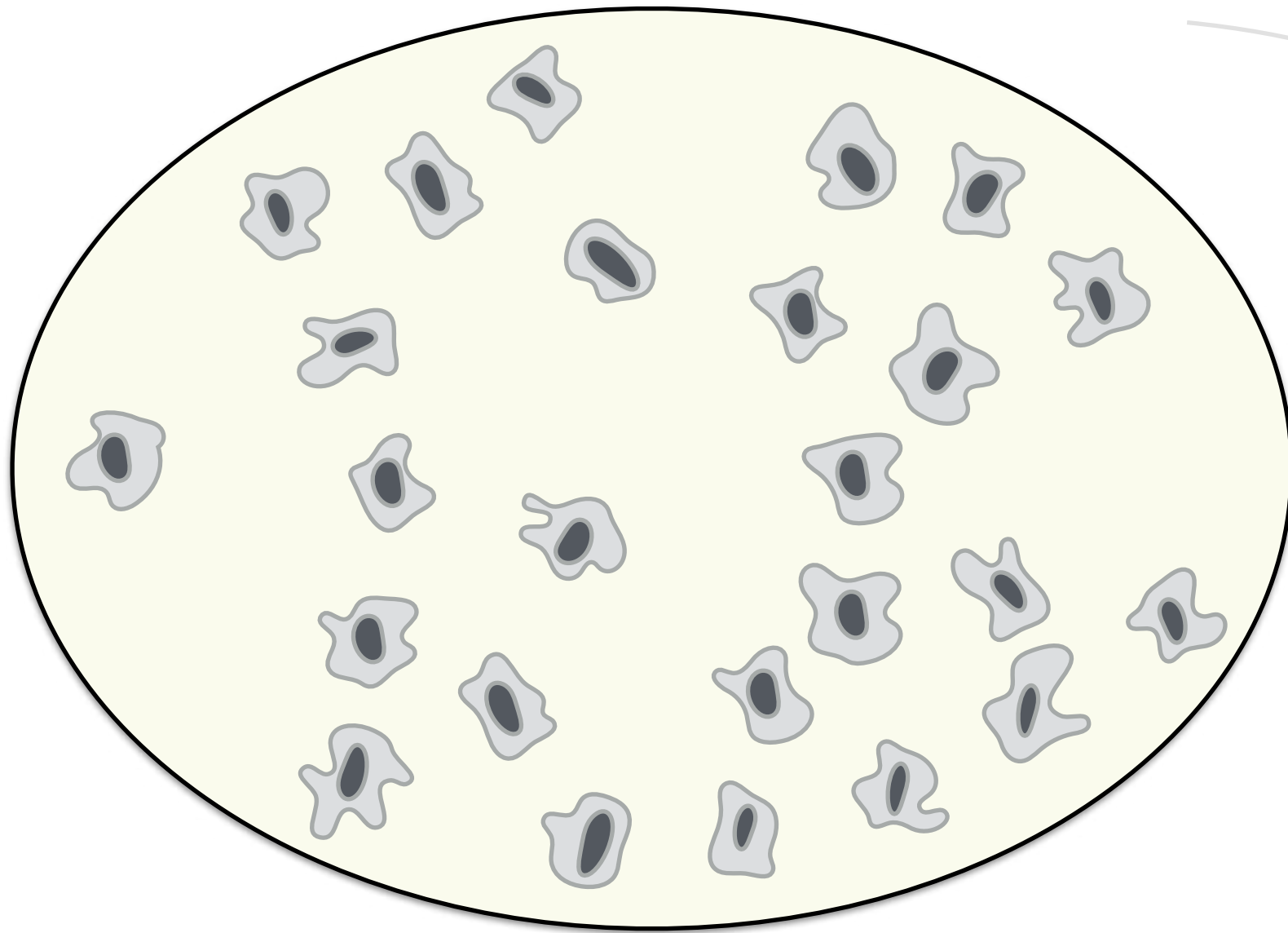
An Example from Single-Cell Genomics

$t = 6h$



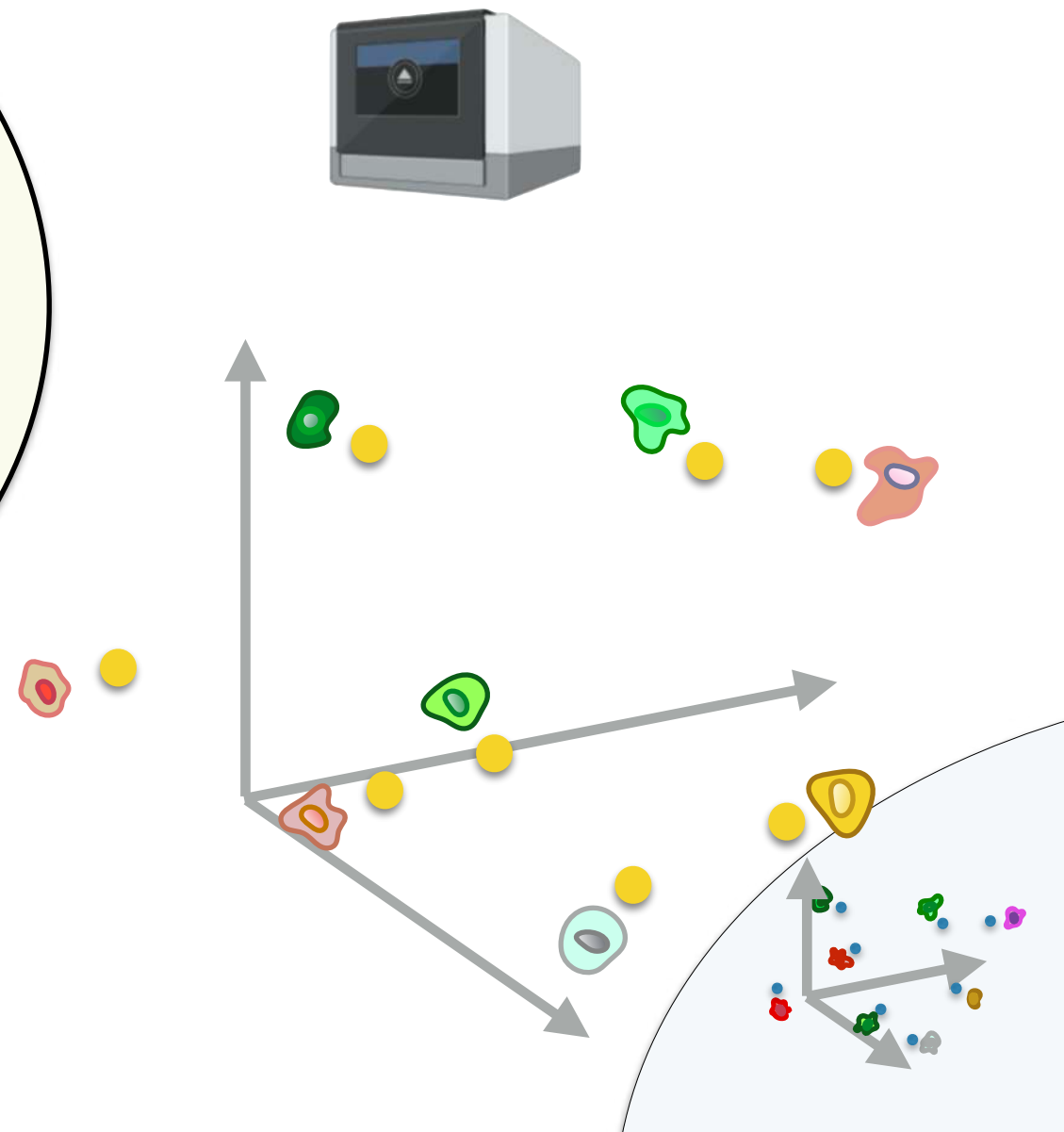
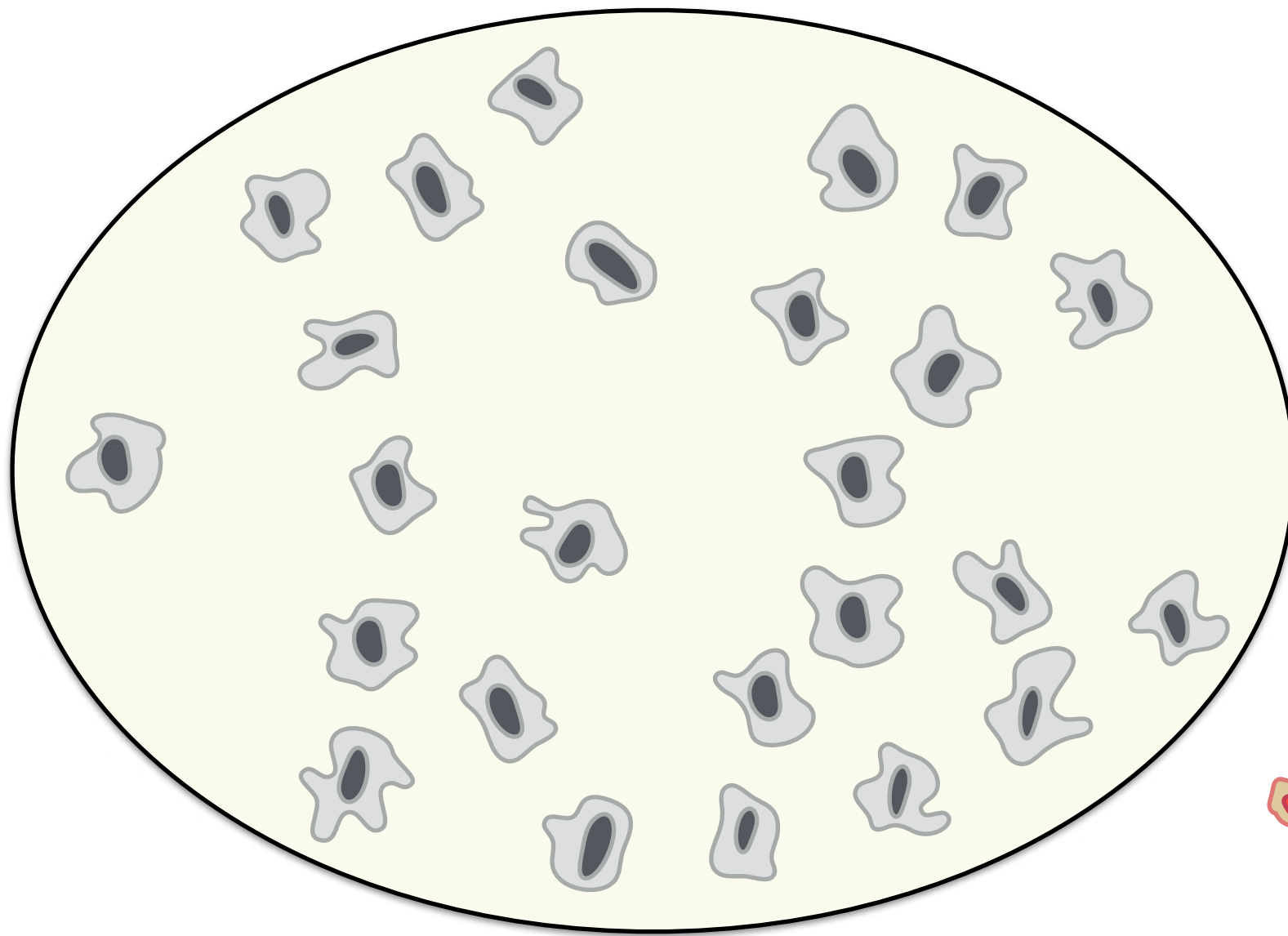
An Example from Single-Cell Genomics

$t = 6h$

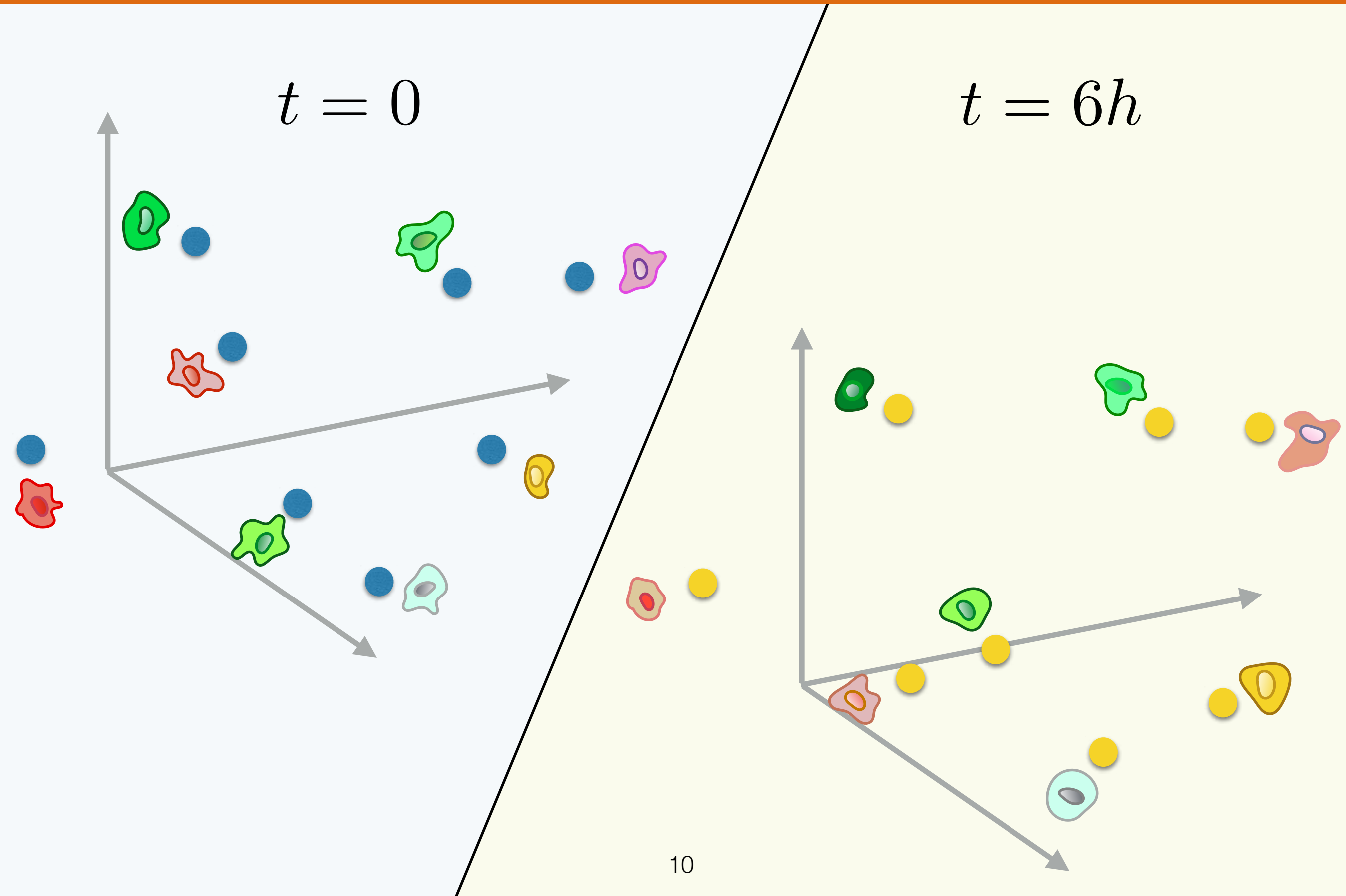


An Example from Single-Cell Genomics

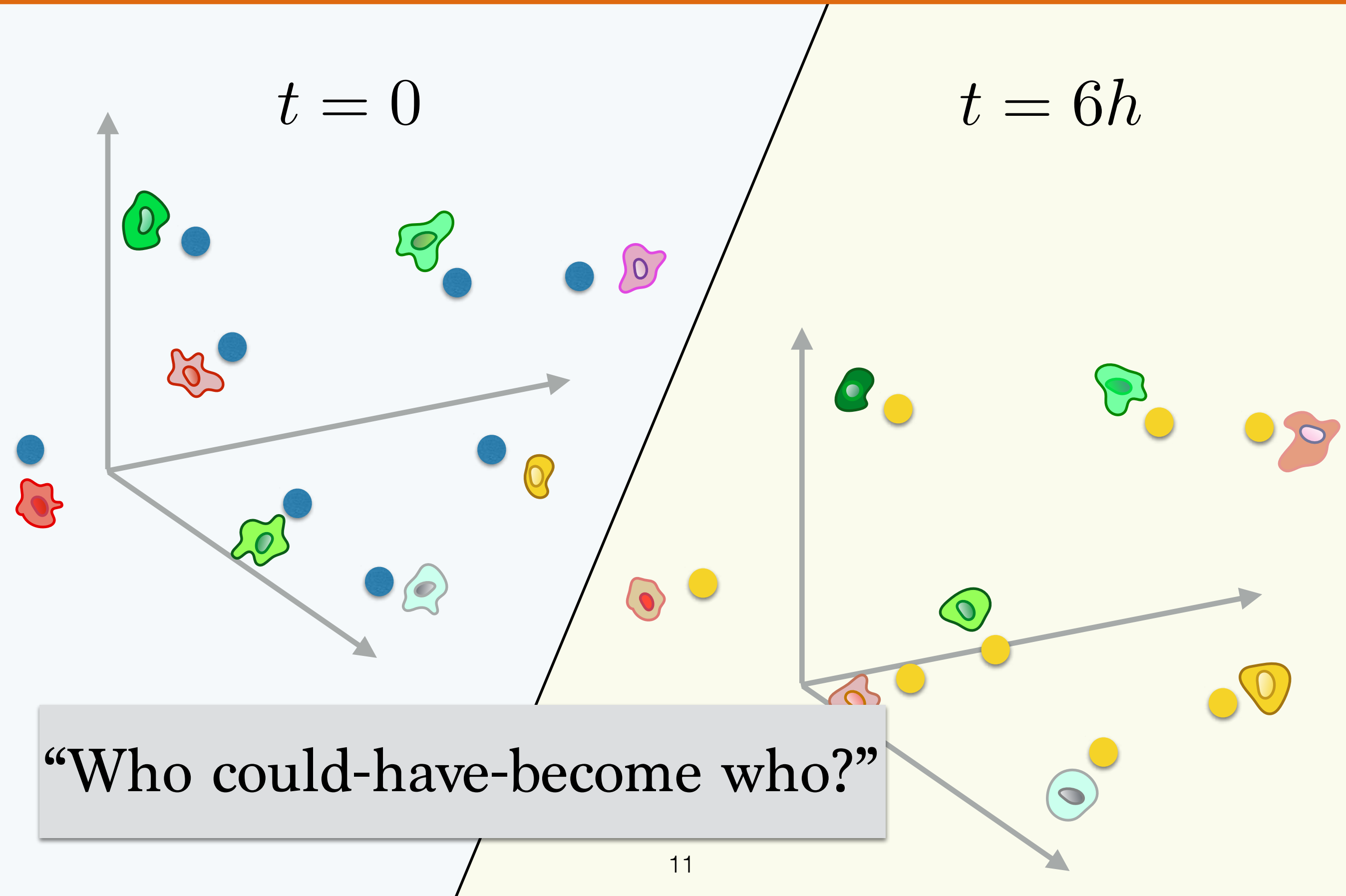
$t = 6h$



Comparing Two Populations

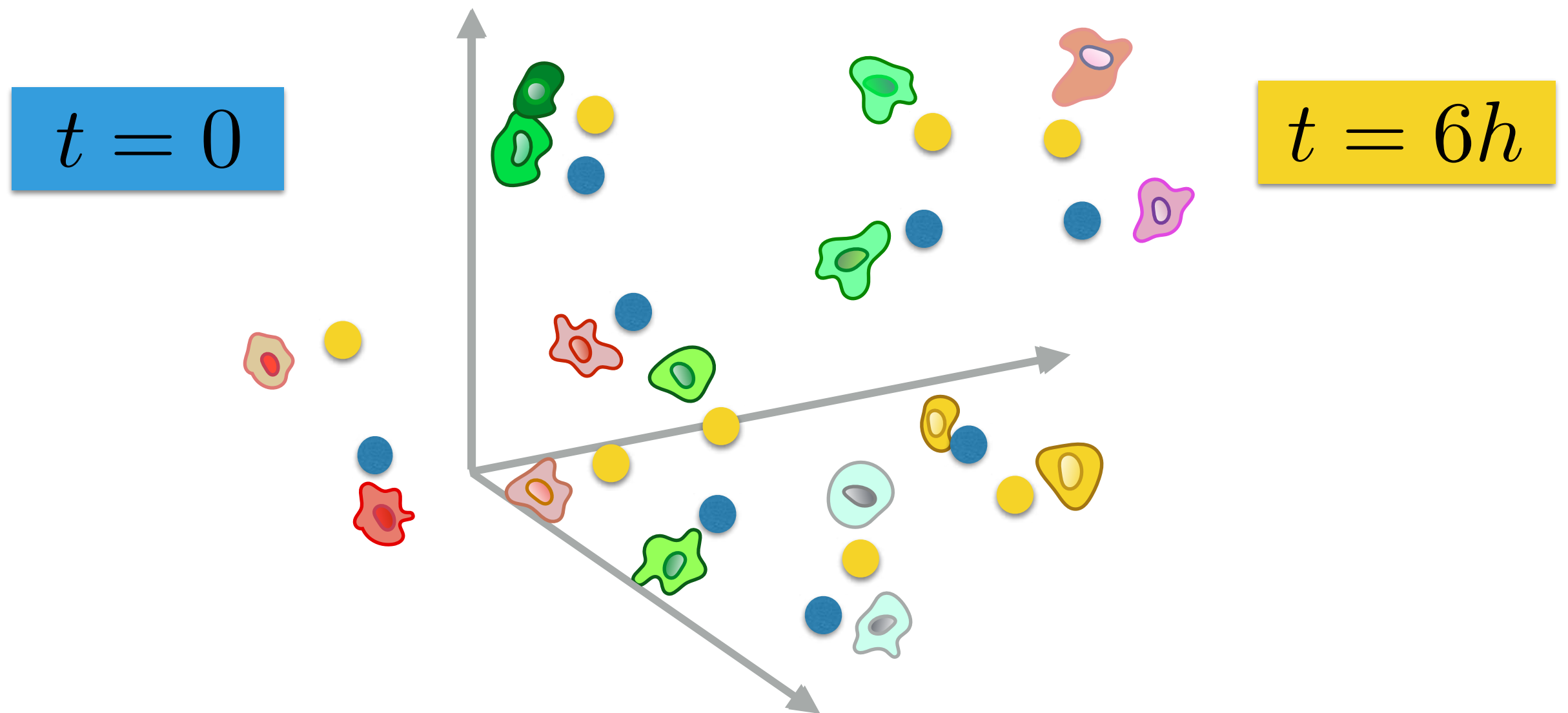


Comparing Two Populations



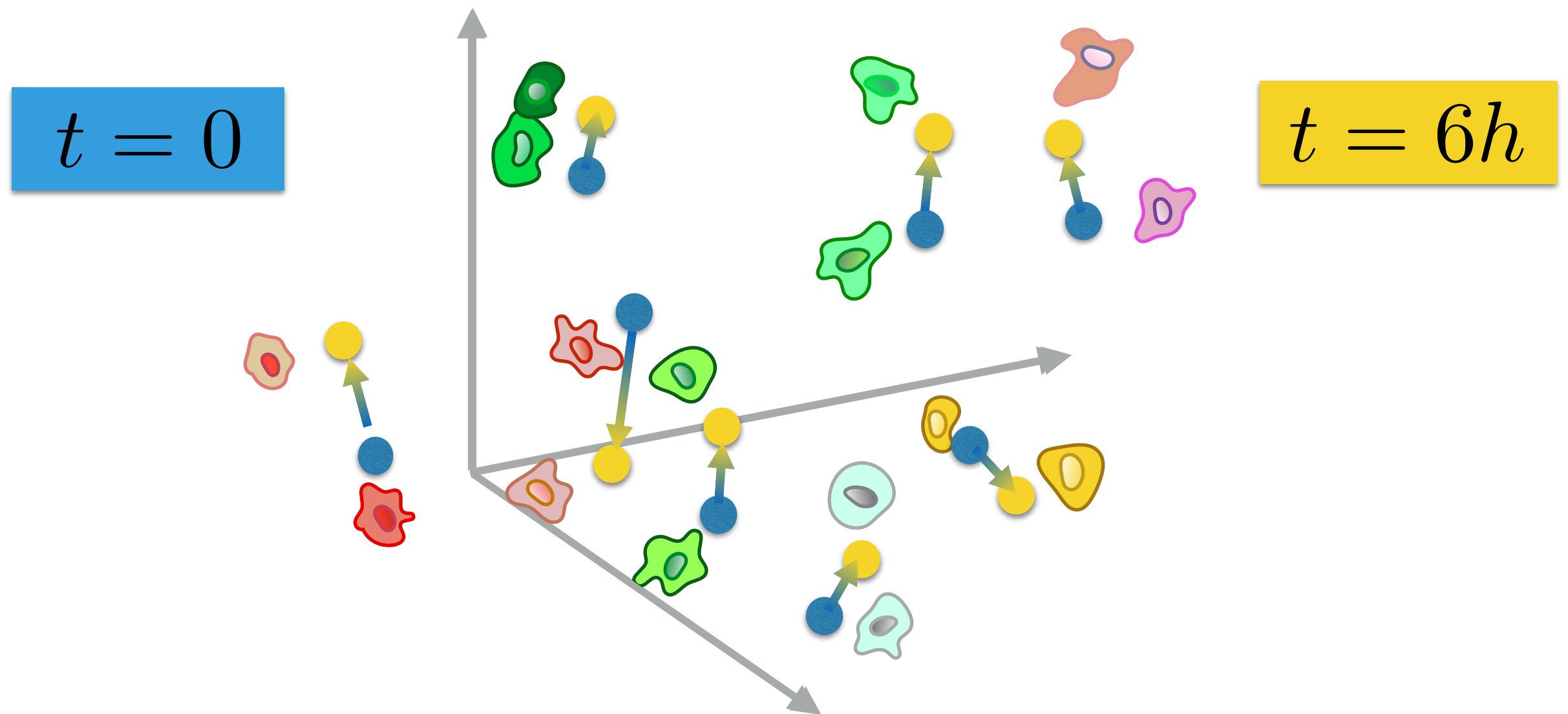
Matching Two Populations

*to answer that question, use an **optimal matching****



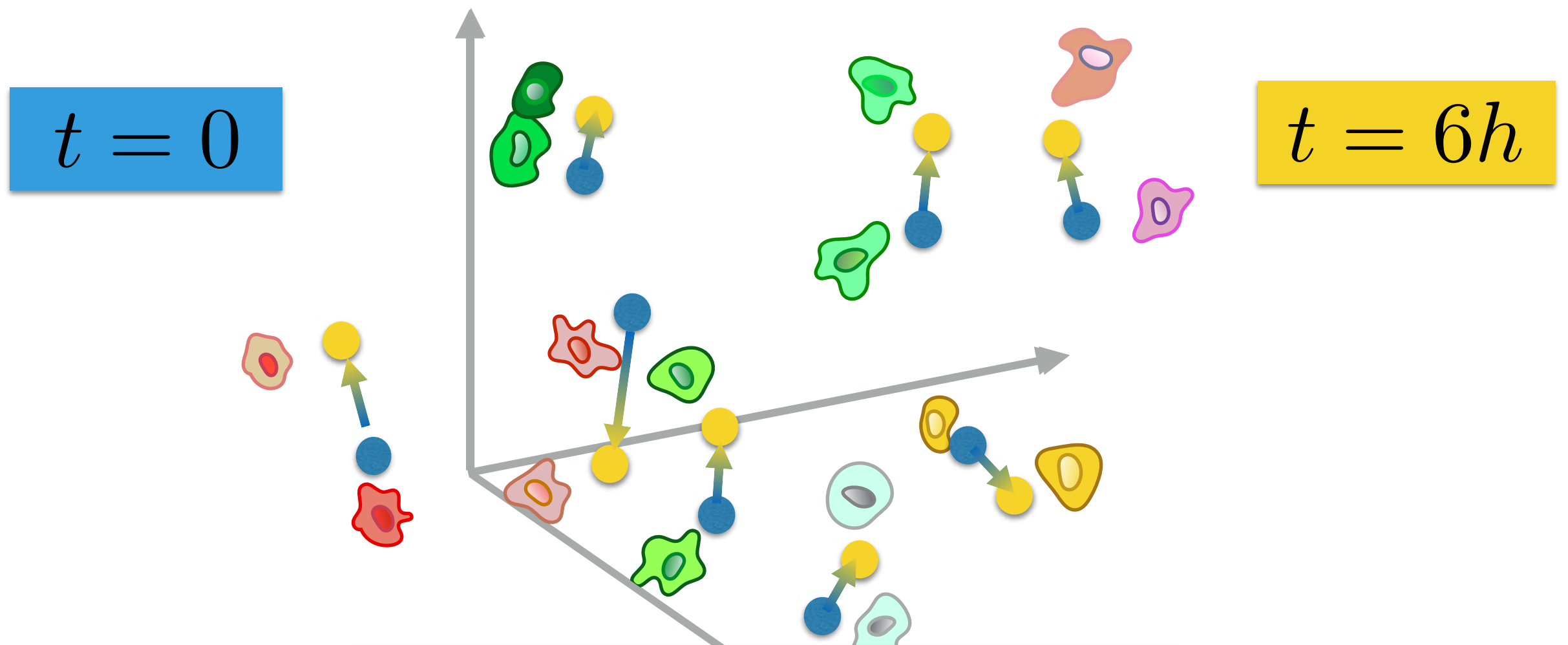
Matching Two Populations

*to answer that question, use an **optimal matching****



Matching Two Populations

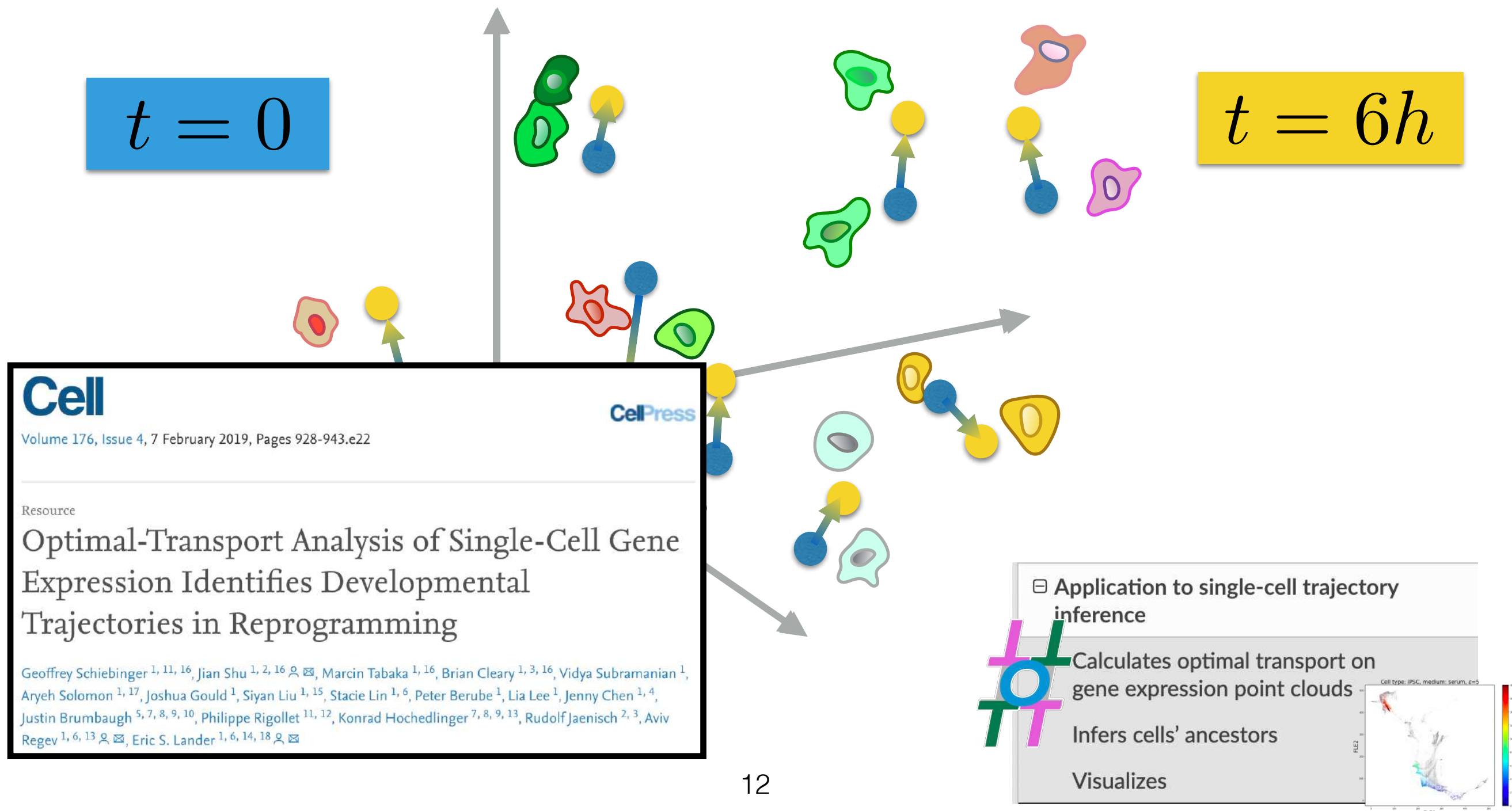
*to answer that question, use an **optimal matching****



$$* \quad \sigma^* = \operatorname{argmin}_{\sigma \in \mathcal{S}(n)} \sum_{i=1}^n d(x_i, y_{\sigma_i})$$

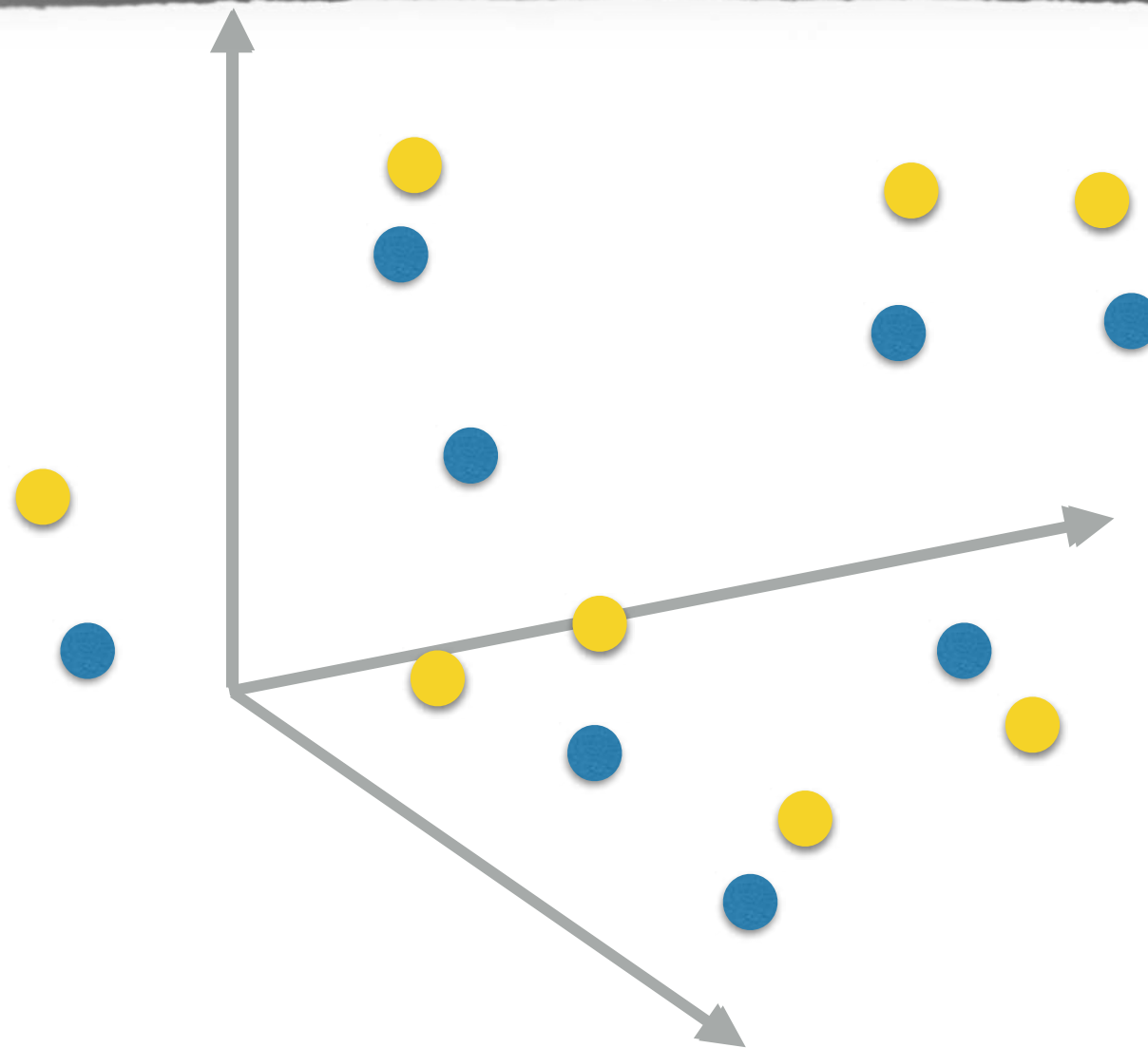
Matching Two Populations

*to answer that question, use an optimal matching**



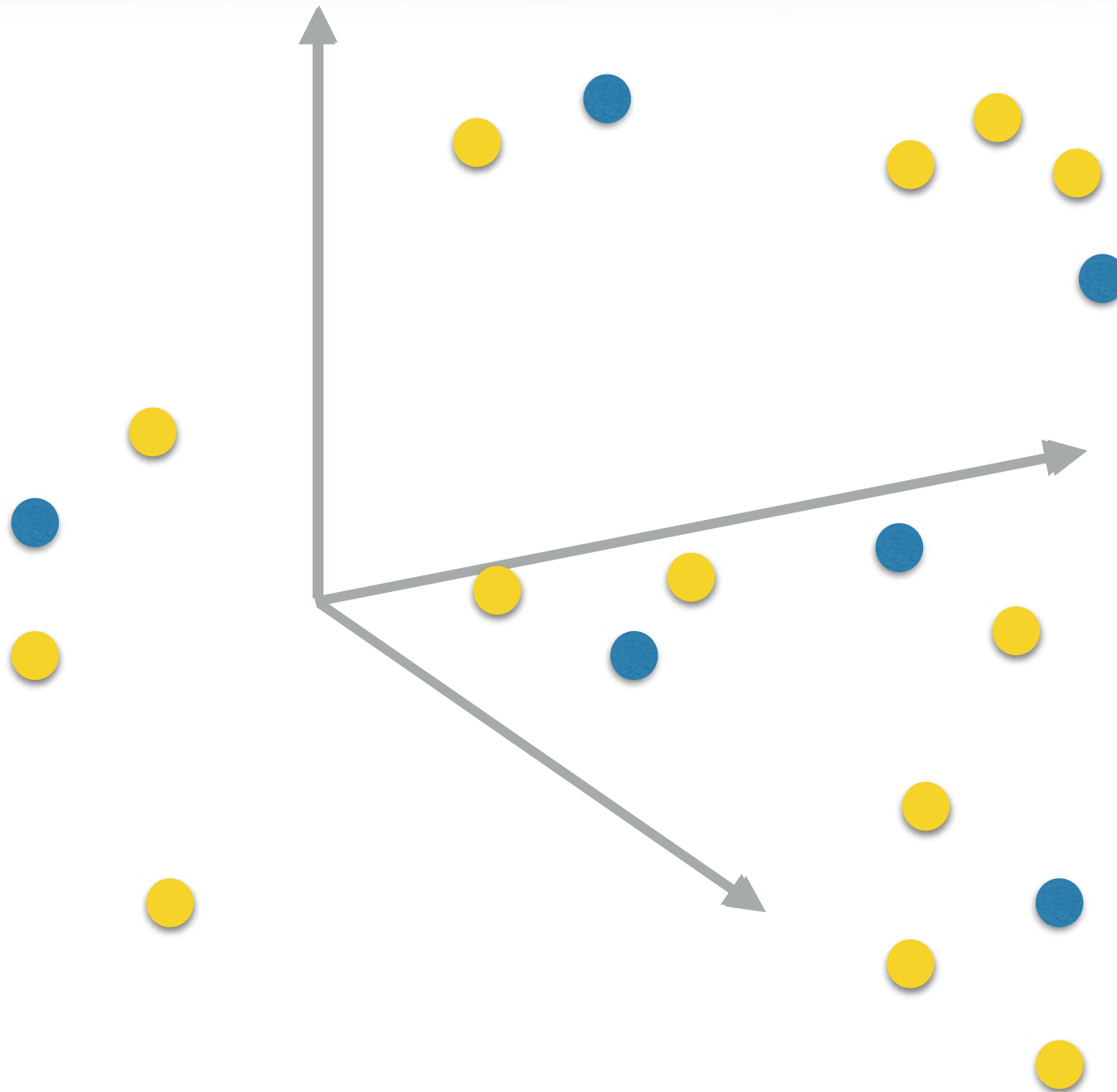
From there, many generalisations

This story describes an ideal experimental setting, notably because this requires the same number of cells.



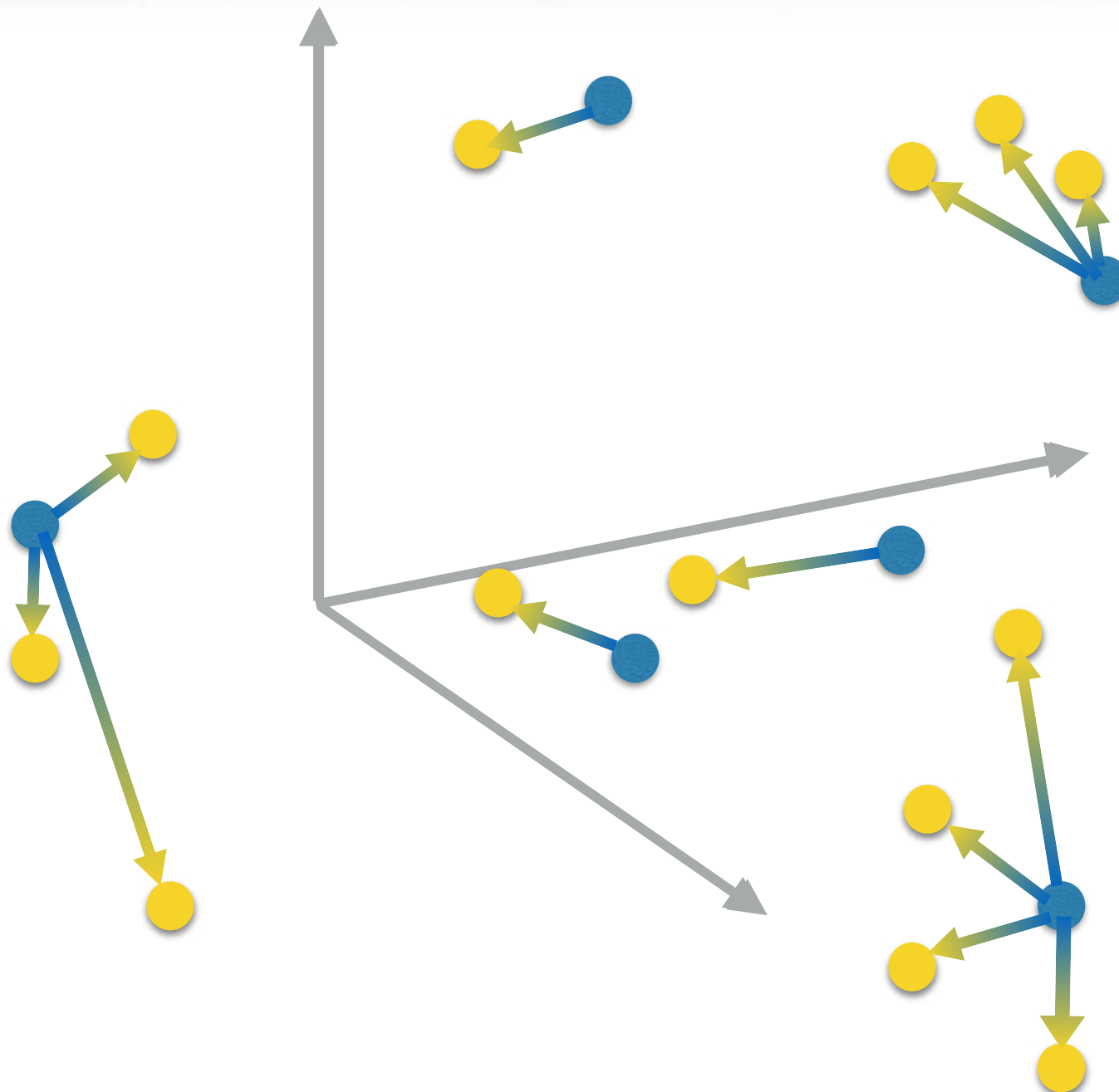
Challenges Down the Road

More likely to happen: *variable number of cells.*
This requires a more versatile notion of matching,



Challenges Down the Road

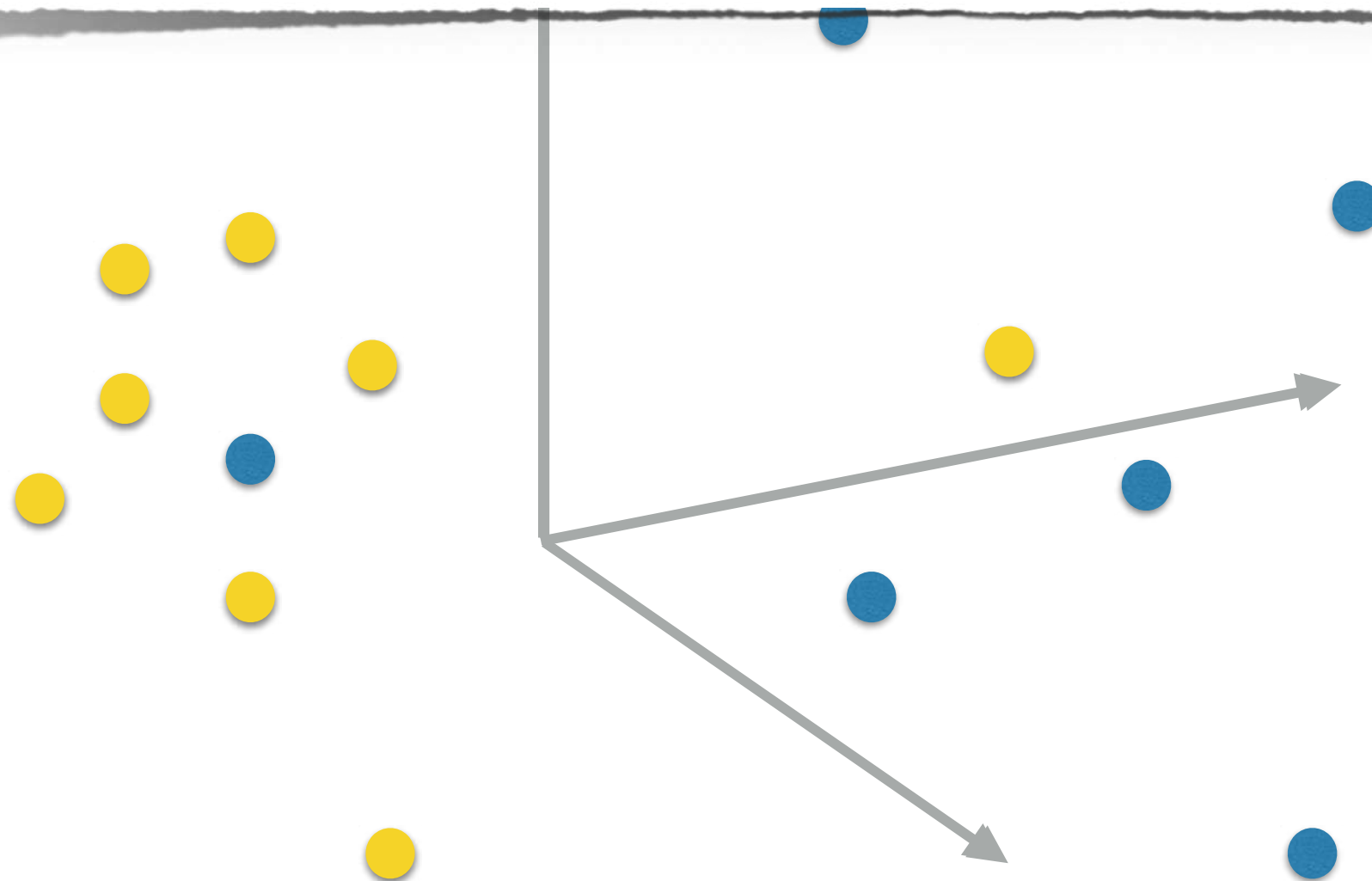
More likely to happen: *variable number of cells.*
This requires a more versatile notion of matching,



Challenges Down the Road

More likely to happen: *variable number of cells.*

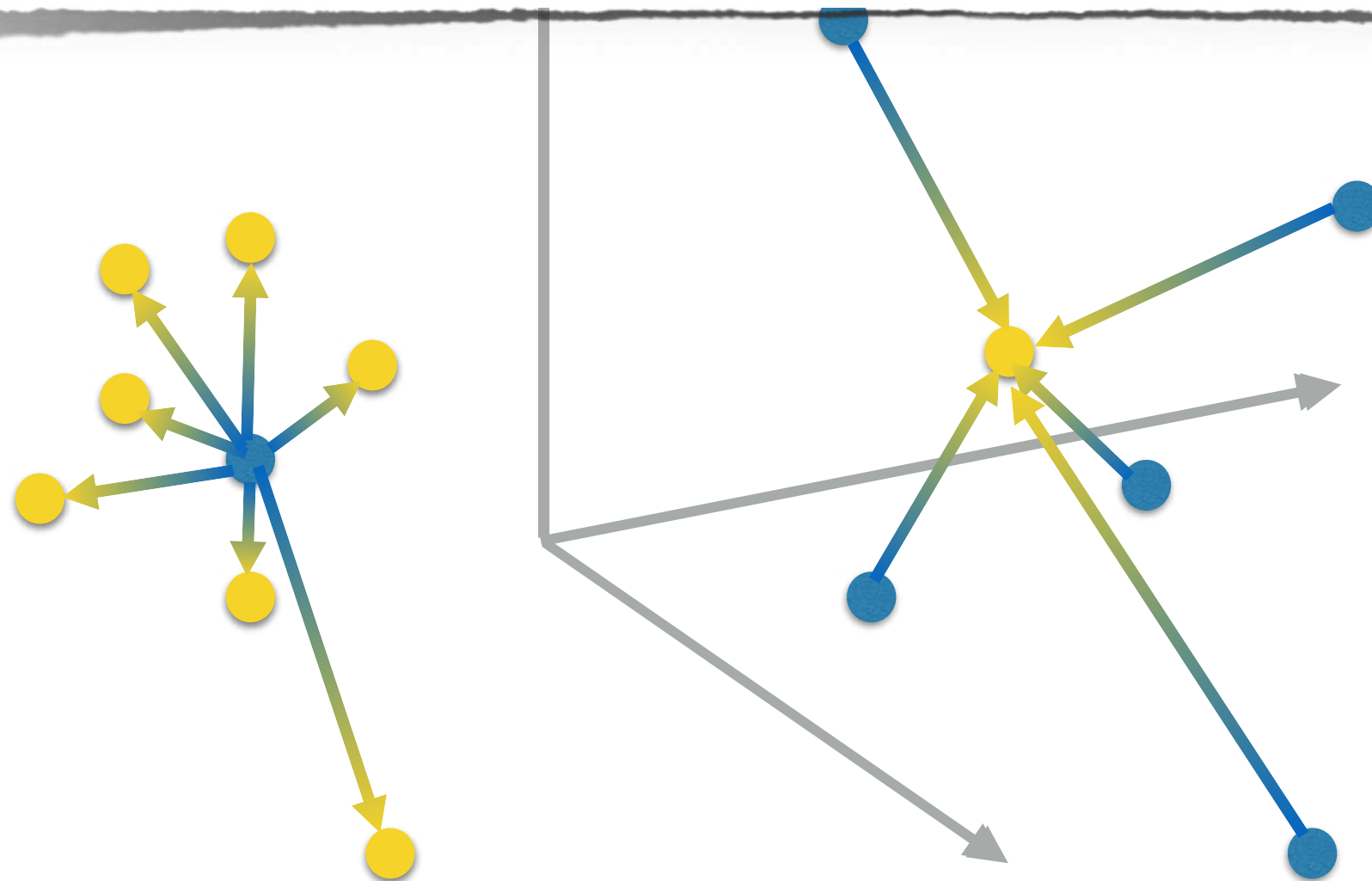
This requires a more versatile notion of matching, one that is more robust than simple nearest-neighbor.



Challenges Down the Road

More likely to happen: *variable number of cells.*

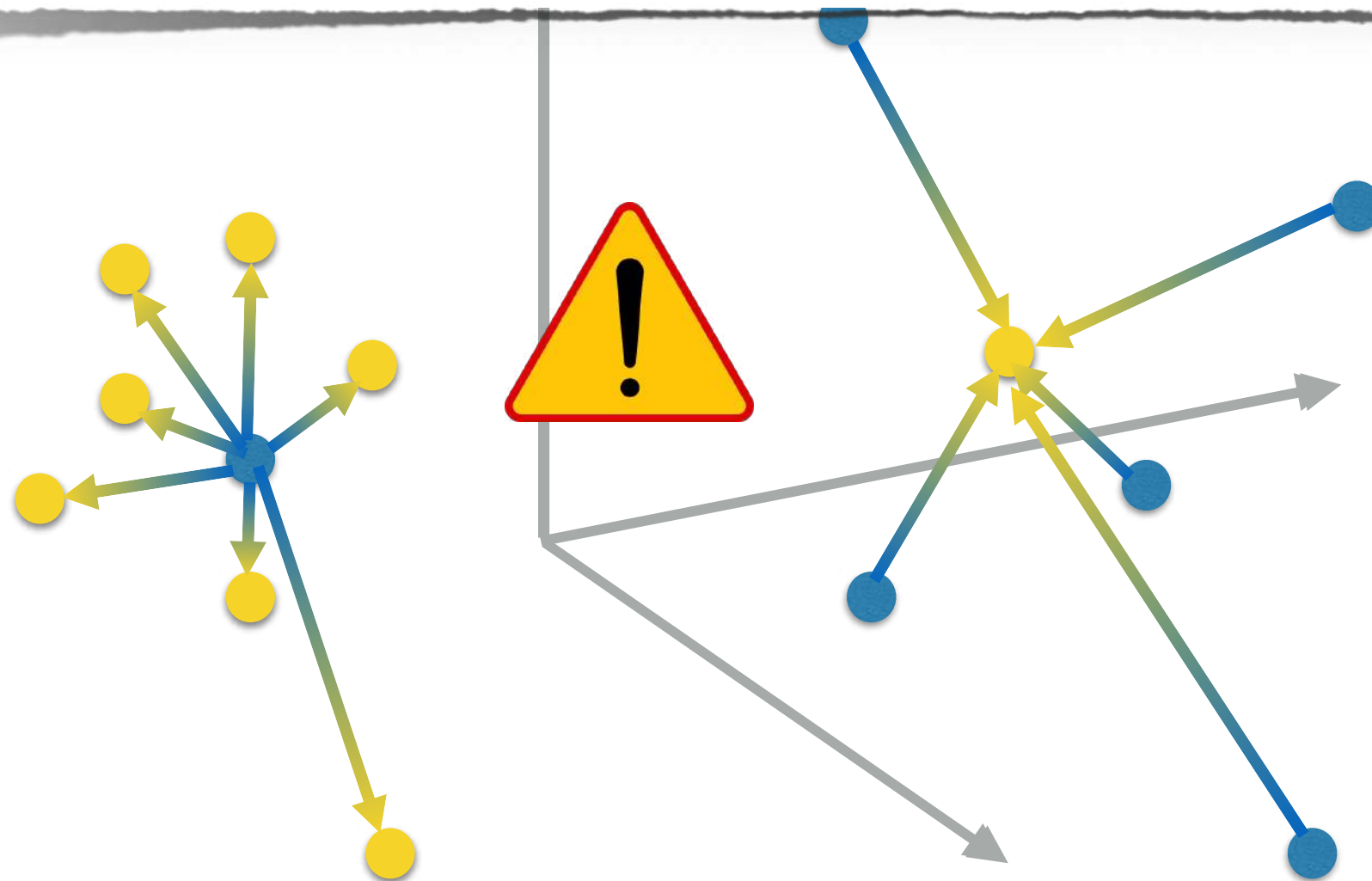
This requires a more versatile notion of matching, one that is more robust than simple nearest-neighbor.



Challenges Down the Road

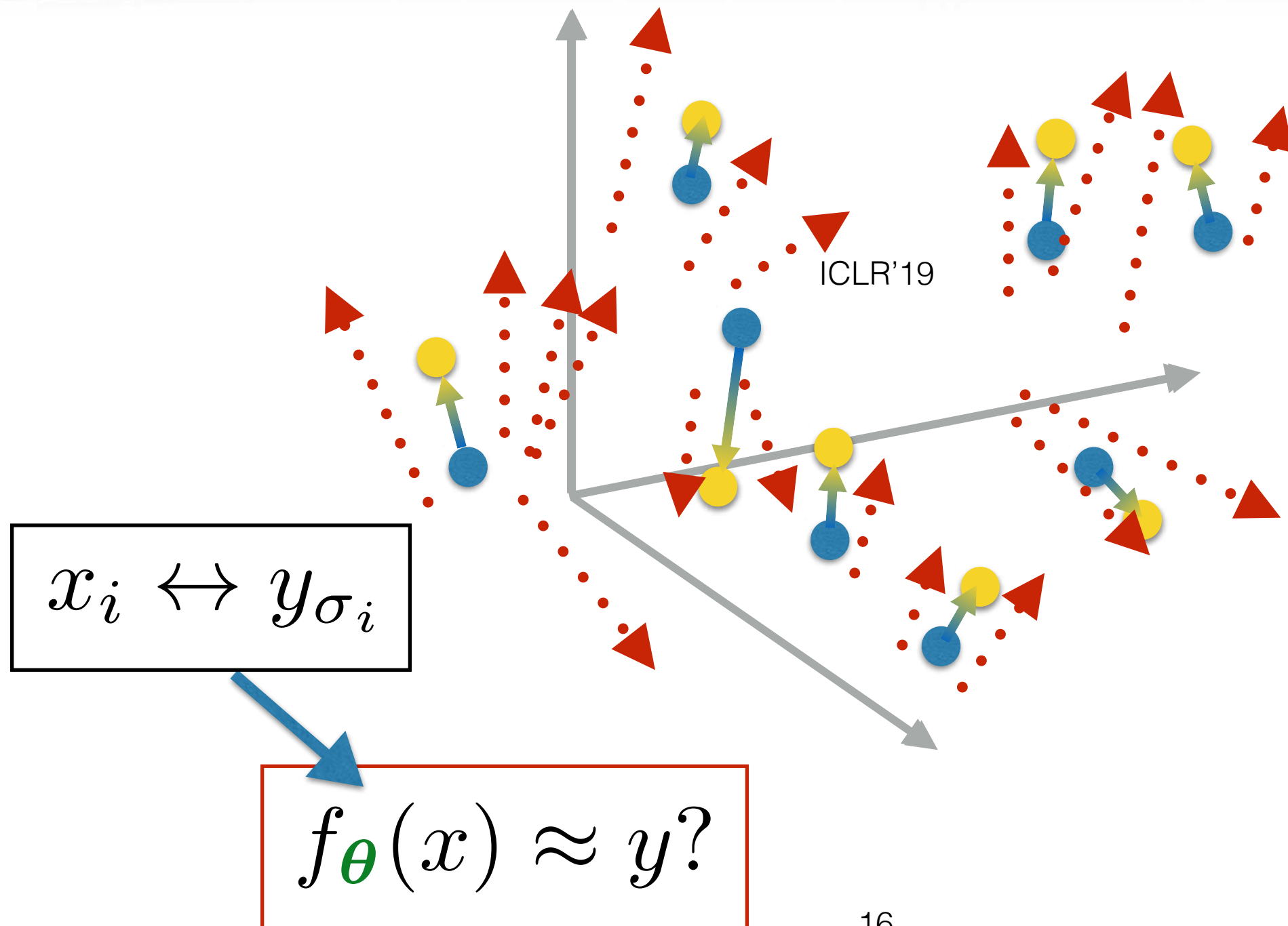
More likely to happen: *variable number of cells.*

This requires a more versatile notion of matching, one that is more robust than simple nearest-neighbor.



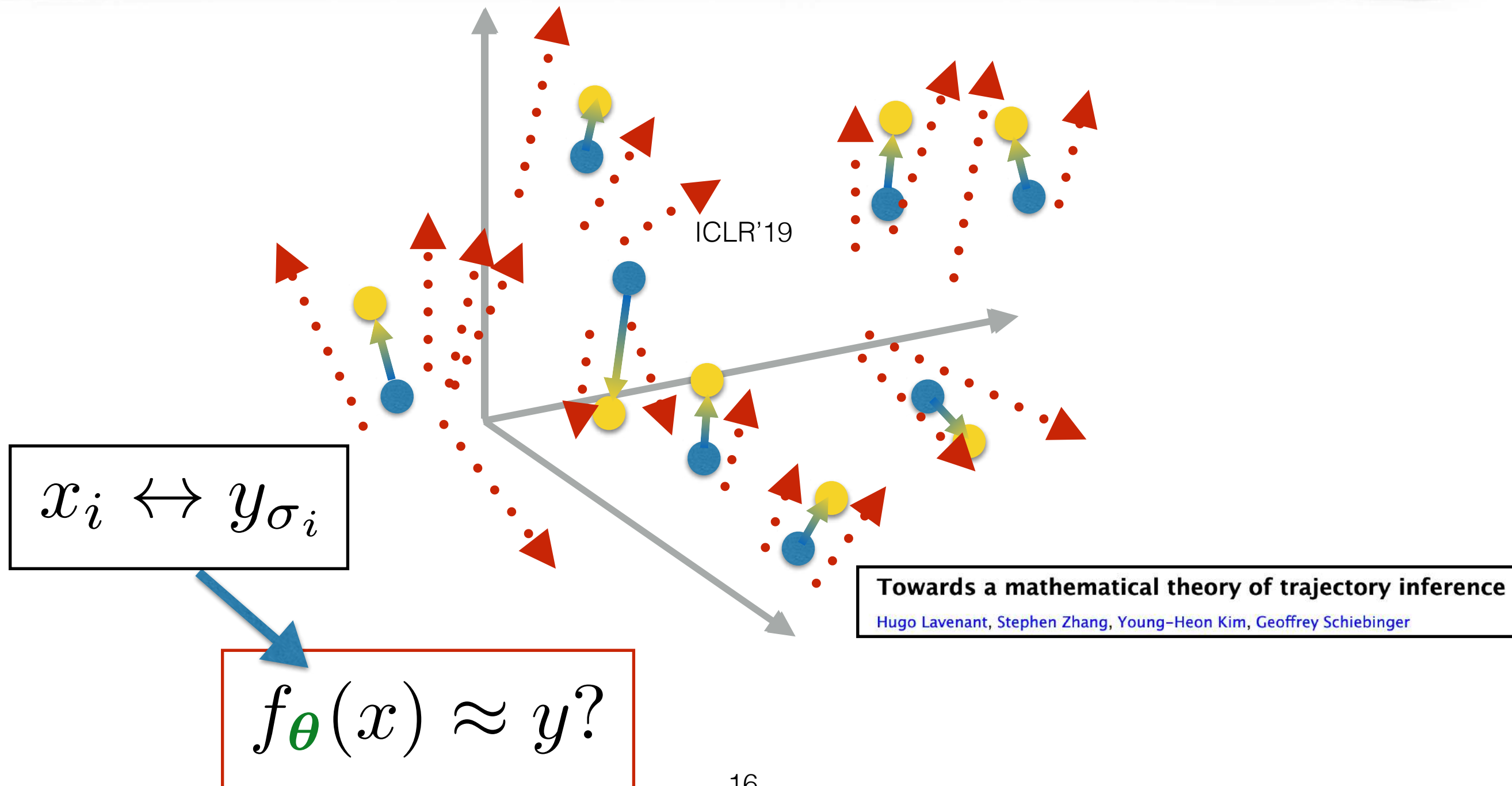
Challenges Down the Road

A broader goal: *infer functions able to provide out of sample mappings, i.e. a continuous extension.*



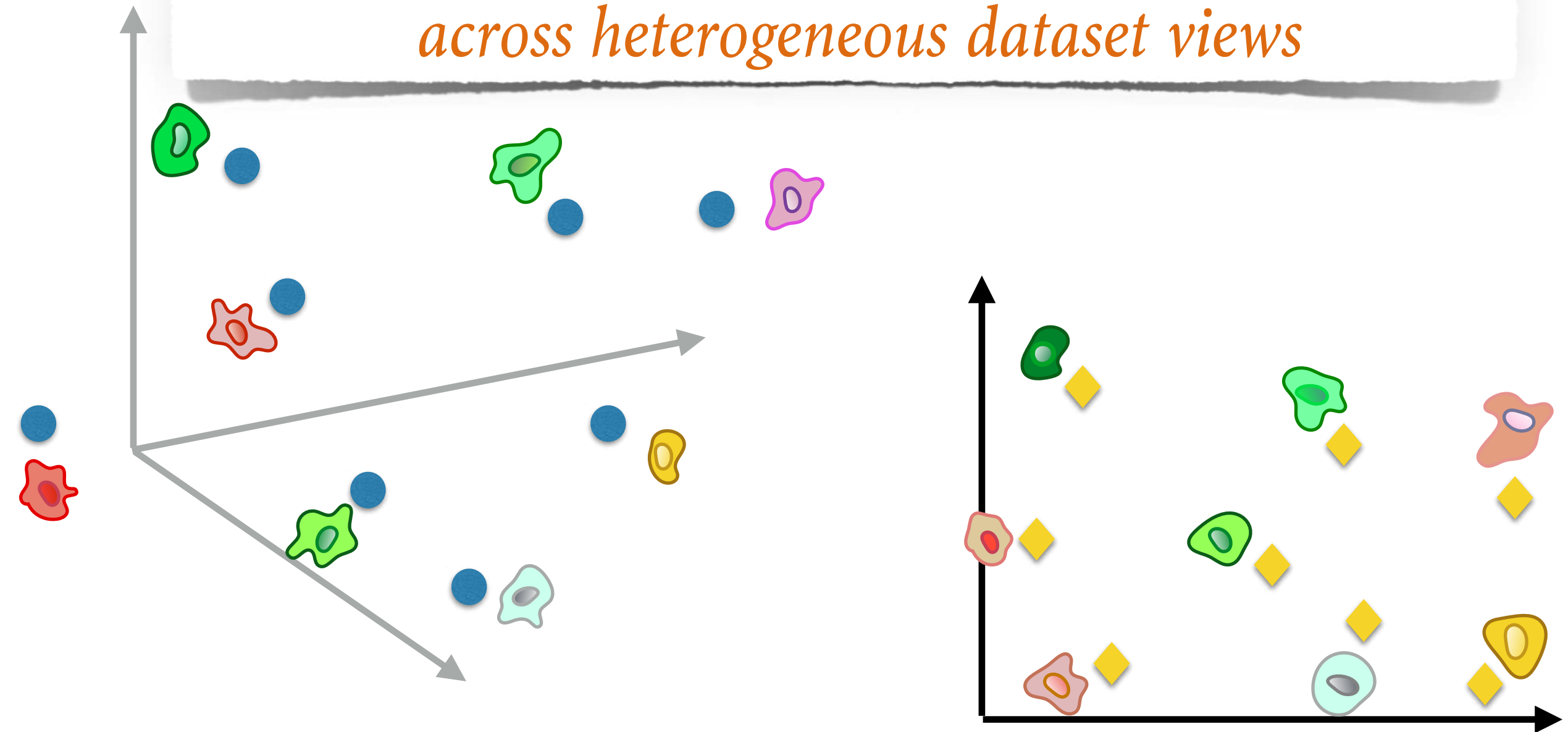
Challenges Down the Road

A broader goal: *infer functions able to provide out of sample mappings, i.e. a continuous extension.*



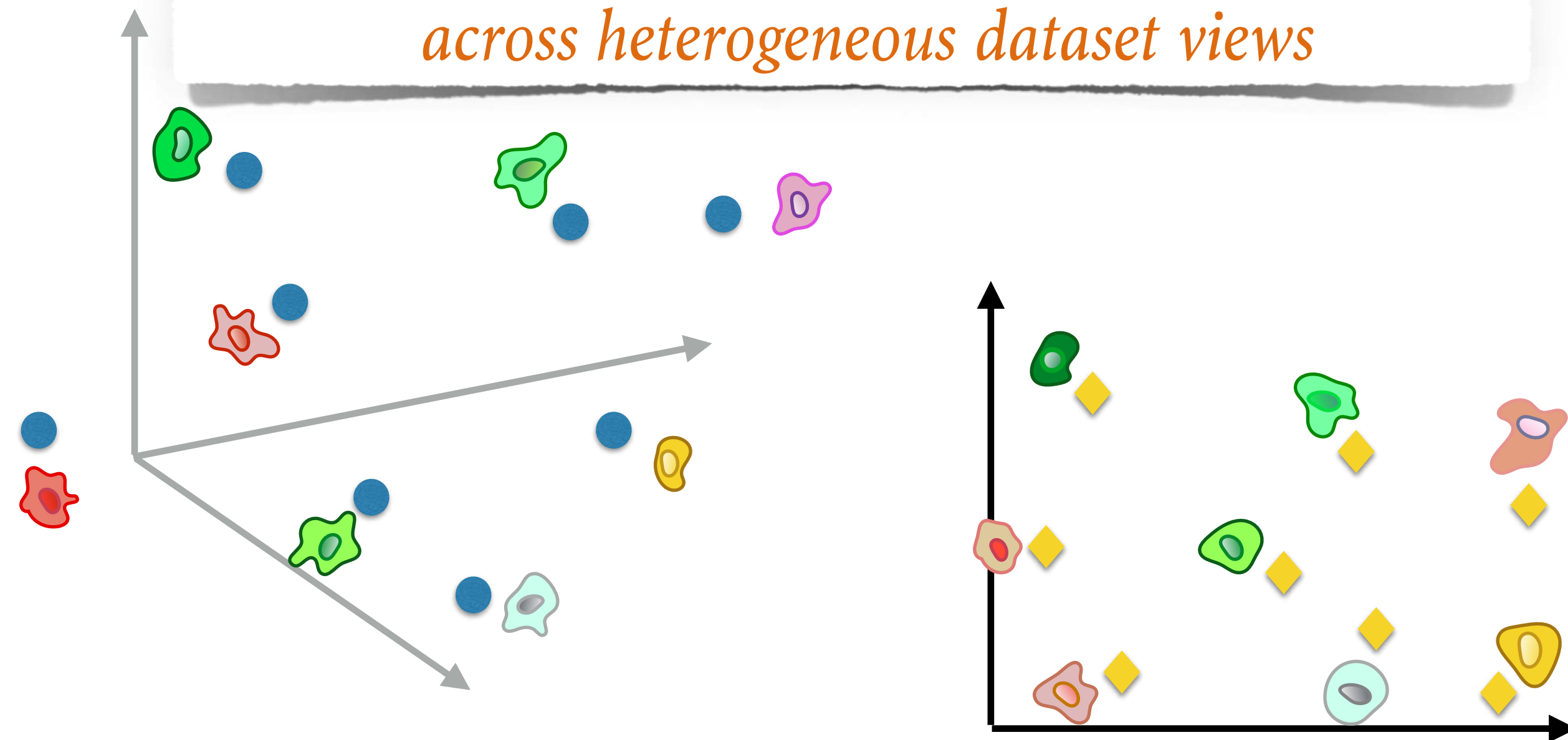
Challenges Down the Road

Add even more flexibility: *match points across heterogeneous dataset views*



Challenges Down the Road

Add even more flexibility: *match points across heterogeneous dataset views*



Gromov-Wasserstein optimal transport to align single-cell multi-omics data

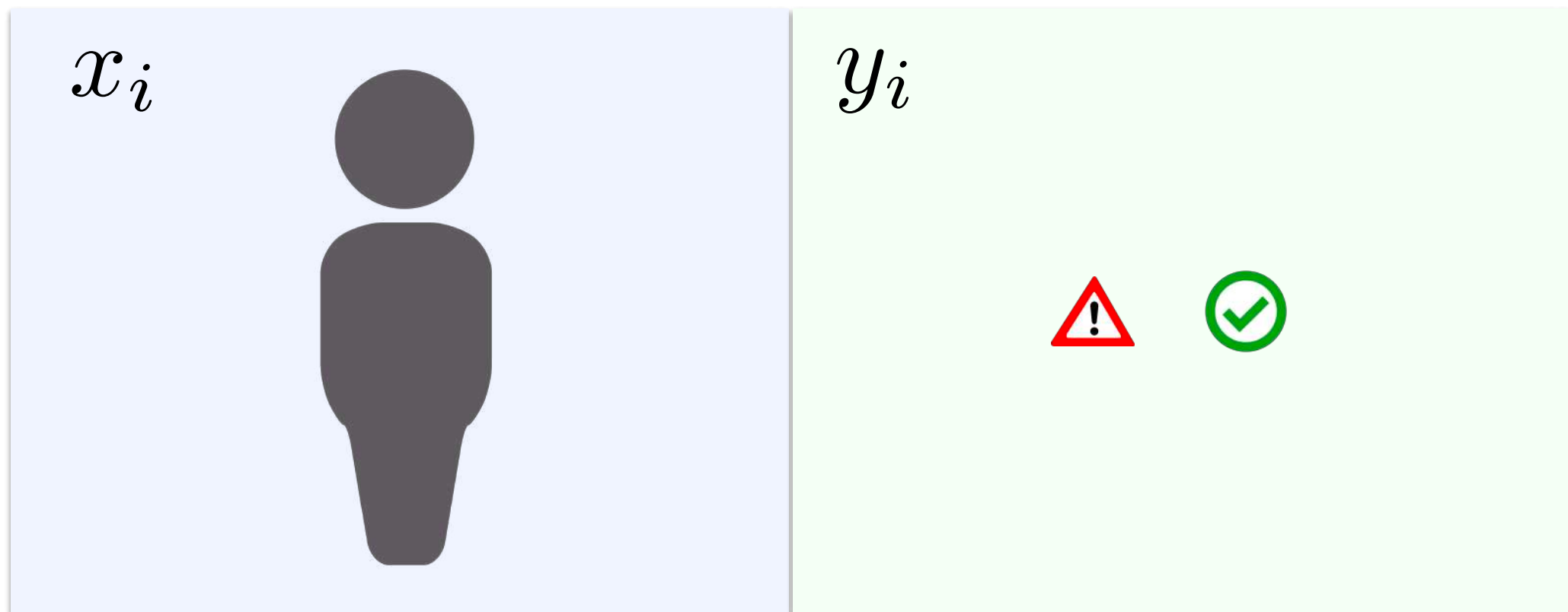
Pinar Demetci, Rebecca Santorella, Björn Sandstede,  William Stafford Noble, Ritambhara Singh

doi: <https://doi.org/10.1101/2020.04.28.066787>

This article is a preprint and has not been certified by peer review [what does this mean?].

An Example from Fairness

Person x_i , characteristic y_i we wish to predict.

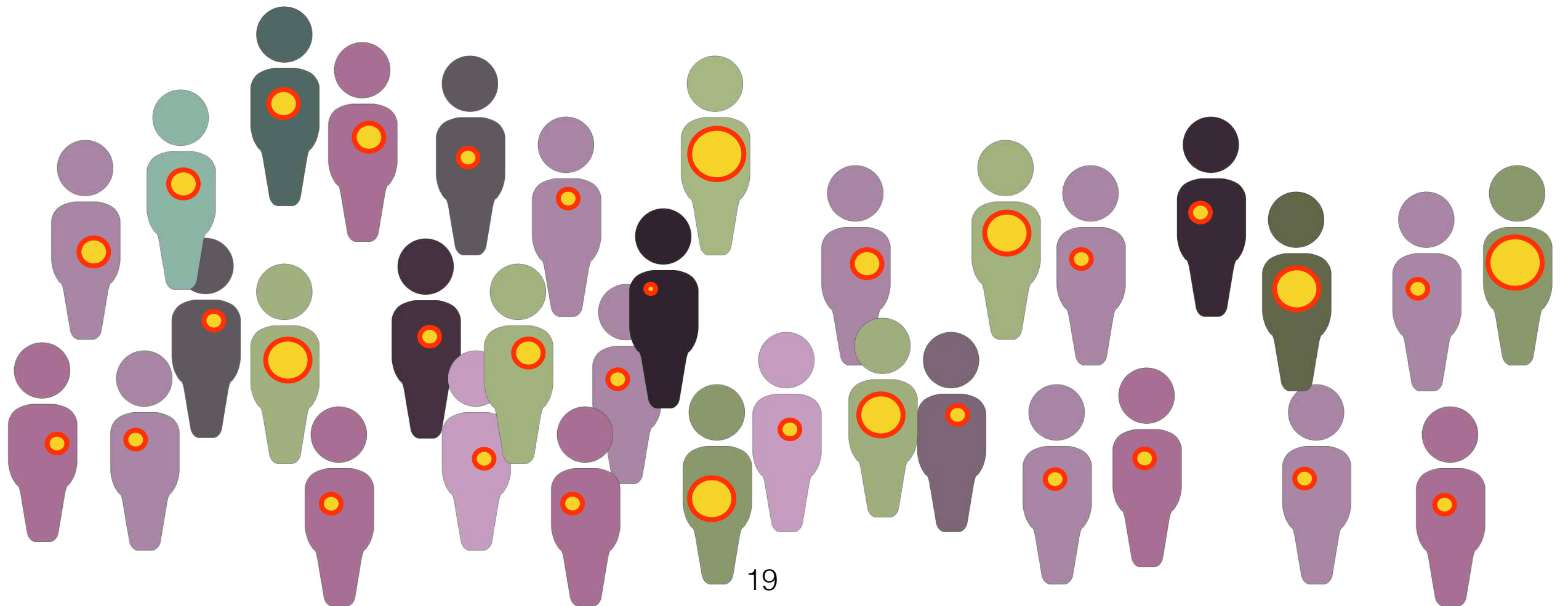


Find θ such that $f_{\theta}(x_i) \approx y_i$
typically $\min_{\theta} \sum_i \ell(f_{\theta}(x_i), y_i)$

An Example from Fairness

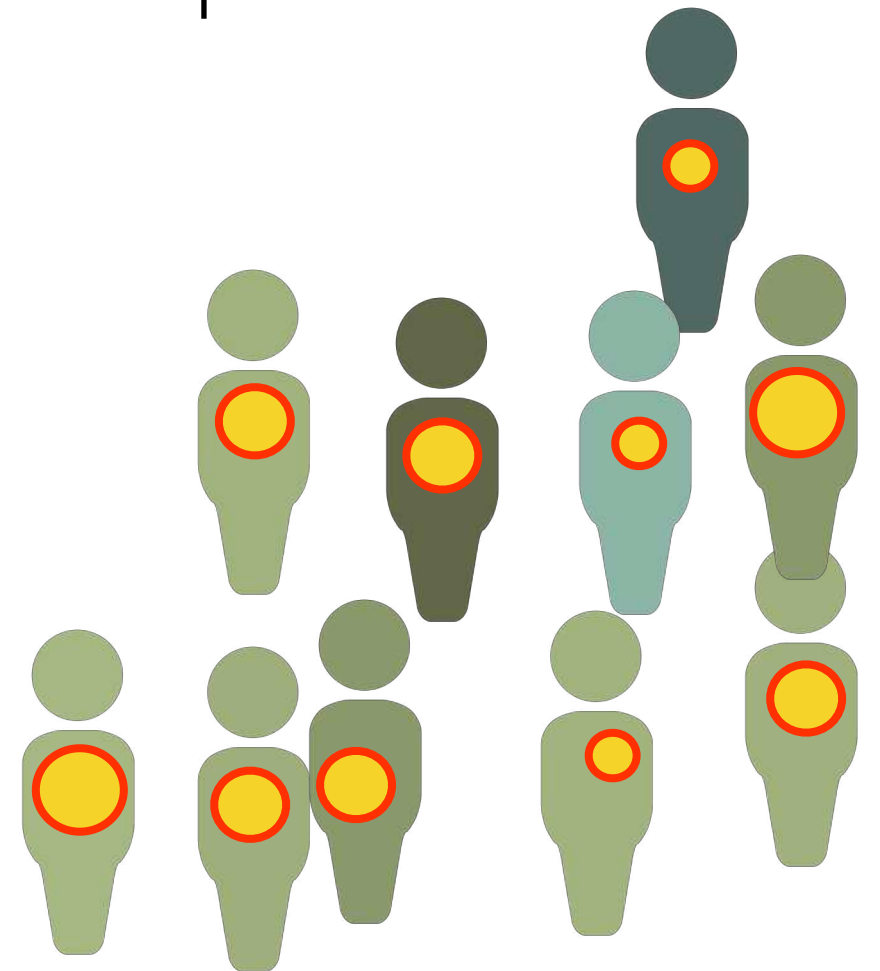
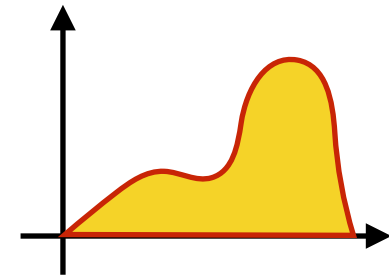
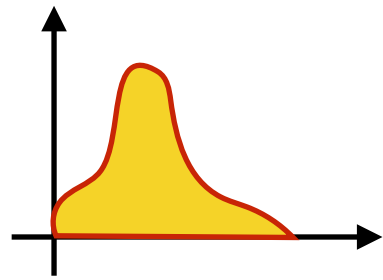
Found θ^* such that $f_{\theta^*}(x_i) \approx y_i$

For each individual i , a different error $\ell(f_{\theta^*}(x_i), y_i)$.



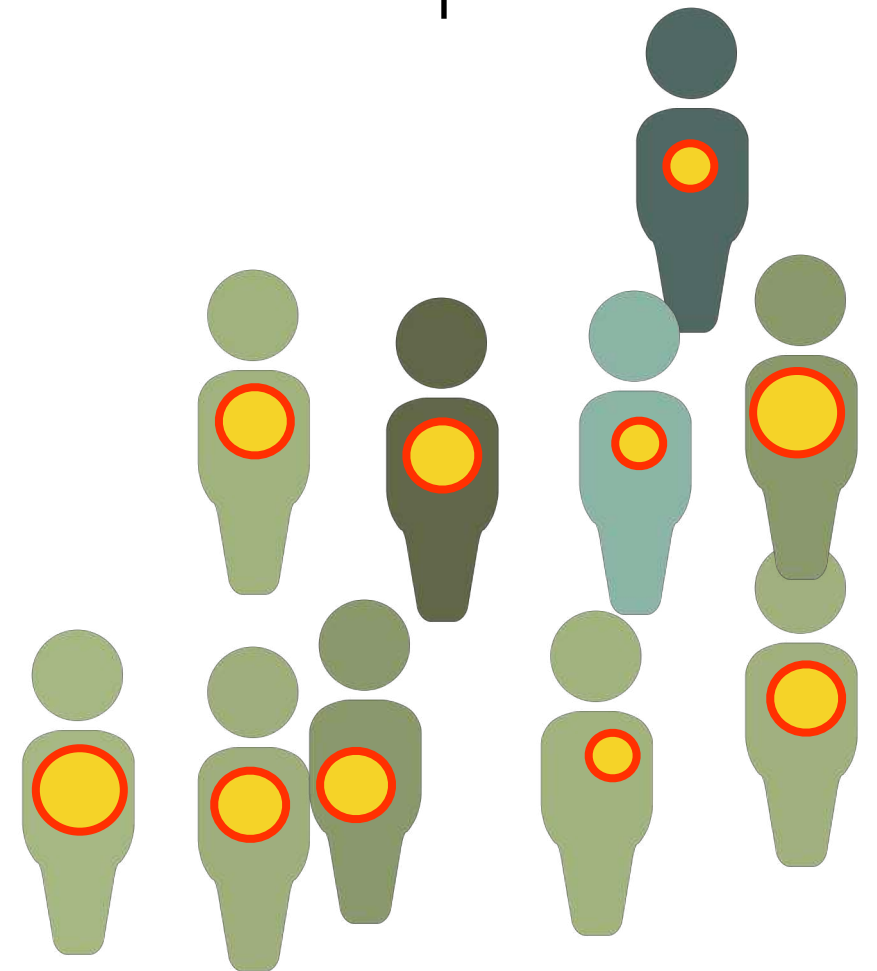
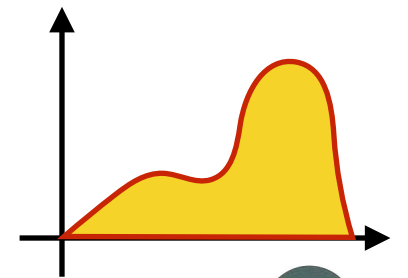
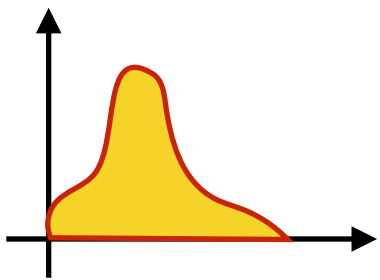
An Example from Fairness

Obviously, θ^* achieved performance by skewing distribution of $\ell(f_{\theta}(x_i), y_i)$



An Example from Fairness

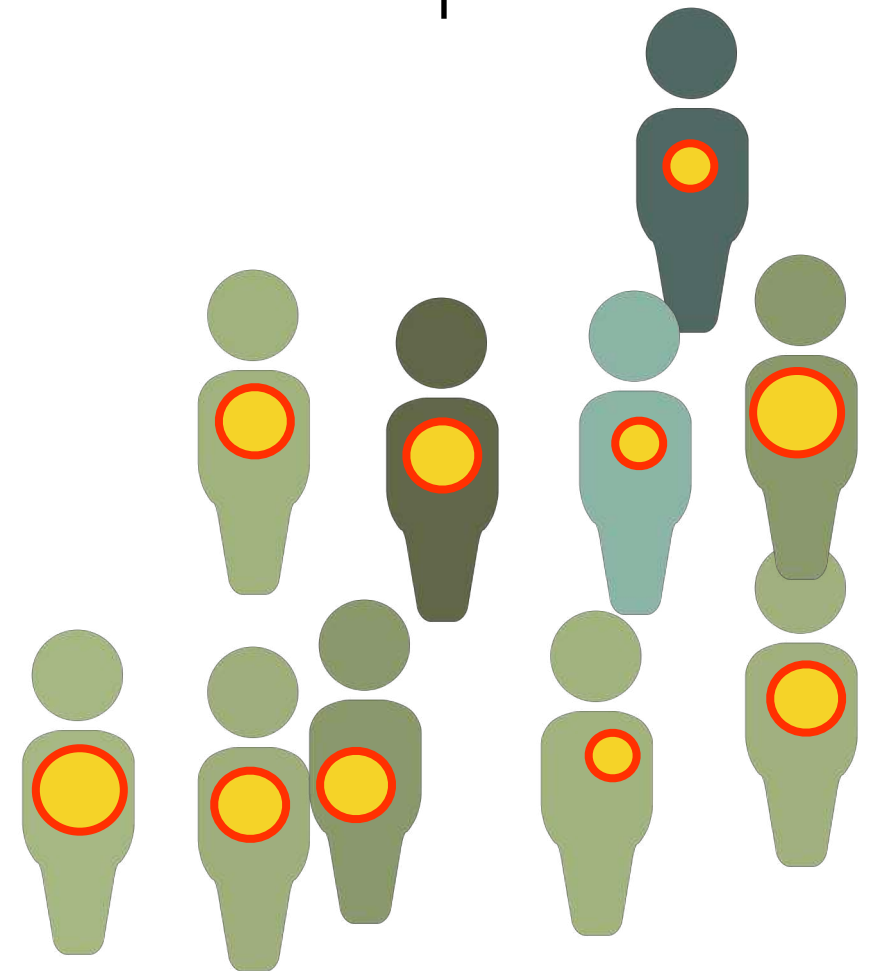
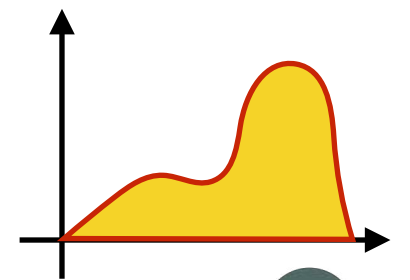
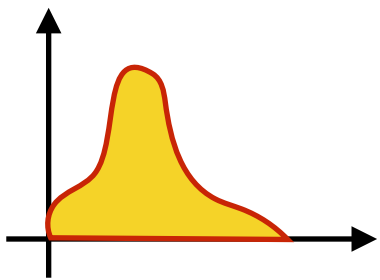
Goal: *enforce fairness mechanisms.*



An Example from Fairness

Goal: *enforce fairness mechanisms.*

Problem: *constraints are distributional, not individual*

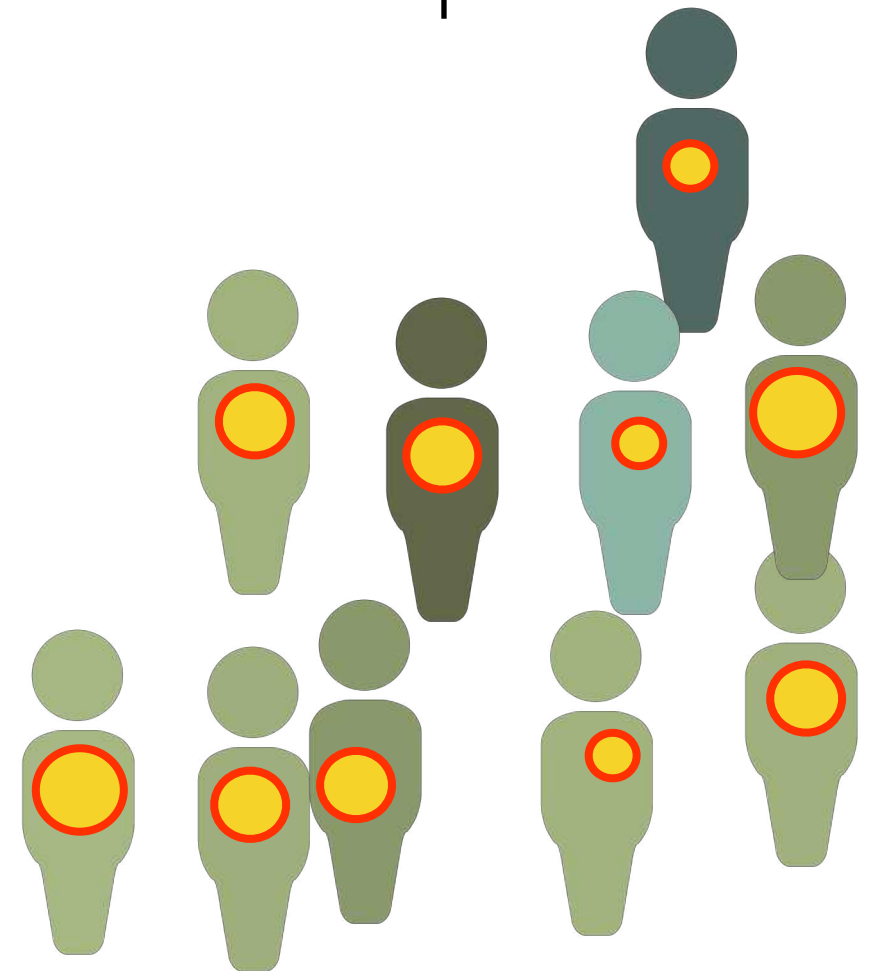
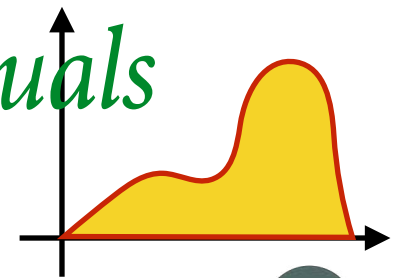
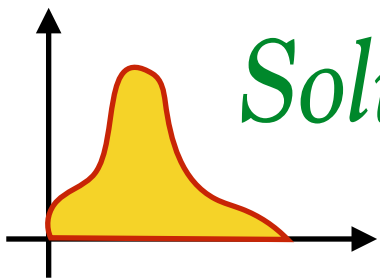


An Example from Fairness

Goal: *enforce fairness mechanisms.*

Problem: *constraints are distributional, not individual*

*Solution: treat similarly **matched** individuals*

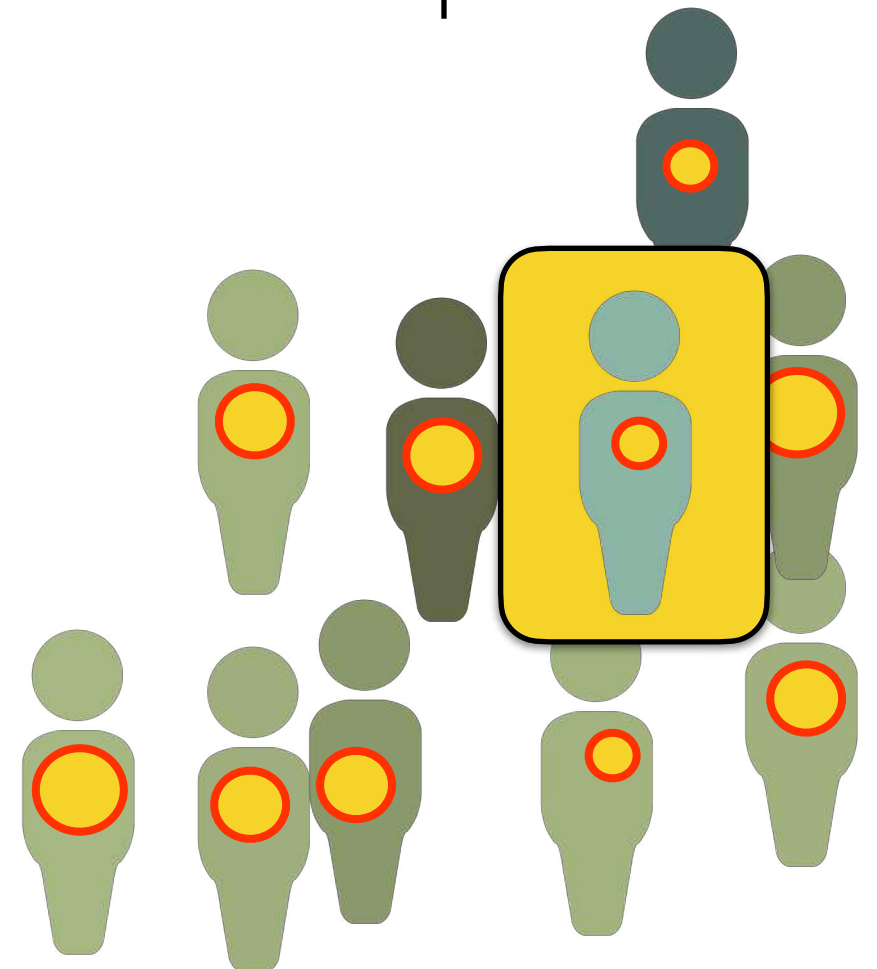
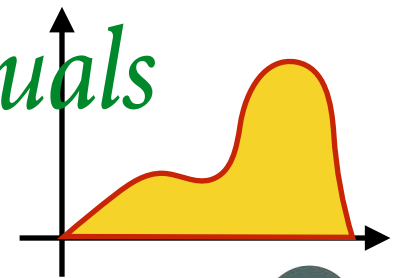
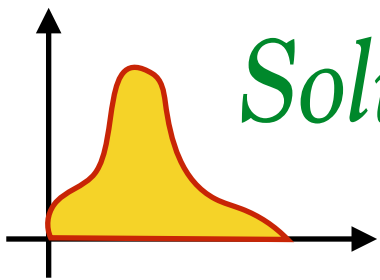


An Example from Fairness

Goal: *enforce fairness mechanisms.*

Problem: *constraints are distributional, not individual*

*Solution: treat similarly **matched** individuals*

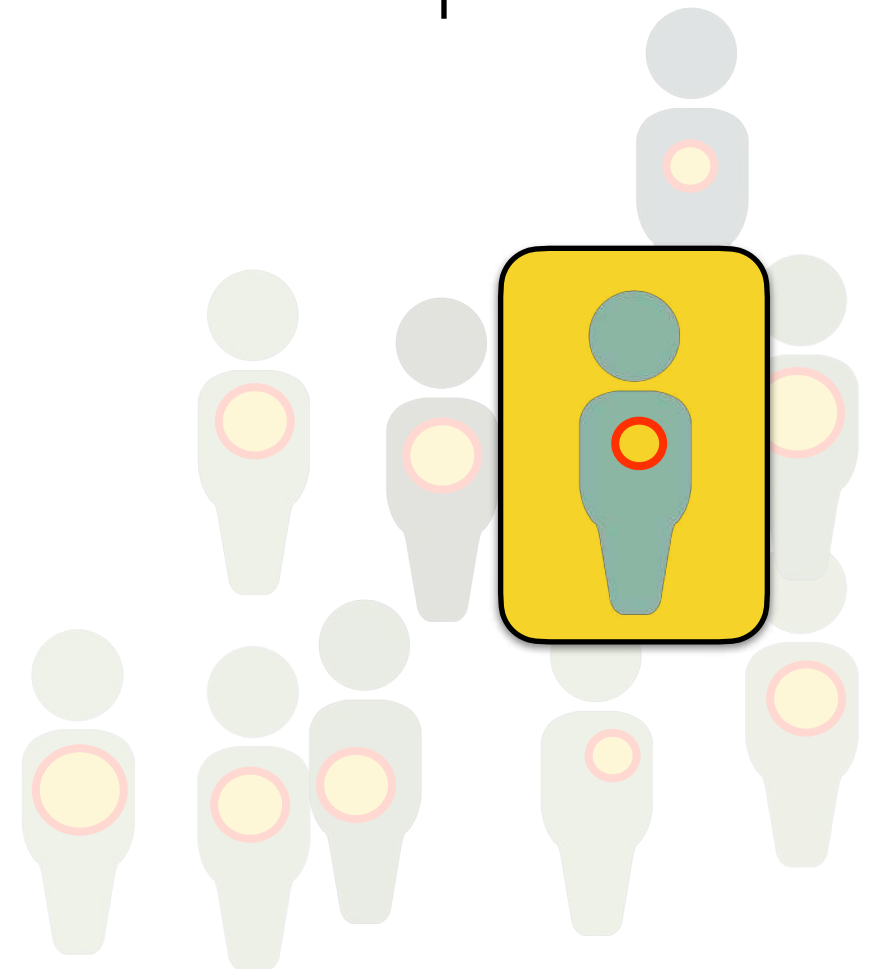
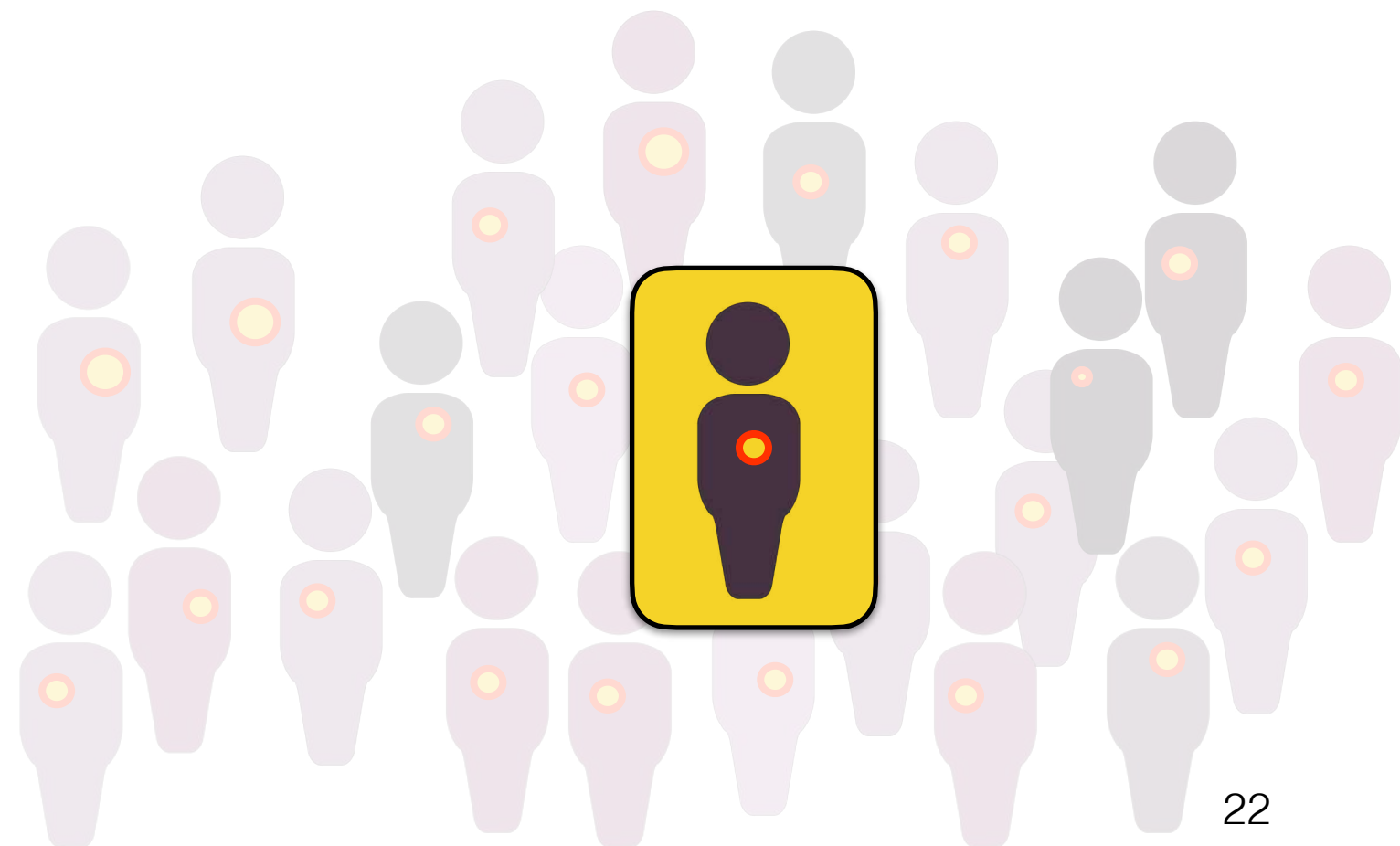
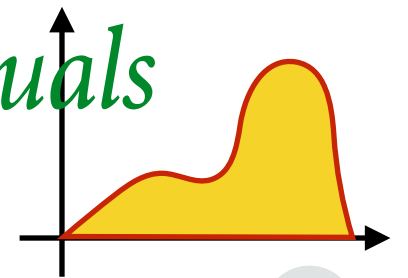
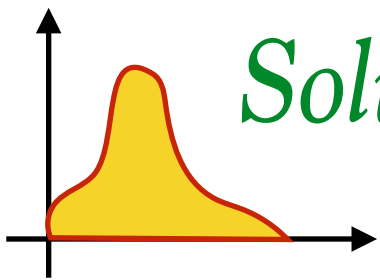


An Example from Fairness

Goal: *enforce fairness mechanisms.*

Problem: *constraints are distributional, not individual*

*Solution: treat similarly **matched** individuals*



An Example from Fairness

Goal: *enforce fairness*

fairness optimal transport

About 72,400 results (0.05 sec; Showing 50 m

Fairness with Overlapping Groups

Forest Yang, Moustapha Cisse, Sanmi Koyejo

hor name ☐ Word matching ☒

Show the calculator

Obtaining fairness using optimal transport theory

[P Gordaliza](#), [E Del Barrio](#), [G Fabrice](#) - ... on *Machine Learning*, 2019 - [proceedings.mlr.press](#)
In the fair classification setup, we recast the links between **fairness** and predictability in terms of probability metrics. We analyze repair methods based on mapping conditional distributions to the Wasserstein barycenter. We propose a Random Repair which yields a ...

☆ [Cited by 55](#) [Related articles](#) [»](#)

FlipTest: fairness testing via optimal transport

[E Black](#), [S Yeom](#), [M Fredrikson](#), [Fairness](#) - ... of the 2020 Conference on Fairness ..., 2020 - [proceedings.mlr.press](#)
We present FlipTest, a black-box technique for uncovering discrimination in classifiers. FlipTest is motivated by the intuitive question: had an individual been of a different protected status, would the model have treated them differently? Rather than relying on causal ...

☆ [Cited by 18](#) [Related articles](#) [All 4 versions](#)

Obtaining fairness using optimal transport theory

[E Del Barrio](#), [F Gamboa](#), [P Gordaliza](#) - *arXiv preprint arXiv* ..., 2018 - [arxiv.org](#)
Statistical algorithms are usually helping in making decisions in many aspects of our lives. But, how do we know if these algorithms are biased and commit unfair discrimination of a particular group of people, typically a minority? **Fairness** is generally studied in a ...

☆ [Cited by 11](#) [Related articles](#) [All 11 versions](#) [»](#)

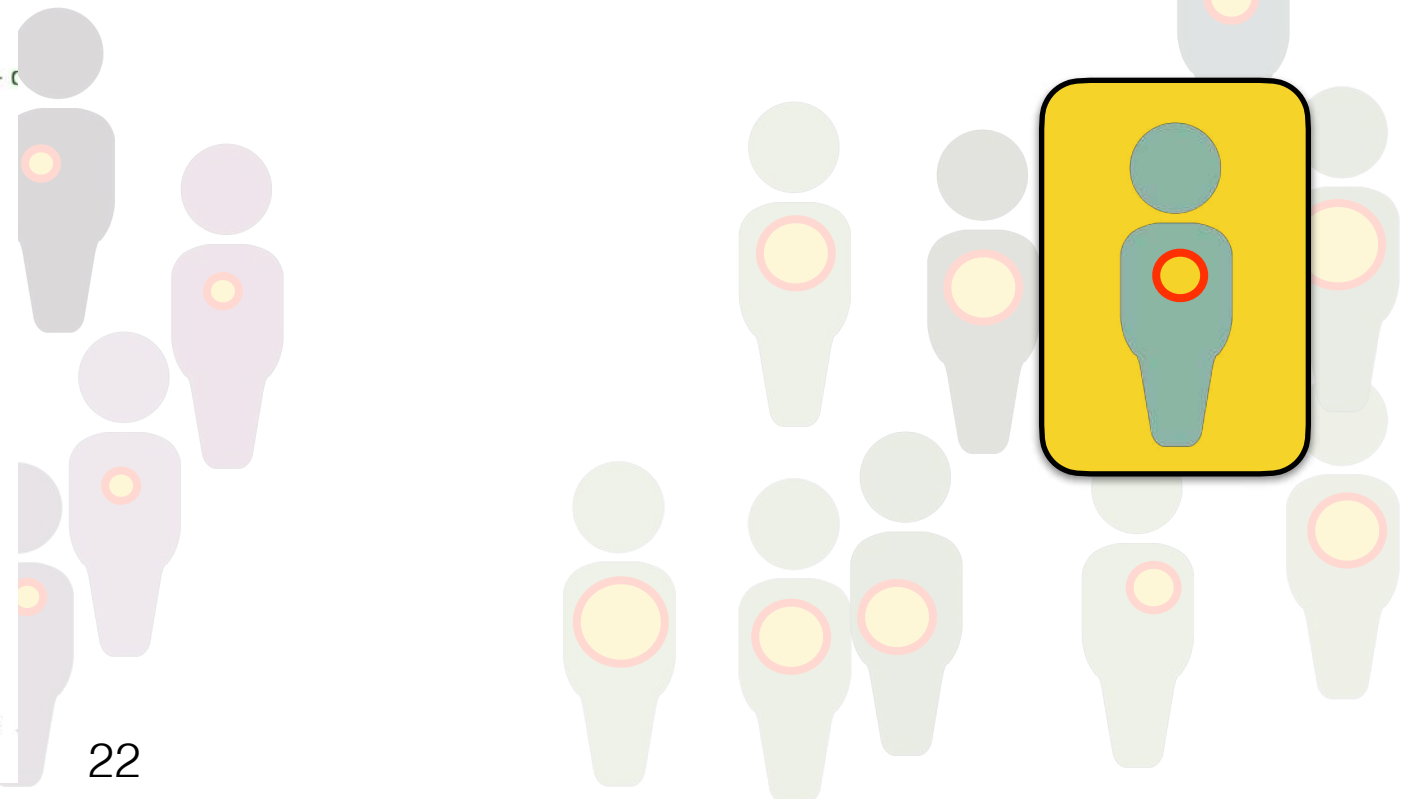
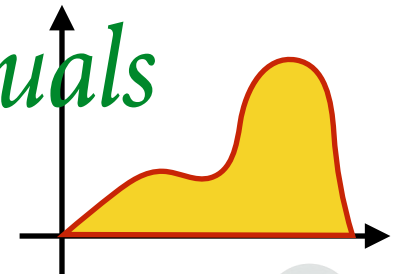
A general approach to fairness with optimal transport

[C Silvia](#), [J Ray](#), [S Tom](#), [P Aldo](#), [J Heinrich](#) - *Proceedings of the AAAI* ..., 2020 - [ojs.aaai.org](#)
We propose a general approach to **fairness** based on transporting distributions corresponding to different sensitive attributes to a common distribution. We use **optimal transport** theory to derive target distributions and methods that allow us to achieve **fairness**

☆ [Cited by 7](#) [Related articles](#) [All 2 versions](#) [»](#)

Problem: *constraints are distributional, not individual*

early matched individuals



An Example from Fairness

Goal: *enforce fairness*

fairness optimal transport

About 72,400 results (0.05 sec; Showing 50 m

Fairness with Overlapping Groups

Forest Yang, Moustapha Cisse, Sanmi Koyejo

hor name ☐ Word matching ☒

Show the calculator

Obtaining fairness using optimal transport theory

[P Gordaliza, E Del Barrio, G Fabrice](#) - ... on Machine Learning, 2019 - proceedings.mlr.press

In the fair classification setup, we recast the links between **fairness** and predictability in terms of probability metrics. We analyze repair methods based on mapping conditional distributions to the Wasserstein barycenter. We propose a Random Repair which yields a ...

☆ [Cited by 55](#) [Related articles](#) [»](#)

FlipTest: fairness testing via optimal transport

[E Black, S Yeom, M Fredrikson, Fairness](#) - ... of the 2020 Conference on Fairness ..., 2020 -

We present FlipTest, a black-box technique for uncovering discrimination in classifiers. FlipTest is motivated by the intuitive question: had an individual been of a different protected status, would the model have treated them differently? Rather than relying on causal ...

☆ [Cited by 18](#) [Related articles](#) [All 4 versions](#)

Obtaining fairness using optimal transport theory

[E Del Barrio, F Gamboa, P Gordaliza](#) - arXiv preprint arXiv ..., 2018 - arxiv.org

Statistical algorithms are usually helping in making decisions in many aspects of our lives. But, how do we know if these algorithms are biased and commit unfair discrimination of a particular group of people, typically a minority? **Fairness** is generally studied in a ...

☆ [Cited by 11](#) [Related articles](#) [All 11 versions](#) [»](#)

A general approach to fairness with optimal transport

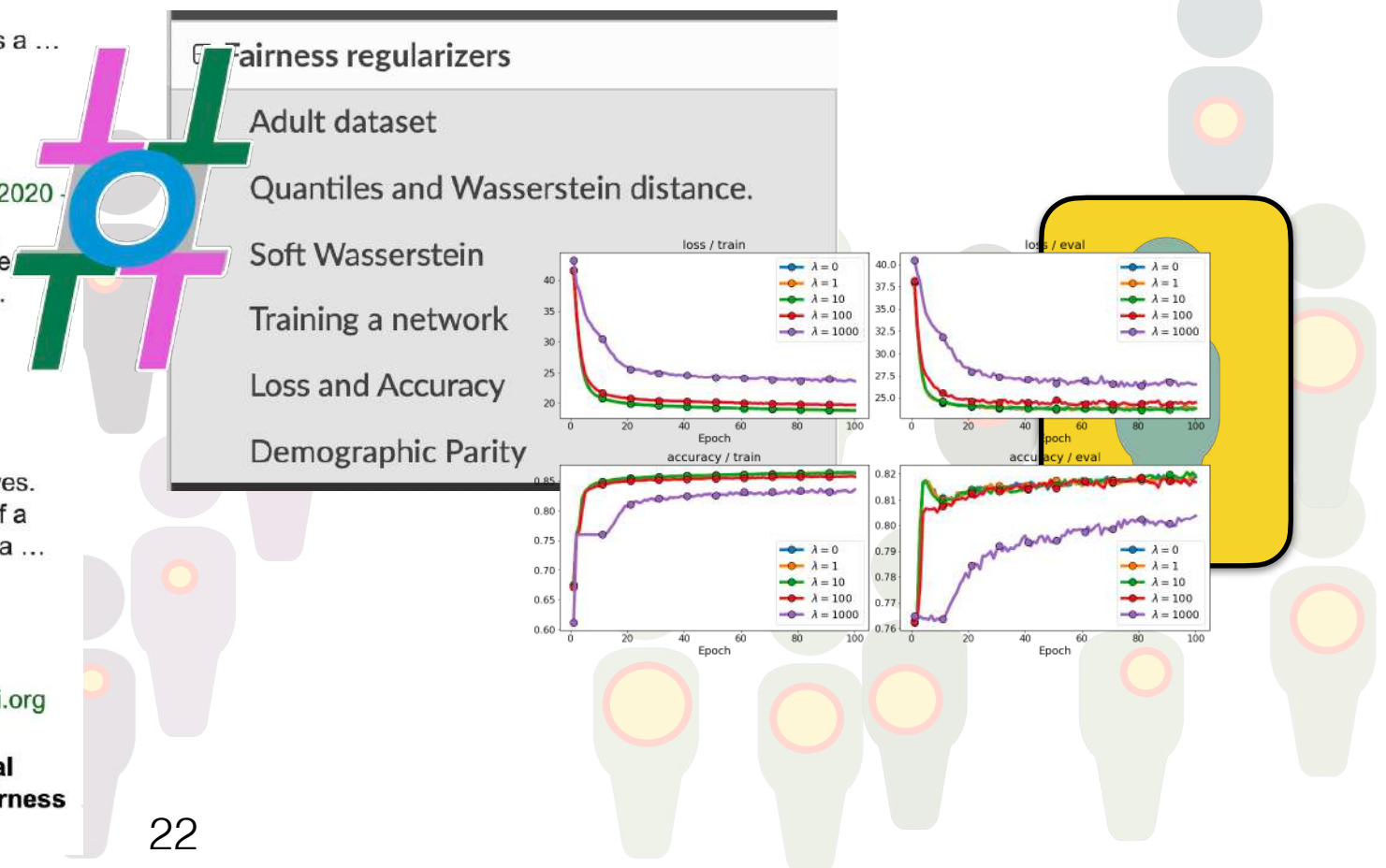
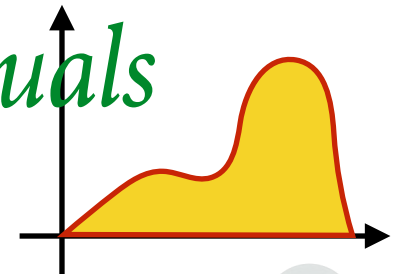
[C Silvia, J Ray, S Tom, P Aldo, J Heinrich](#) - Proceedings of the AAAI ..., 2020 - ojs.aaai.org

We propose a general approach to **fairness** based on transporting distributions corresponding to different sensitive attributes to a common distribution. We use **optimal transport** theory to derive target distributions and methods that allow us to achieve **fairness**

☆ [Cited by 7](#) [Related articles](#) [All 2 versions](#) [»](#)

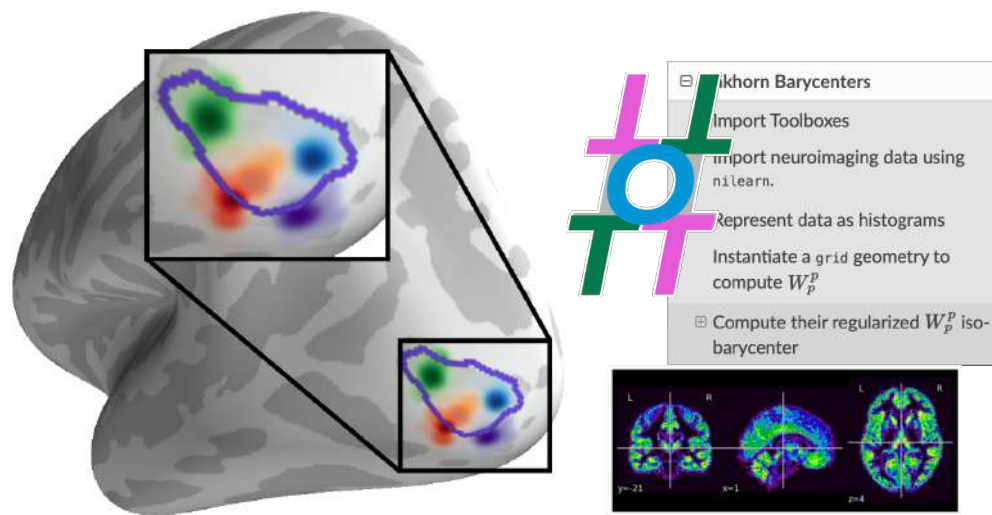
Problem: *constraints are distributional, not individual*

erly matched individuals



Optimal matchings are needed everywhere!

Goal: *predict from brain activation maps*

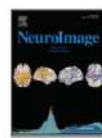


Problem: *no two persons have the same brain, need to match first*



NeuroImage

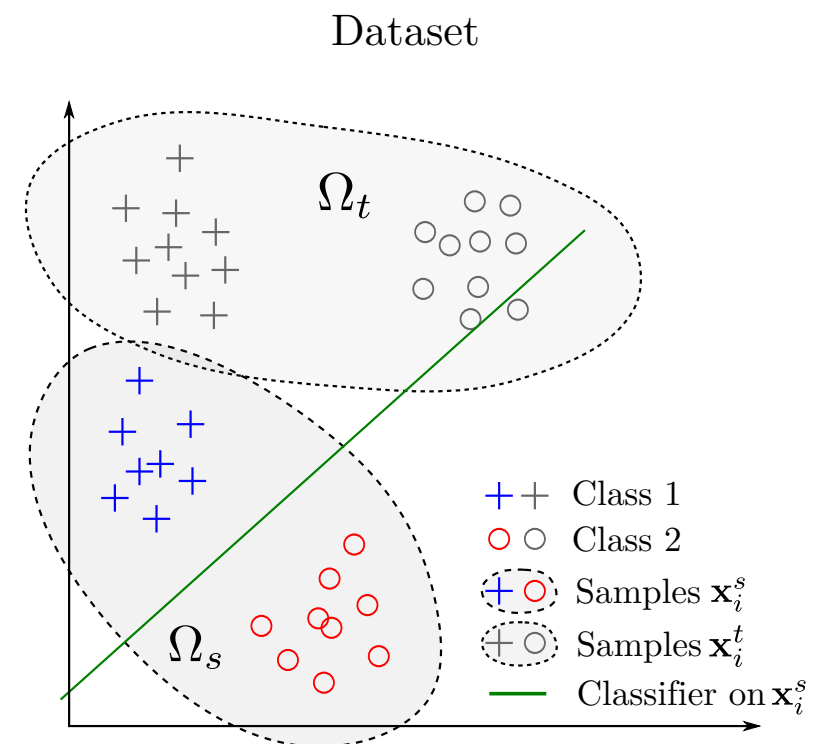
Volume 220, 15 October 2020, 116847



Multi-subject MEG/EEG source imaging with sparse multi-task regression

Hicham Janati ^{a, c, *}, Thomas Bazeille ^{a, *}, Bertrand Thirion ^{a, *}, Marco Cuturi ^{b, c}, Alexandre Gramfort ^{a, *}

Goal: *domain adaptation*

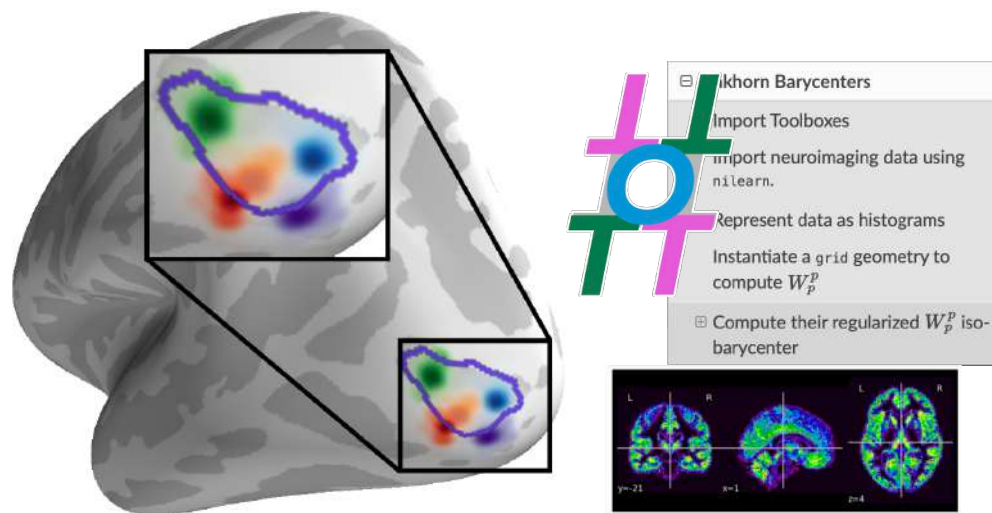


Problem: *train and test data must be matched first*



Optimal matchings are needed everywhere!

Goal: *predict from brain activation maps*

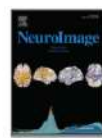


Problem: *no two persons have the same brain, need to match first*



NeuroImage

Volume 220, 15 October 2020, 116847

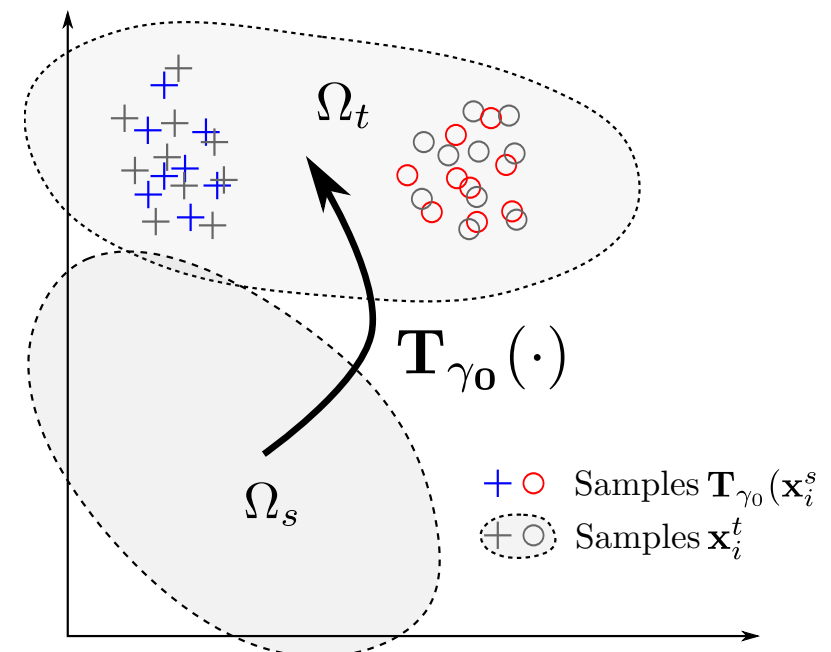


Multi-subject MEG/EEG source imaging with sparse multi-task regression

Hicham Janati ^{a, c, *}, Thomas Bazeille ^a, Bertrand Thirion ^a, Marco Cuturi ^{b, c}, Alexandre Gramfort ^a

Goal: *domain adaptation*

Optimal transport

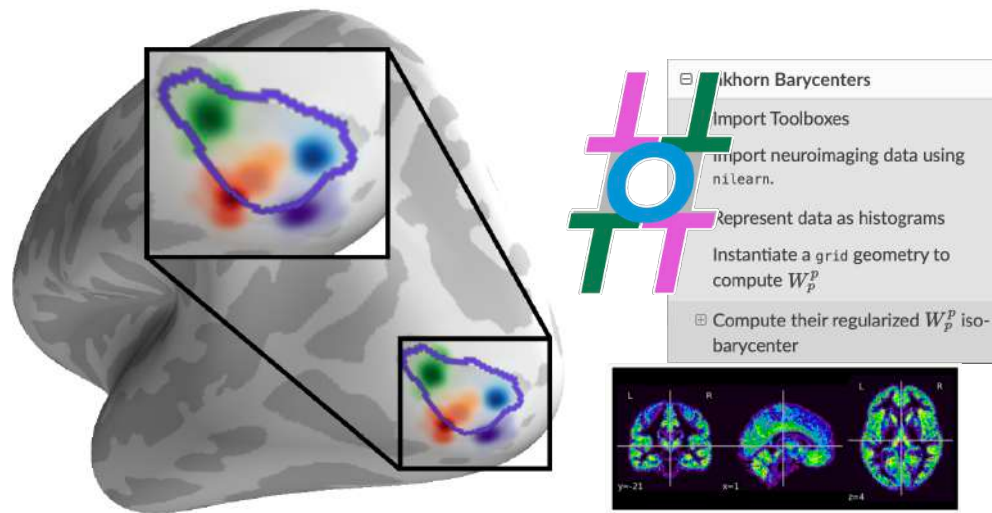


Problem: *train and test data must be matched first*



Optimal matchings are needed everywhere!

Goal: *predict from brain activation maps*

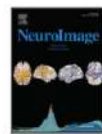


Problem: *no two persons have the same brain, need to match first*



NeuroImage

Volume 220, 15 October 2020, 116847

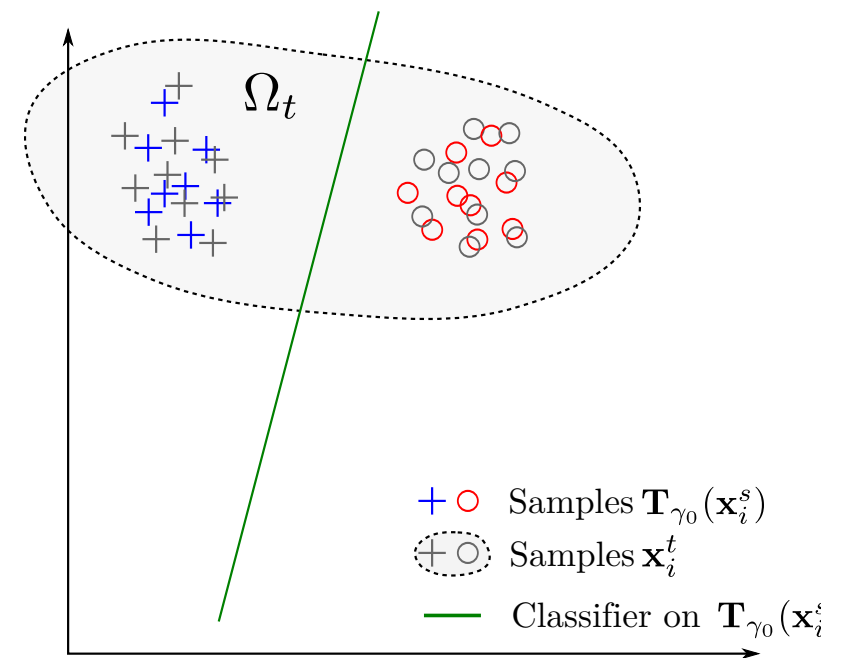


Multi-subject MEG/EEG source imaging with sparse multi-task regression

Hicham Janati ^{a, c, *}, Thomas Bazeille ^a, Bertrand Thirion ^a, Marco Cuturi ^{b, c}, Alexandre Gramfort ^a

Goal: *domain adaptation*

Classification on transported samples



Problem: *train and test data must be matched first*

IEEE TRANSACTIONS ON PATTERN ANALYSIS AND MACHINE INTELLIGENCE, VOL. X, NO. X, JANUARY XX

Optimal Transport for Domain Adaptation

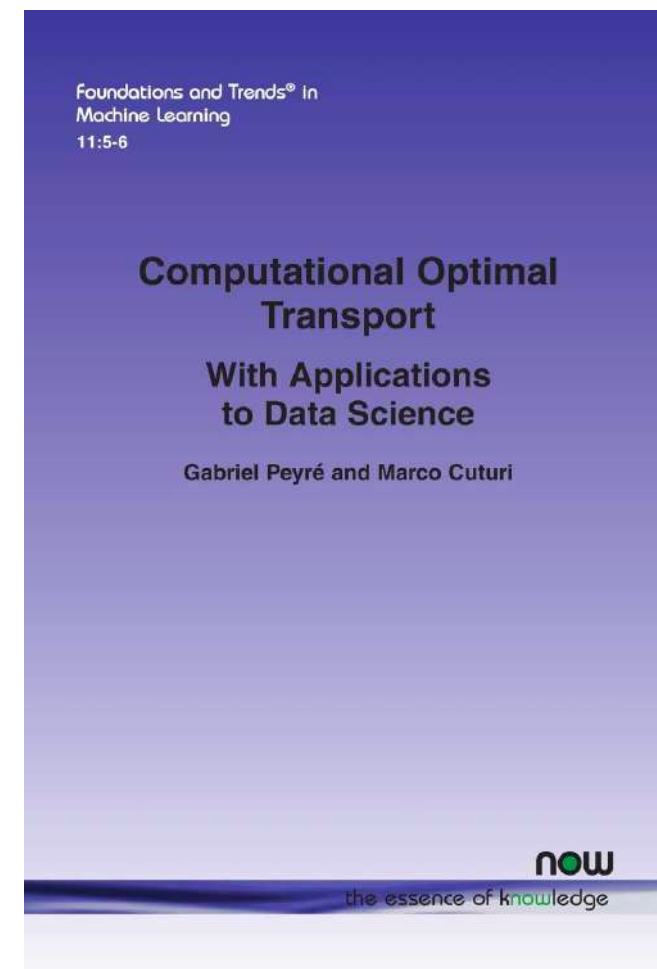
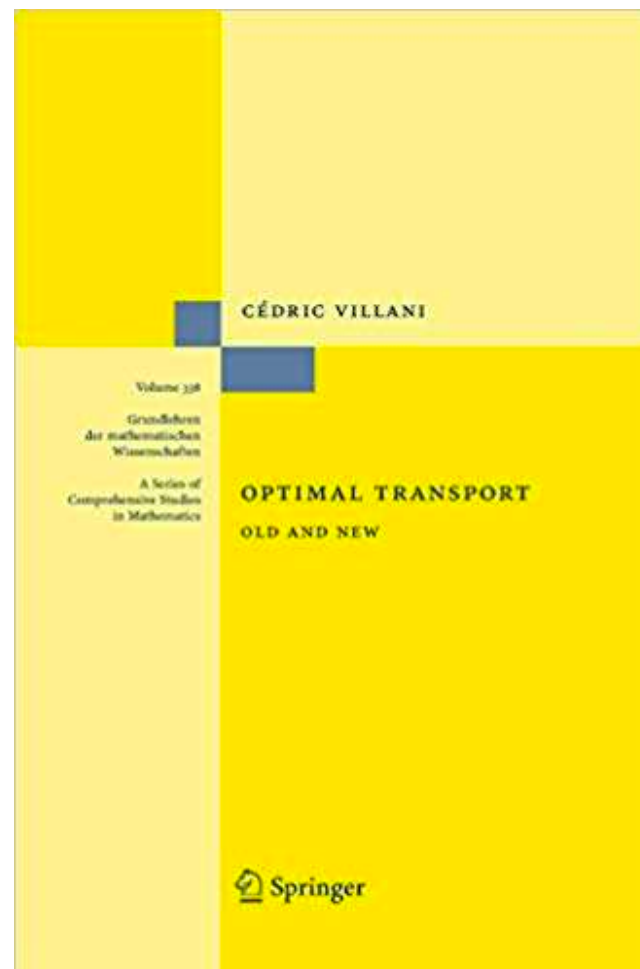
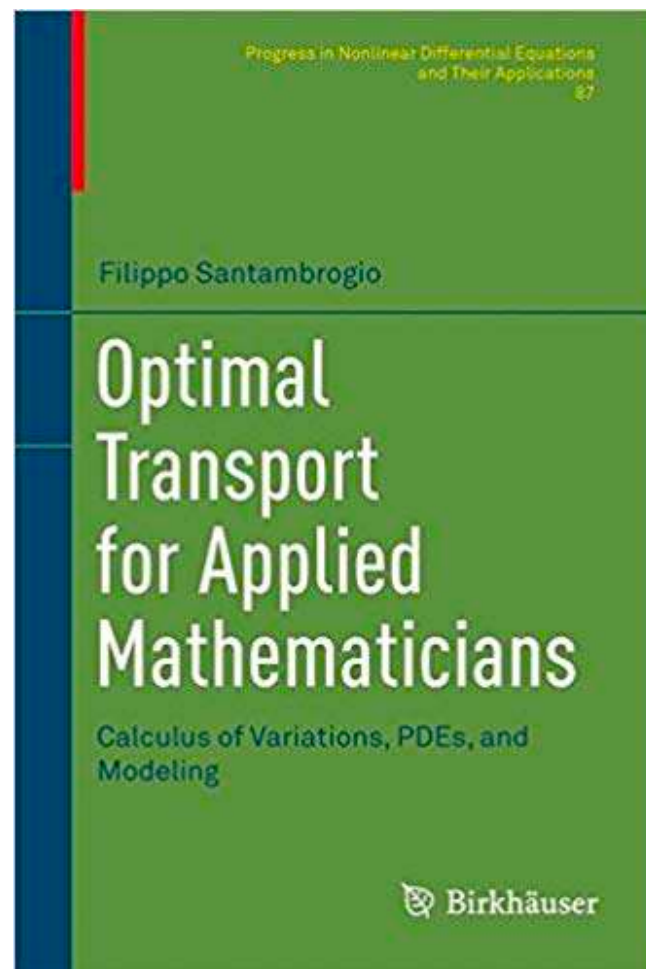
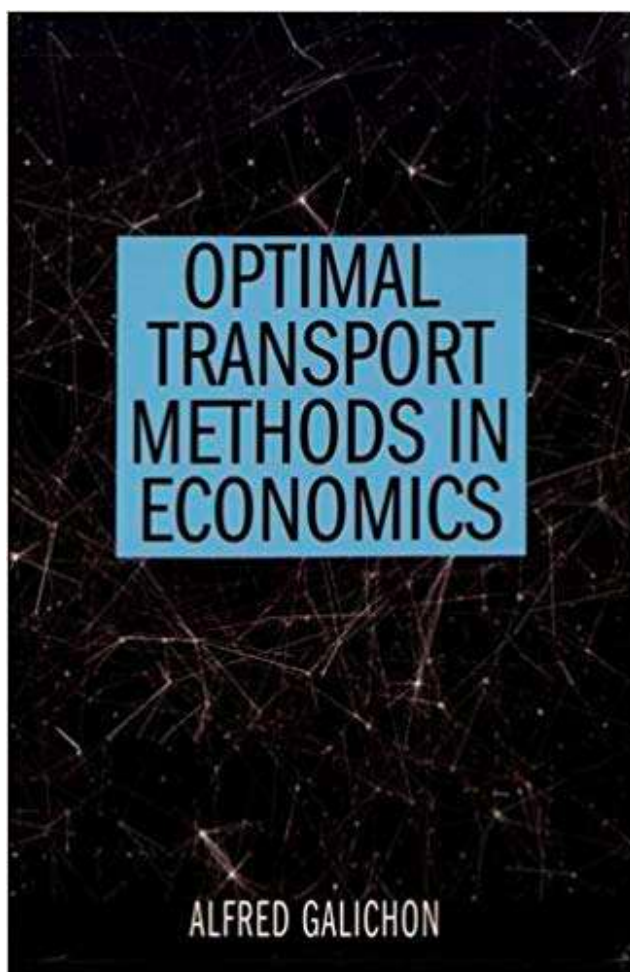
Nicolas Courty, Rémi Flamary, Devis Tuia, *Senior Member, IEEE*,
Alain Rakotomamonjy, *Member, IEEE*

In all these problems...

- An **intermediate** *matching/mapping* step is needed.
- That step is only **intermediate** in the sense that other *advanced learning procedures* build on top of it.
- Wishlist for a *matching* framework:
 - *guided by rich theory*, to inform new algorithms,
 - *versatile*, to accommodate several settings,
 - *scalable*, to handle large *dim/n* datasets,
 - *differentiable*, for end-to-end learning,
 - *robust*, to perform well in most settings.

That theory is Optimal Transport

Matchings happen everywhere, so many people have (re)discovered it: economics, *applied* maths, *pure* maths, graphics, and, increasingly ML people!



Why is it popular across so many fields?

It has a bit of everything! theory provides theorems; algorithms translate them in tangible results; applications guide new developments.



Monge



Kantorovich



Koopmans

Nobel'75



Dantzig



McCann



Otto



Brenier



Caffarelli



Gangbo



Ambrosio



Villani

Fields'10

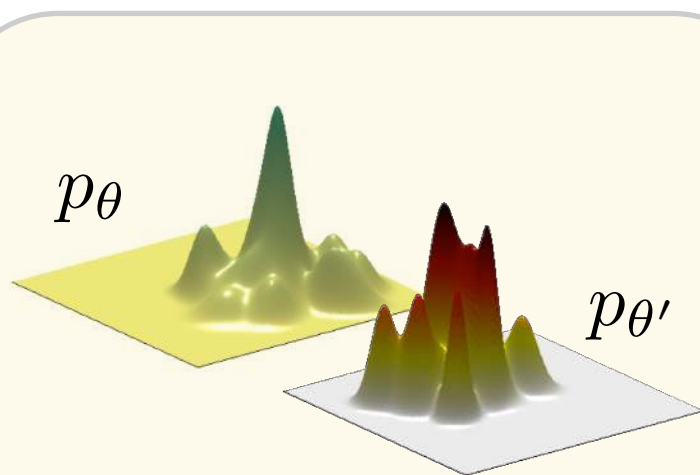


Figalli

Fields'18

Optimal Transport in Data Sciences

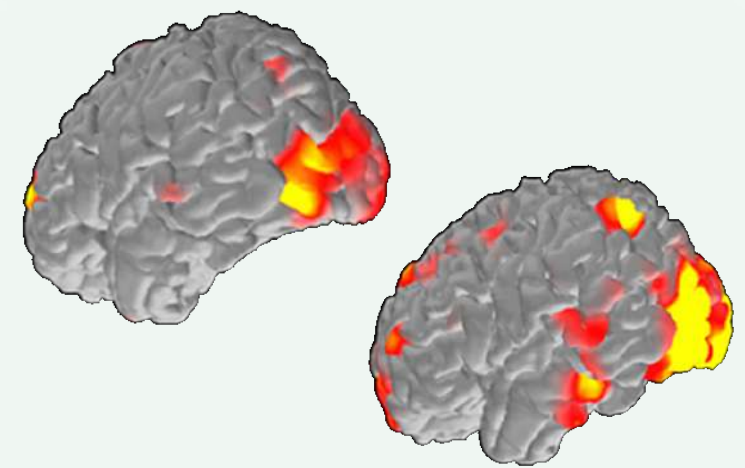
Advertised as the natural way to define geometry, through matching, for **probability measures**, when supported on a geometric space.



Statistical Models

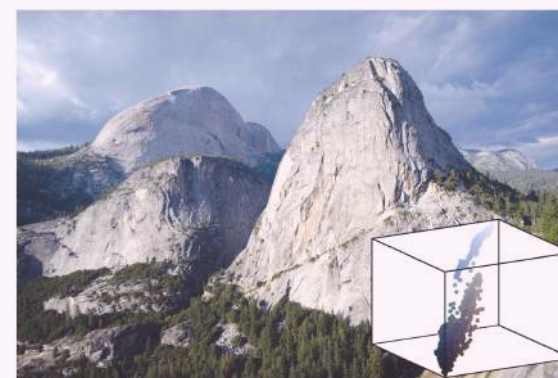
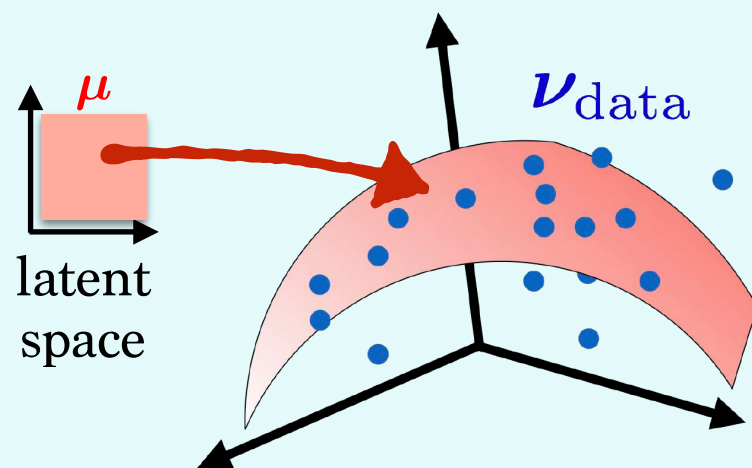


Bags of features



Brain Activation Maps

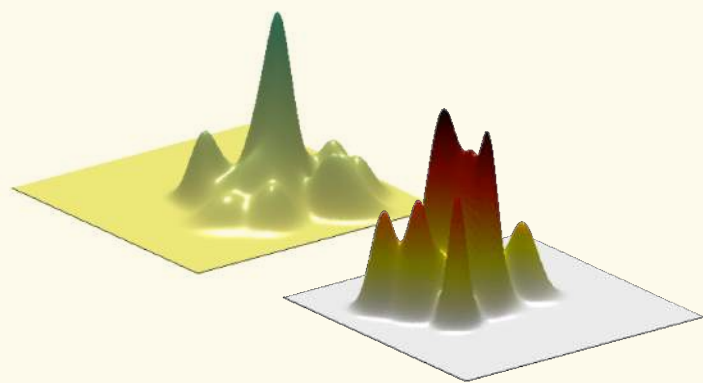
Generative Models vs. data



Color Histograms

Optimal Transport in Data Sciences

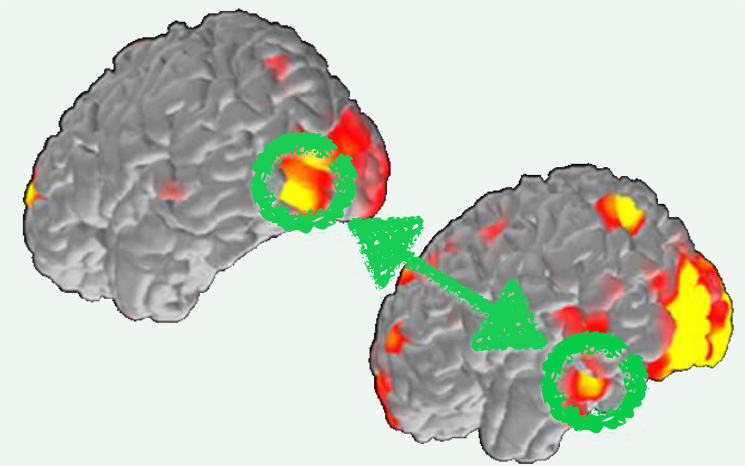
Advertised as the natural way to define geometry, through matching, for **probability measures**, when supported on a geometric space.



Statistical Models

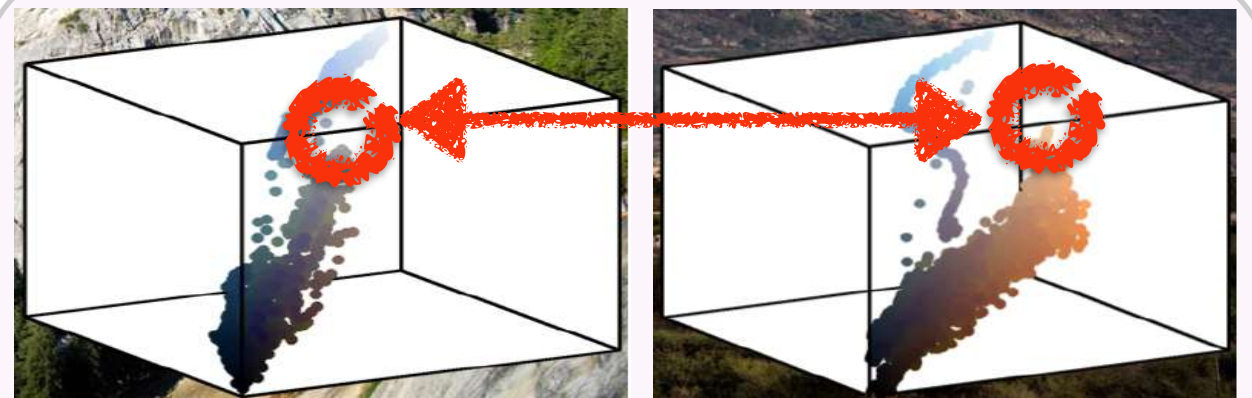
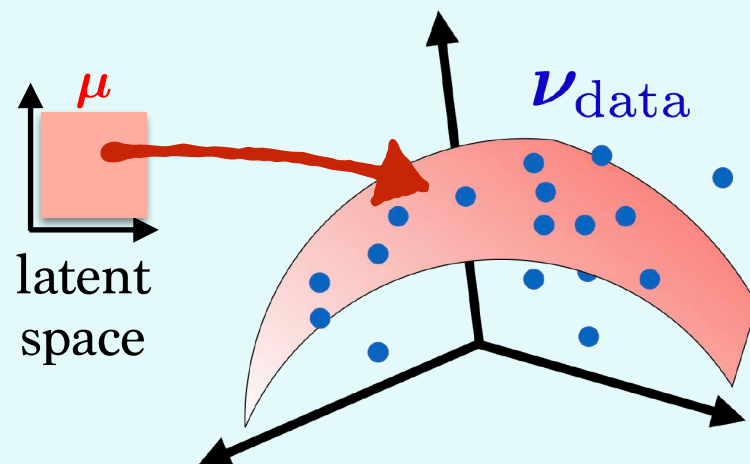


Bags of Features



Brain Activation Maps

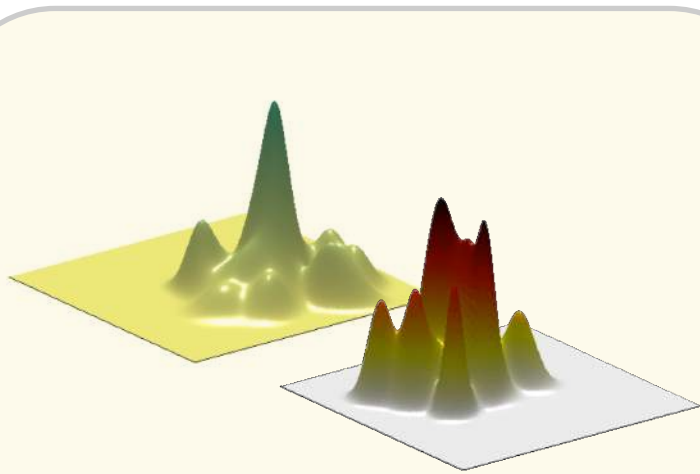
Generative Models vs. Data



Color Histograms

Optimal Transport in Data Sciences

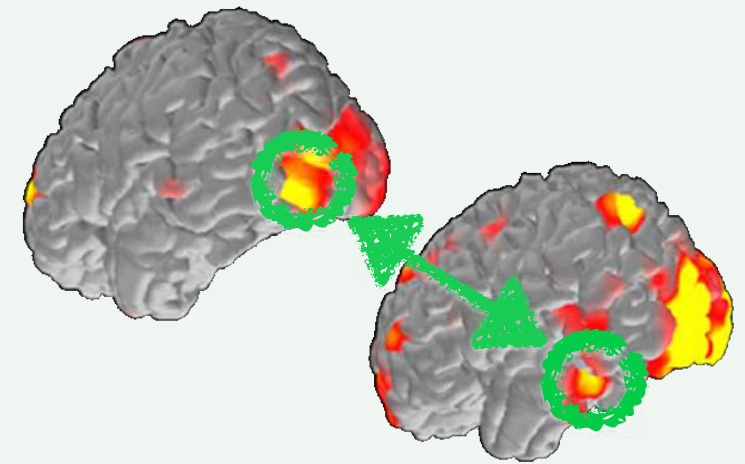
Advertised as the natural way to define geometry,
through matching, for **probability measures**,
when supported on a geometric space.



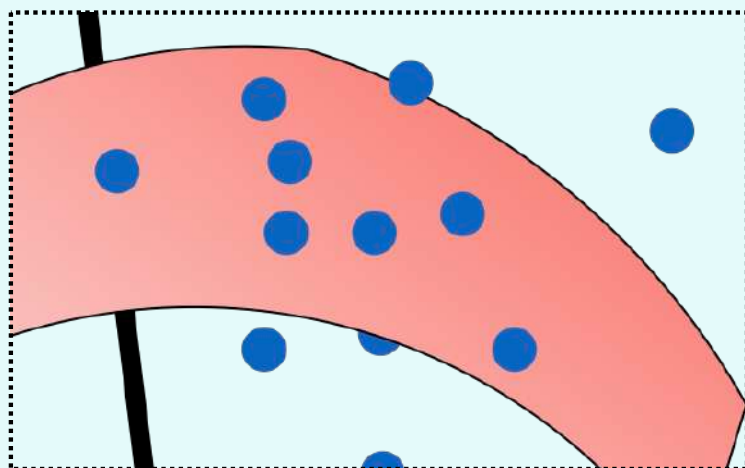
Statistical Models



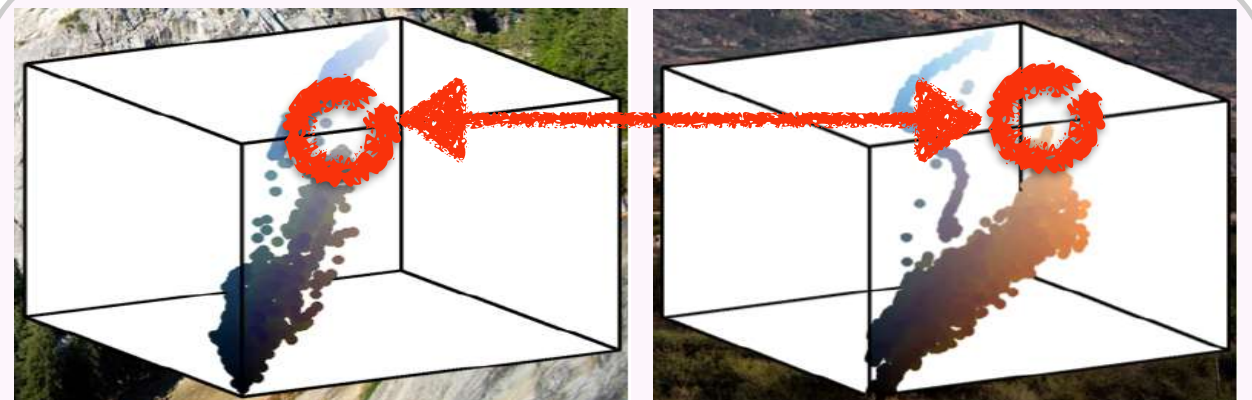
Bags of Features



Brain Activation Maps



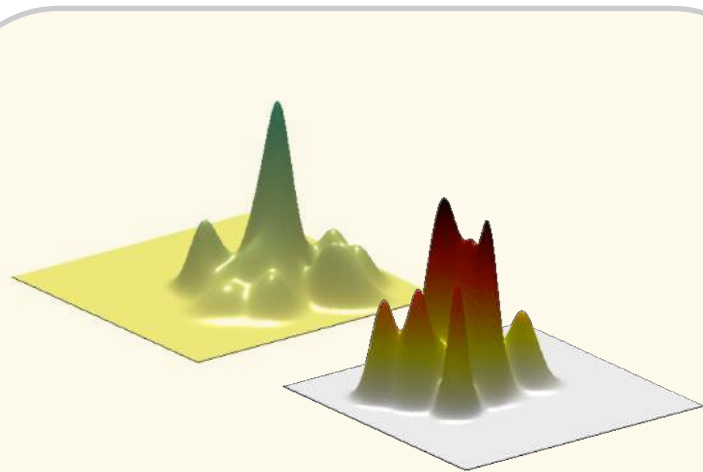
Generative Models vs. Data



Color Histograms

Optimal Transport in Data Sciences

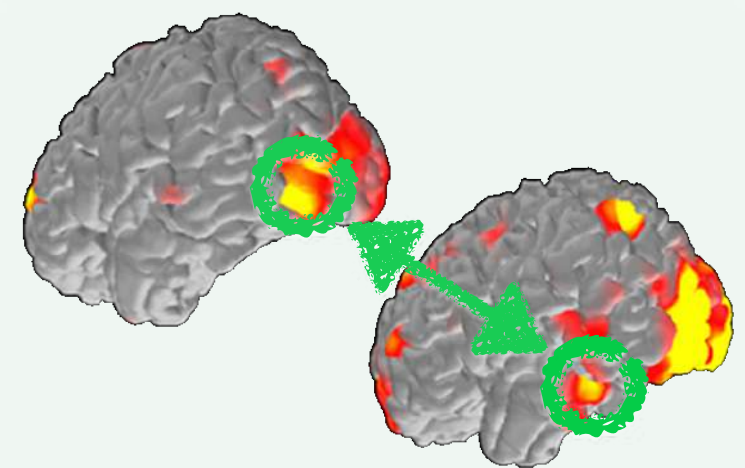
Advertised as the natural way to define geometry, through matching, for **probability measures**, when supported on a geometric space.



Statistical Models

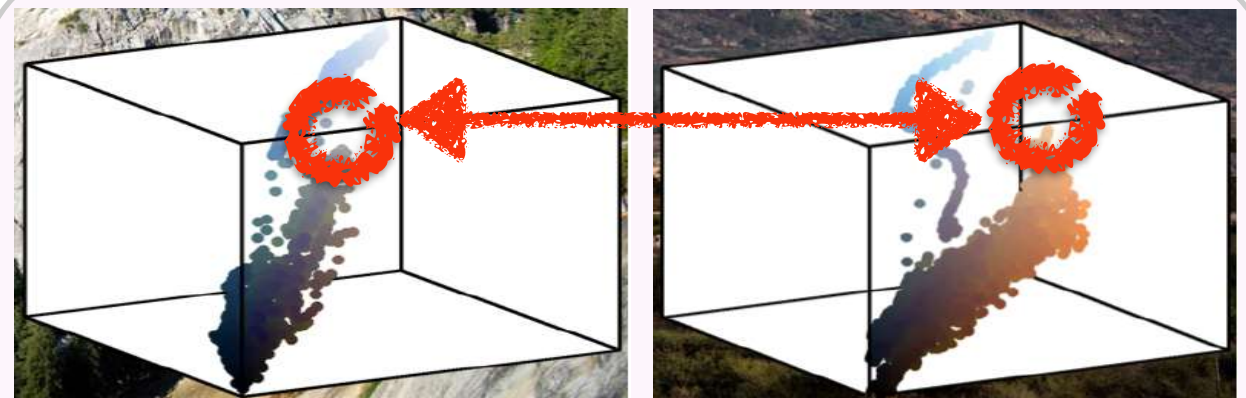
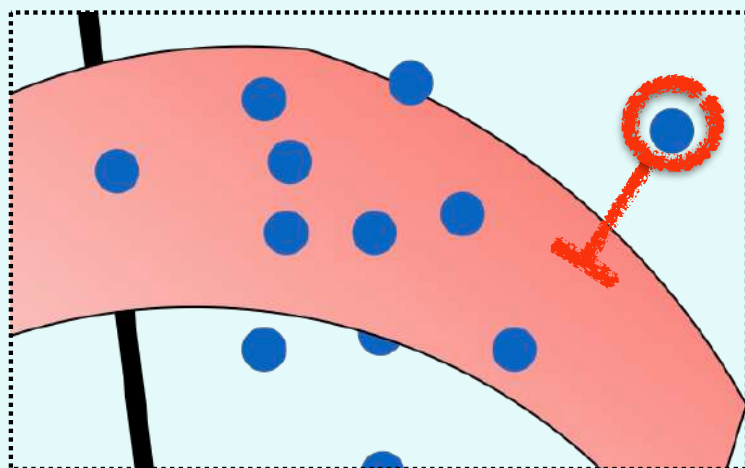


Bags of Features



Brain Activation Maps

Generative Models vs. Data



Color Histograms

Short Course Outline

1. Introduction to optimal transport
2. Computing OT exactly
3. Computing OT for data science

Introduction to OT

- Two examples: moving earth & soldiers
- Monge problem, Kantorovich problem
- OT as geometry, OT as a loss function

Origins: Monge Problem (1781)



Gaspard Monge (1746 - 1818)

Origins: Monge Problem (1781)



Paris ▾



Join

Search

Travel feed: Paris Hotels **Things to do** Restaurants Flights Holiday Homes Shopping Car Hire ...

Europe > France > Ile-de-France > Paris > Things to do in Paris > Place Monge

Place Monge, Paris: Address, Place Monge Reviews: 4/5

Place Monge

84 Reviews

#323 of 2 272 things to do in Paris

Shopping, Flea & Street Markets

Pl. Monge, Paris, France

Save Share

Review Highlights

“Good local market.”

Place Monge market is one of our regular haunts when staying in Paris, it has a good range of... [read more](#)



Reviewed 26 September 2018

johngl8492UH , Middle Park via mobile

“One of the most amazing streets we have be...”

We loved place monge. All the shops slide their products out to the streets, it has a real Parisian... [read more](#)



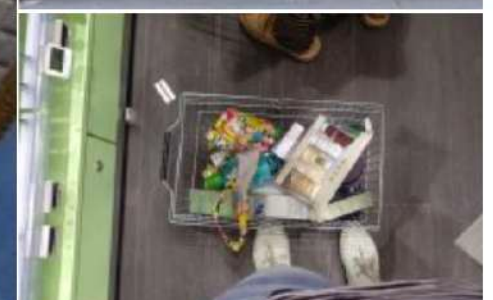
Reviewed 11 October 2018

Steve L , Pacific Coast Australia, Australia
via mobile

[Read all 84 reviews](#)



All photos (35)



Origins: Monge Problem (1781)



66 MÉMOIRES DE L'ACADÉMIE ROYALE

M É M O I R E

S U R L A

T H É O R I E D E S D É B L A I S

E T D E S R E M B L A I S.

Par M. M O N G E.

LORSQU'ON doit transporter des terres d'un lieu dans un autre, on a coutume de donner le nom de *Déblai* au volume des terres que l'on doit transporter, & le nom de *Remblai* à l'espace qu'elles doivent occuper après le transport.

Origins: Monge Problem (1781)



66 MÉMOIRES DE L'ACADÉMIE ROYALE

M É M O I R E

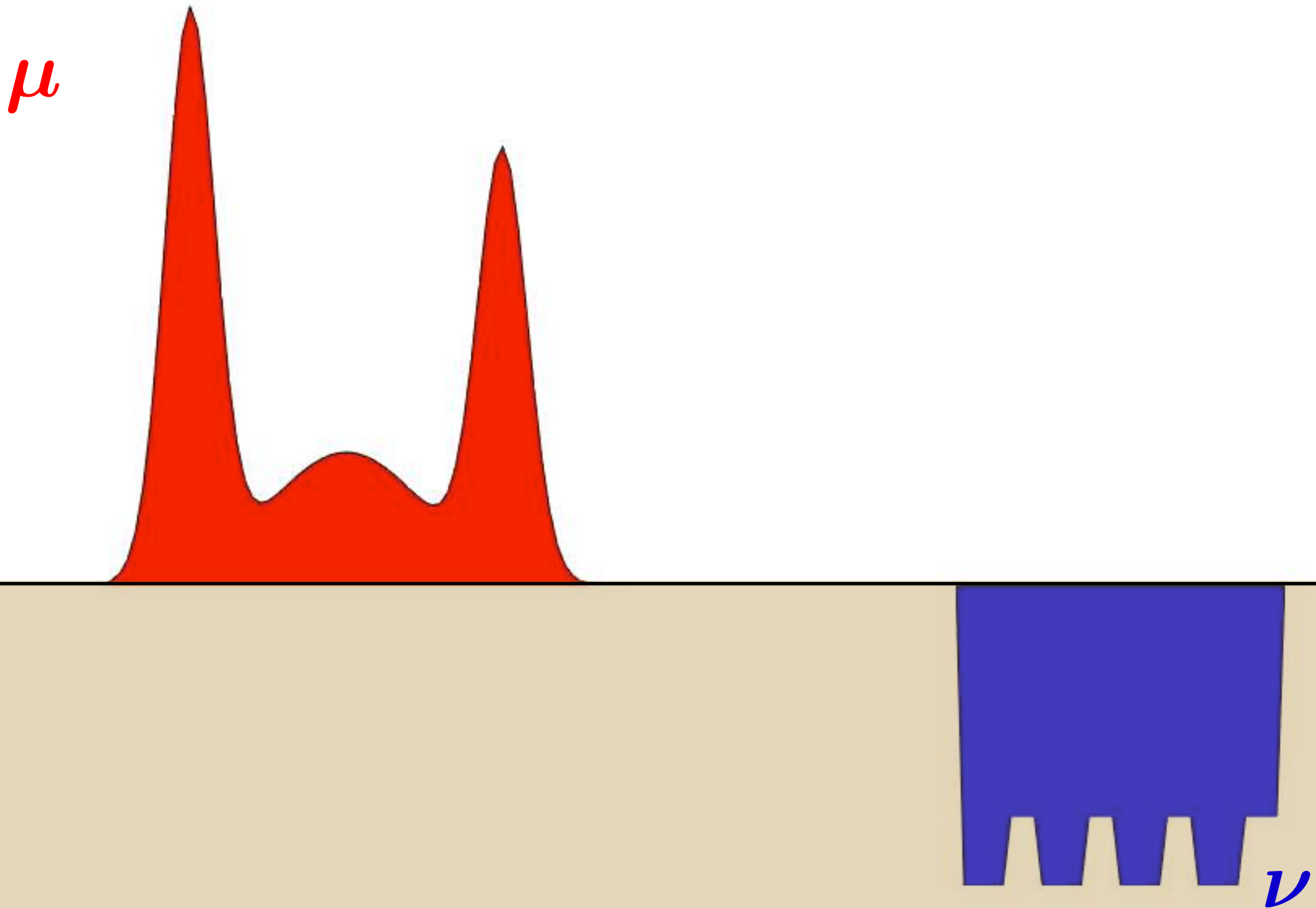
S U R L A

T H É O R I E D E S D É B L A I S

*When one has to bring earth
from one place to another...*

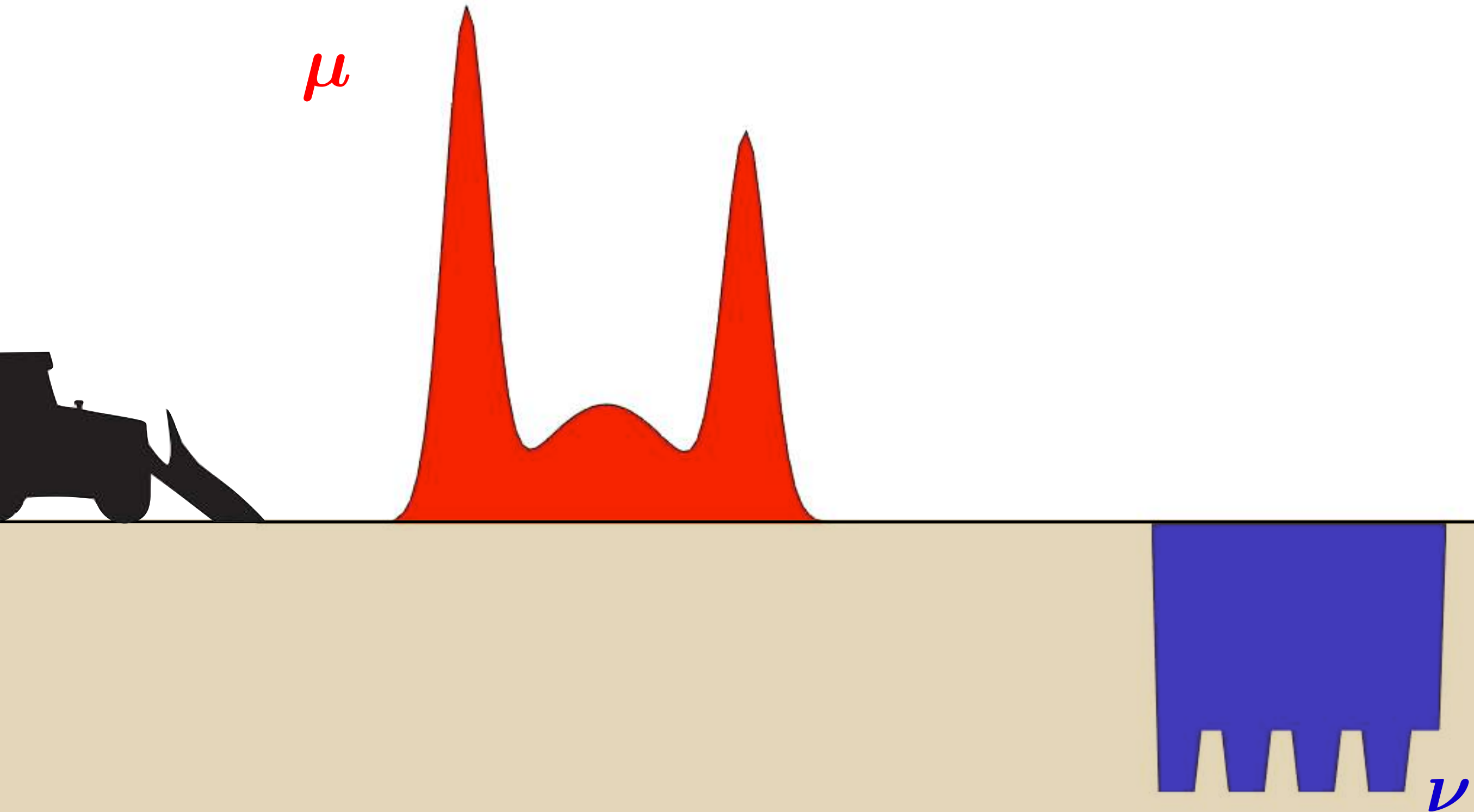
LORSQU'ON doit transporter des terres d'un lieu dans un autre, on a coutume de donner le nom de *Déblai* au volume des terres que l'on doit transporter, & le nom de *Remblai* à l'espace qu'elles doivent occuper après le transport.

Origins: Monge Problem



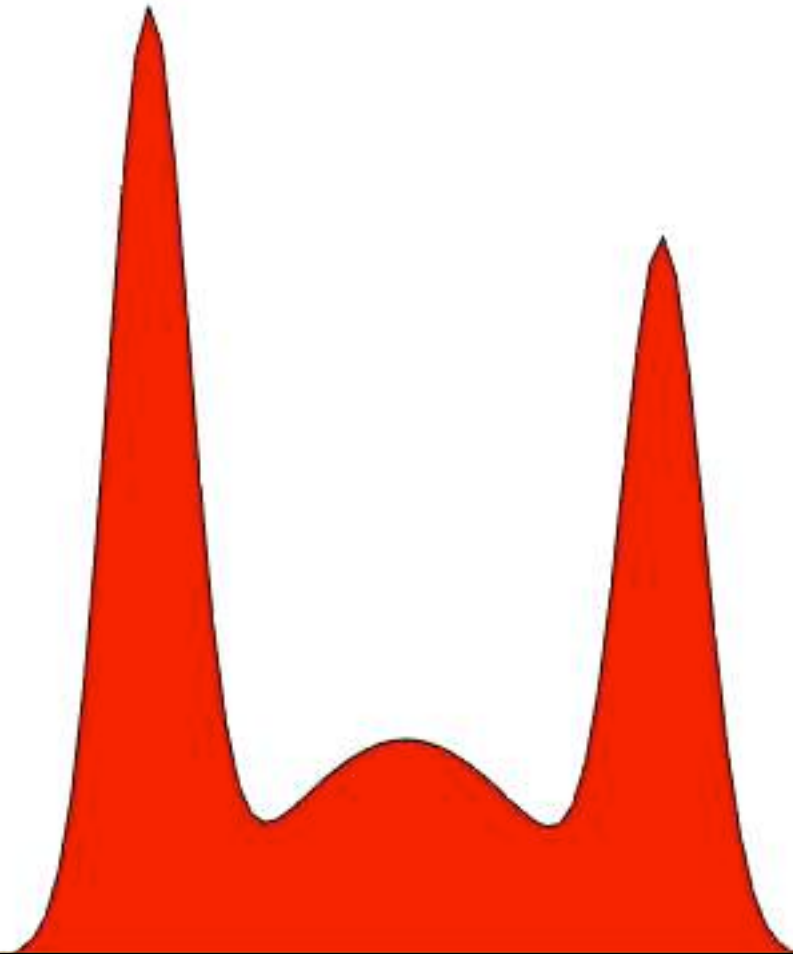
Origins: Monge Problem

In the 21st Century...



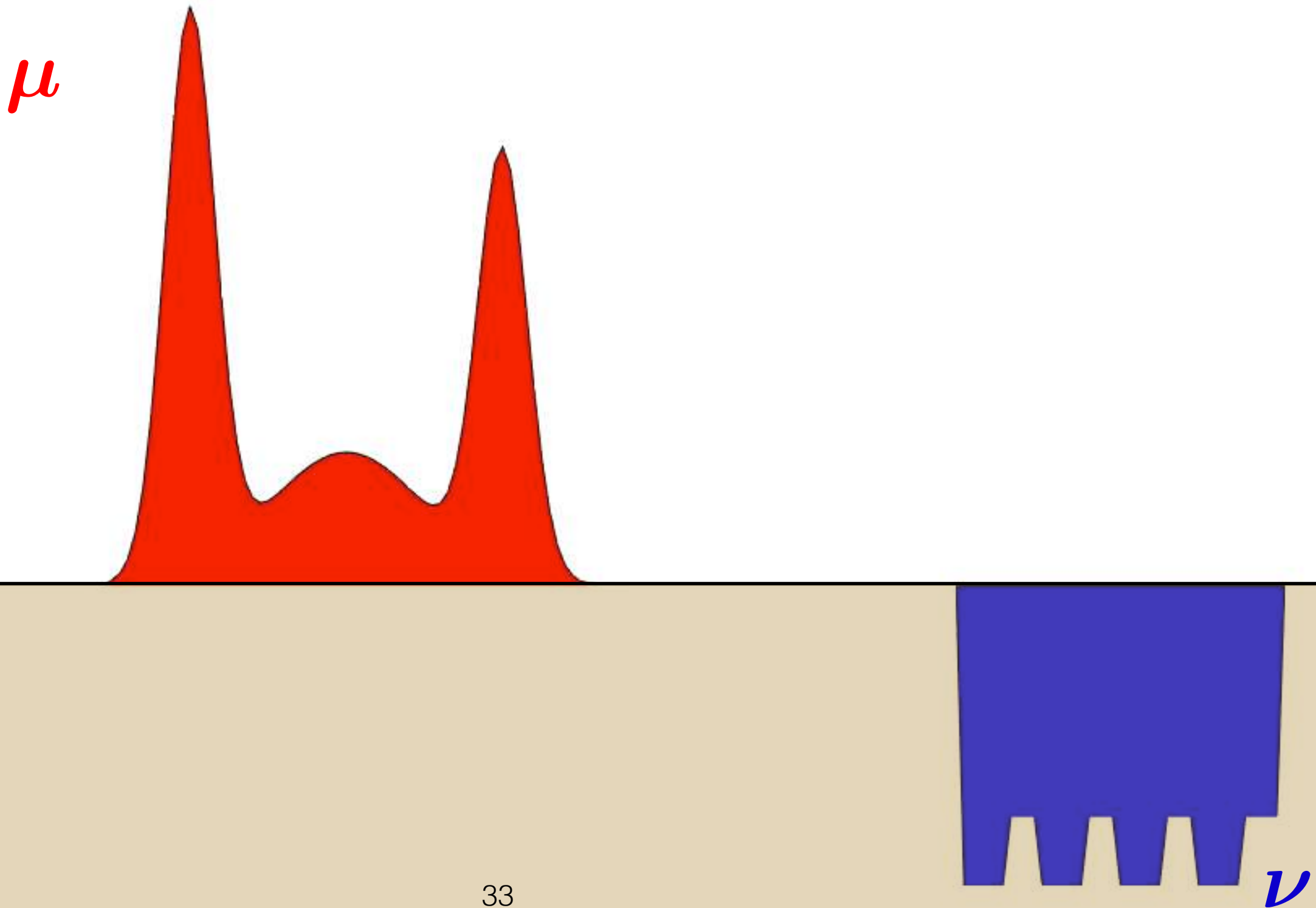
Origins: Monge Problem

In the 21st Century...



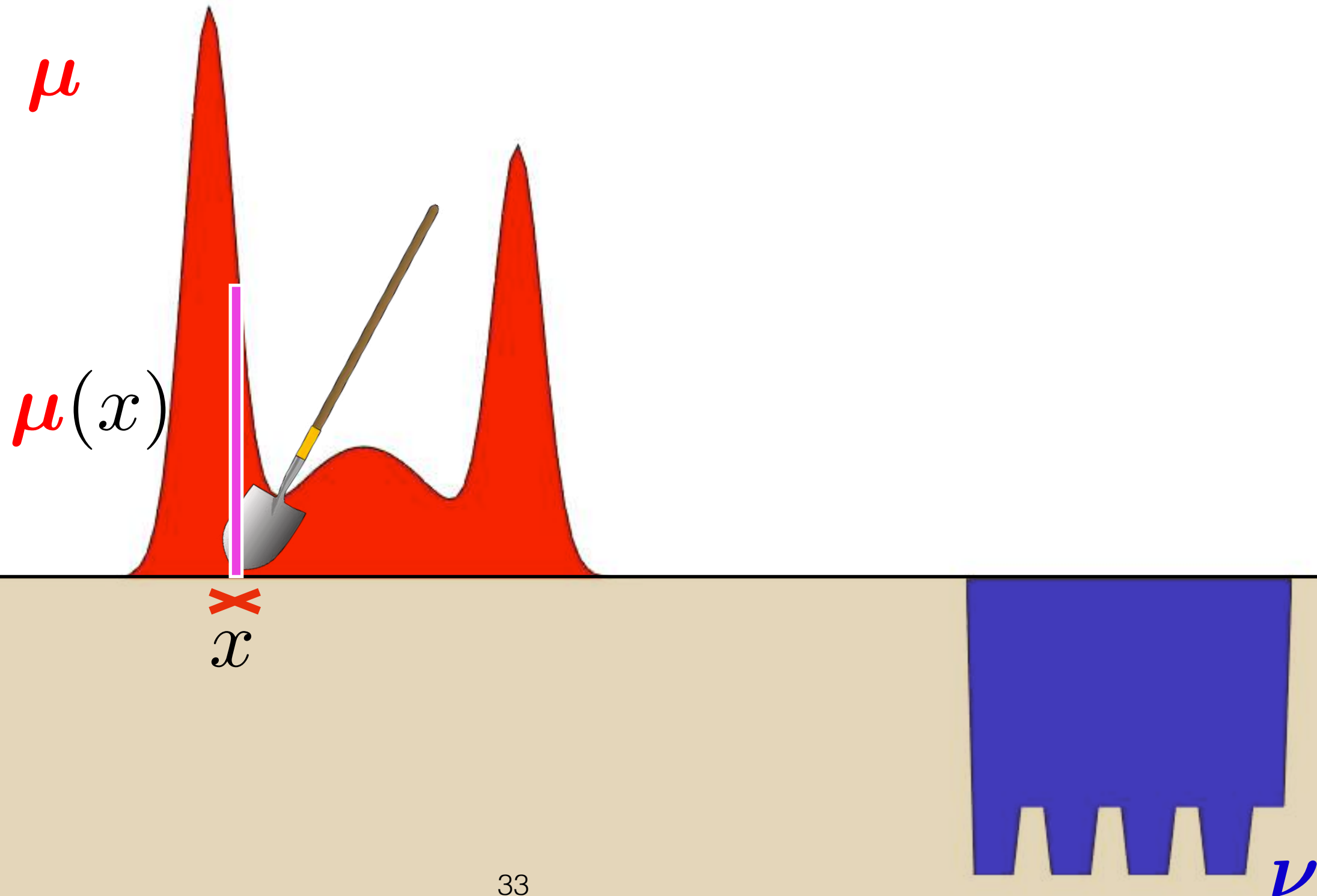
Origins: Monge's Problem

In 1781 however...



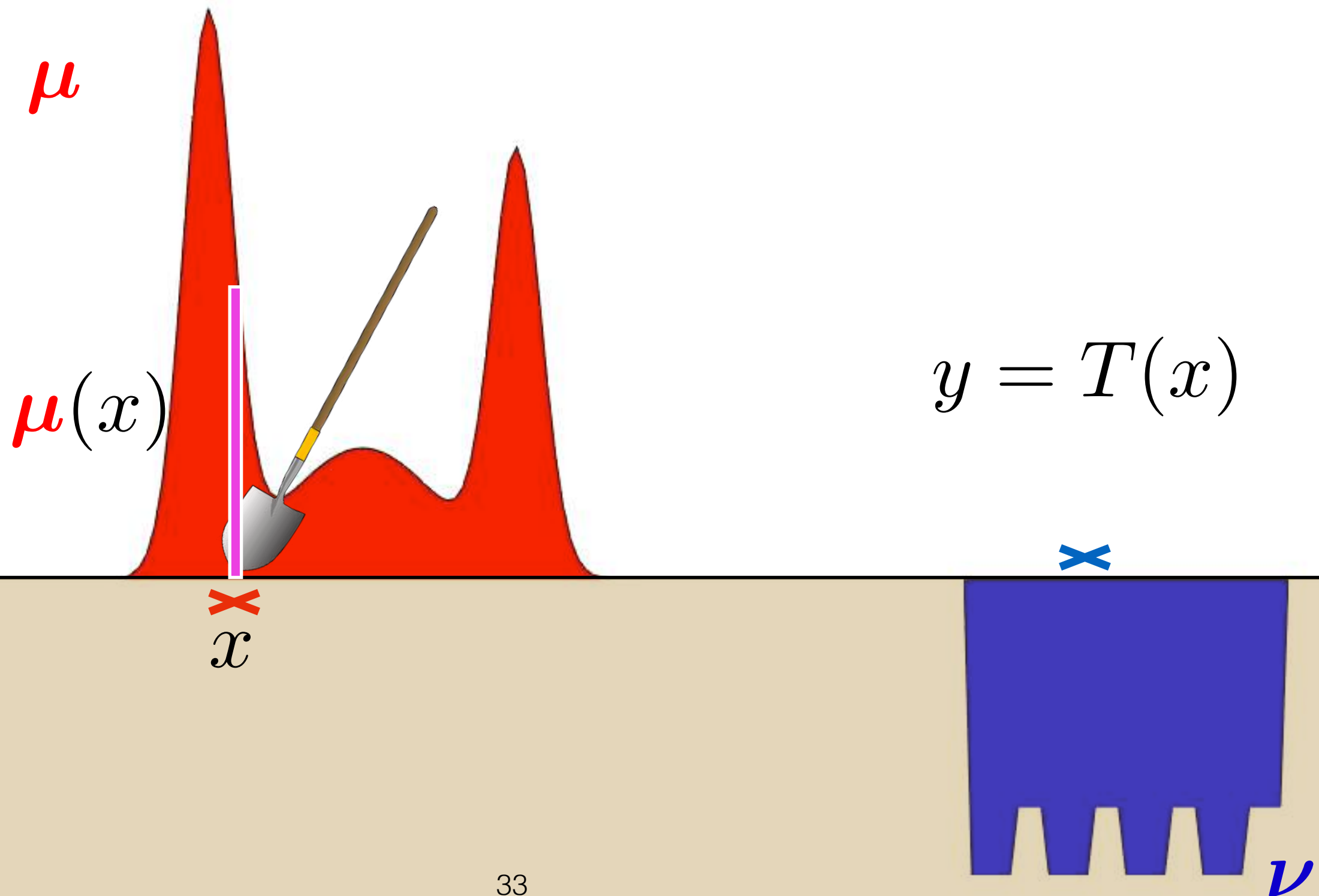
Origins: Monge's Problem

In 1781 however...



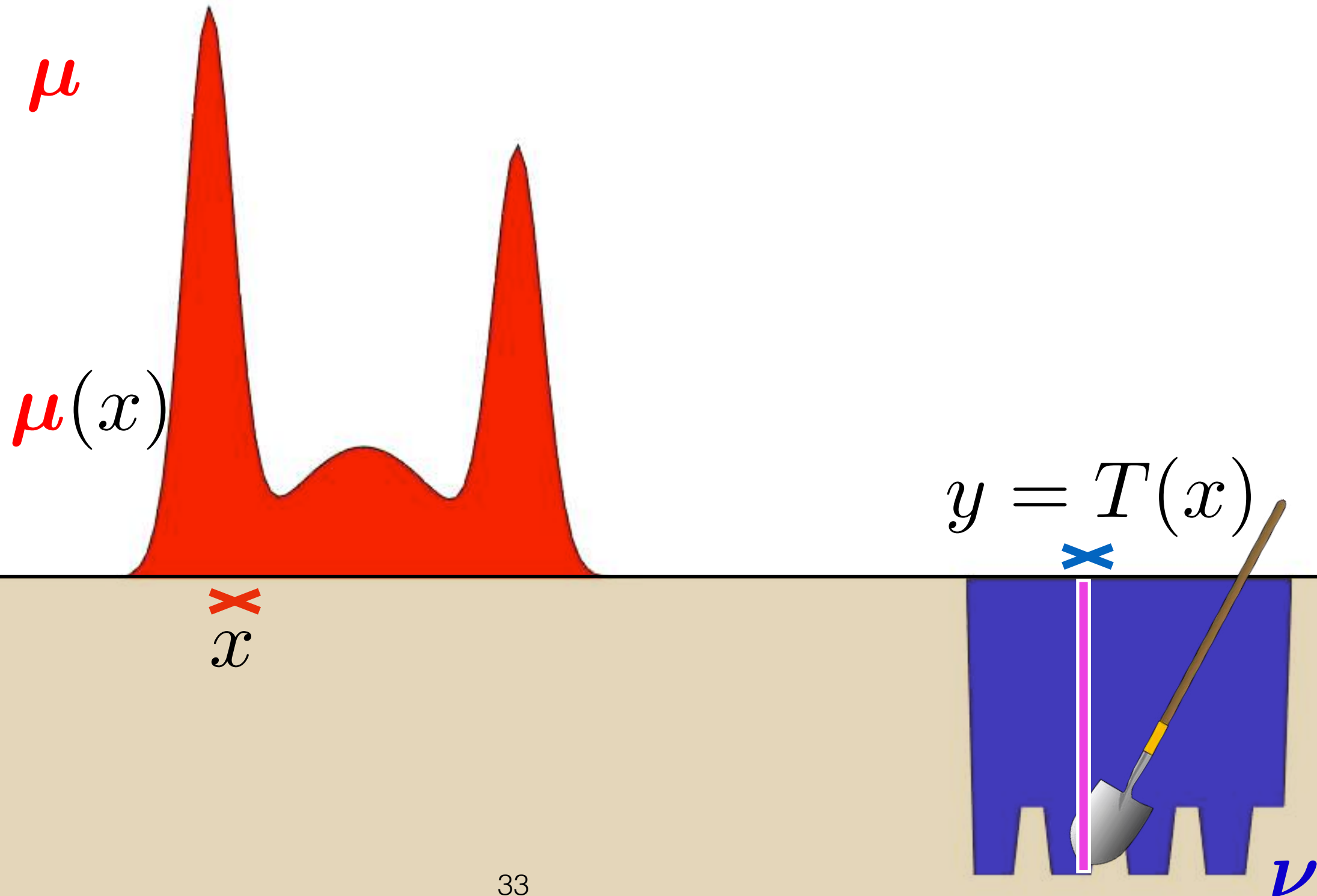
Origins: Monge's Problem

In 1781 however...



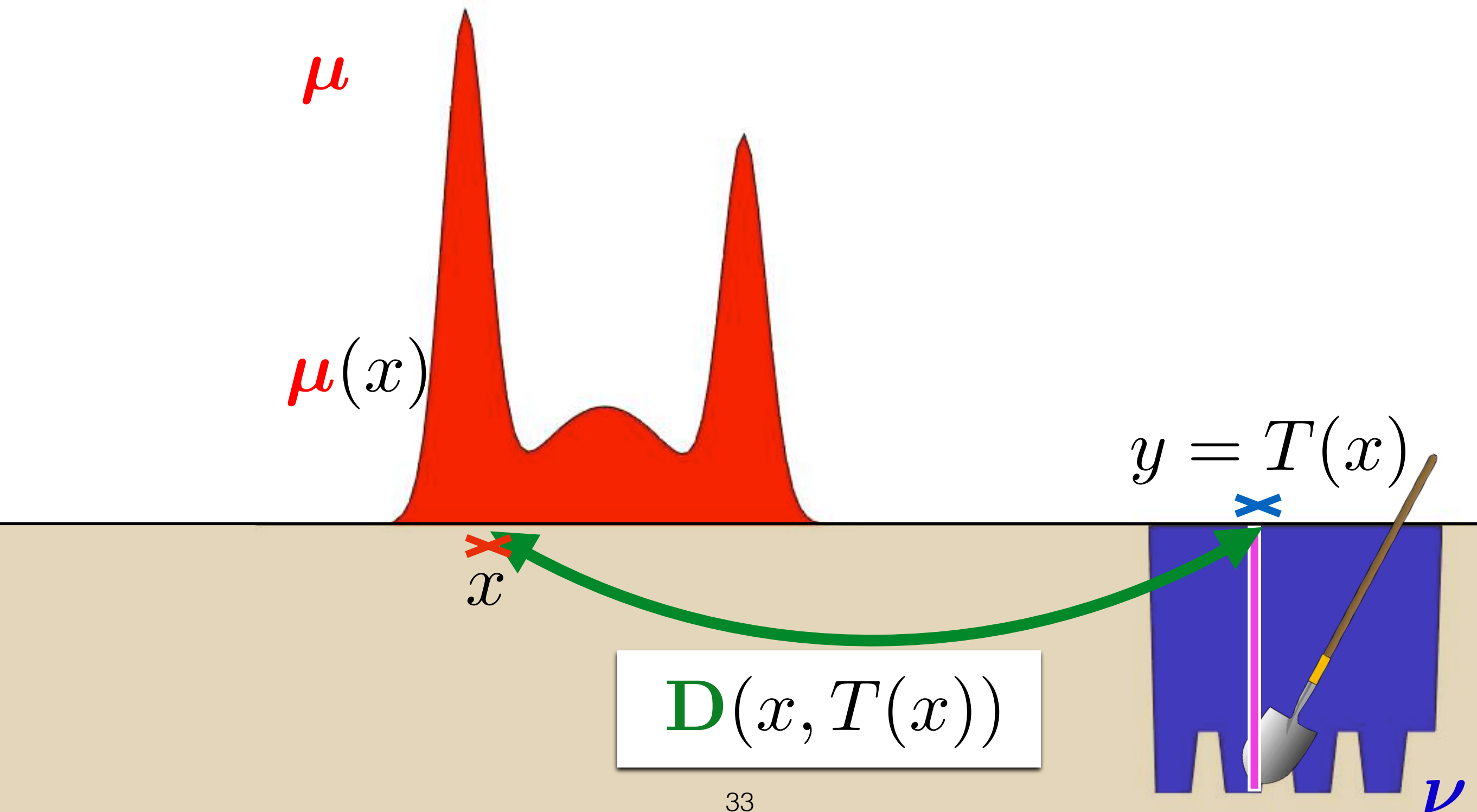
Origins: Monge's Problem

In 1781 however...



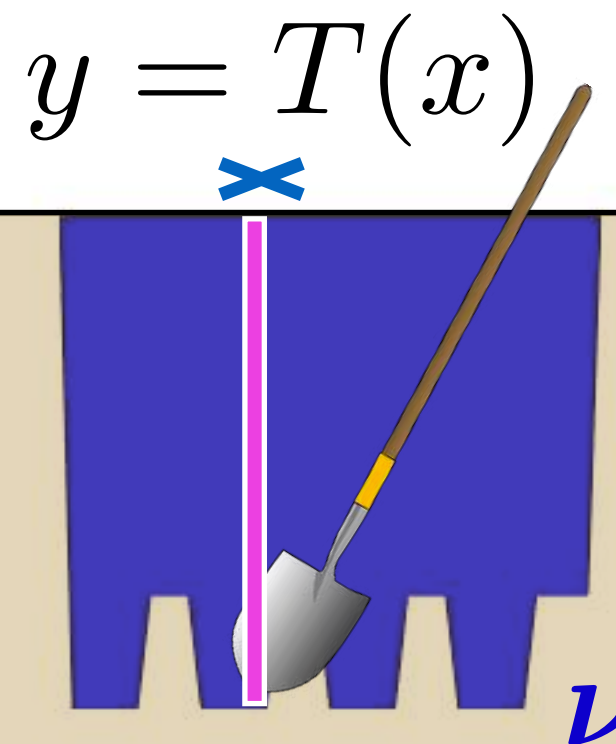
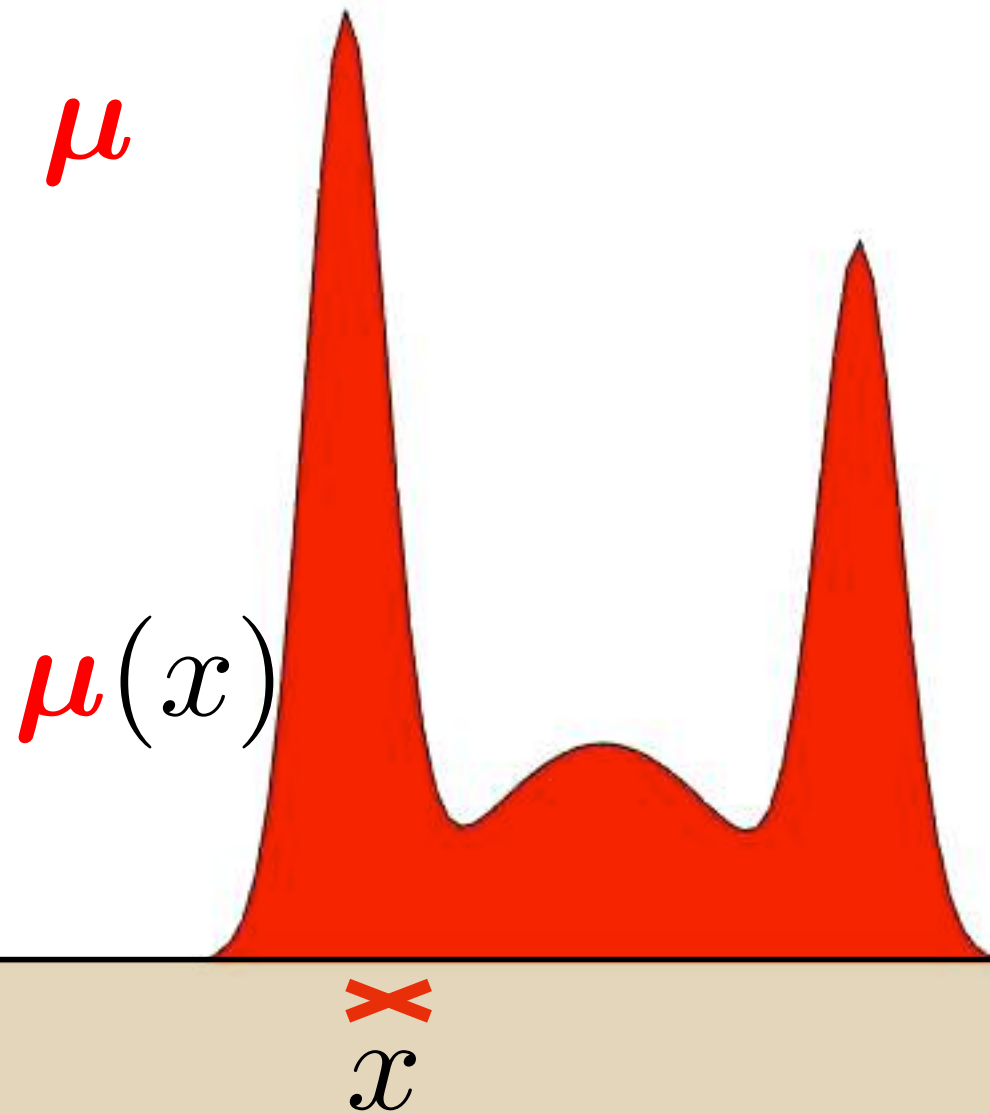
Origins: Monge's Problem

In 1781 however...



Origins: Monge's Problem

In 1781 however...

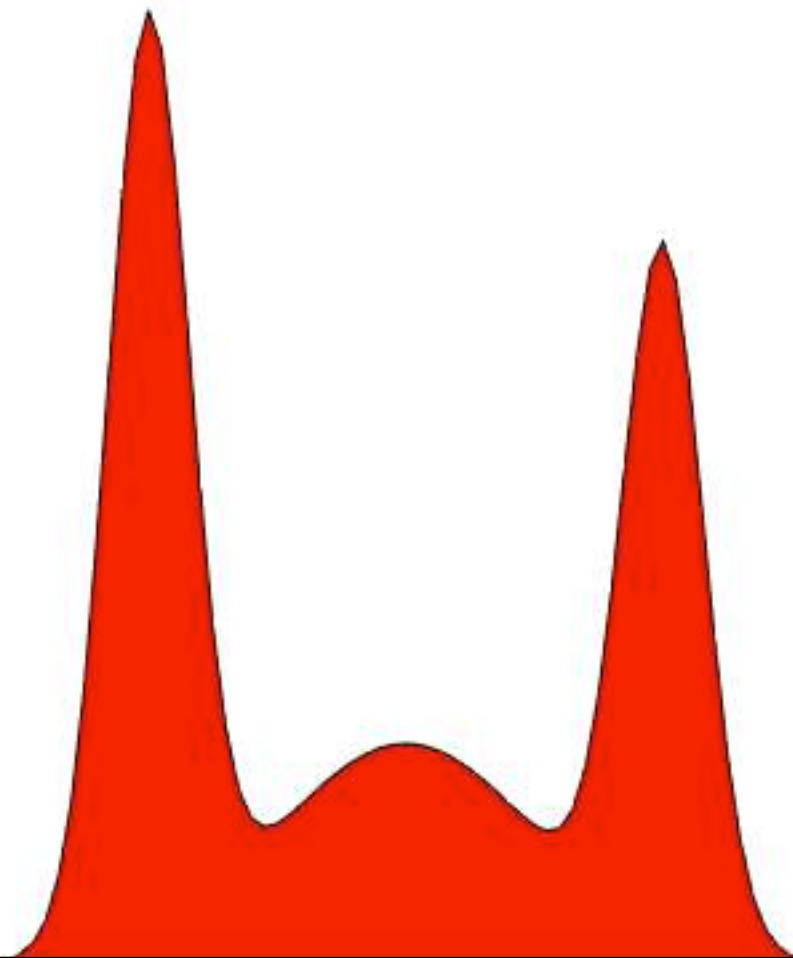


work: $\mu(x) D(x, T(x))$

Origins: Monge's Problem

T must map red to blue.

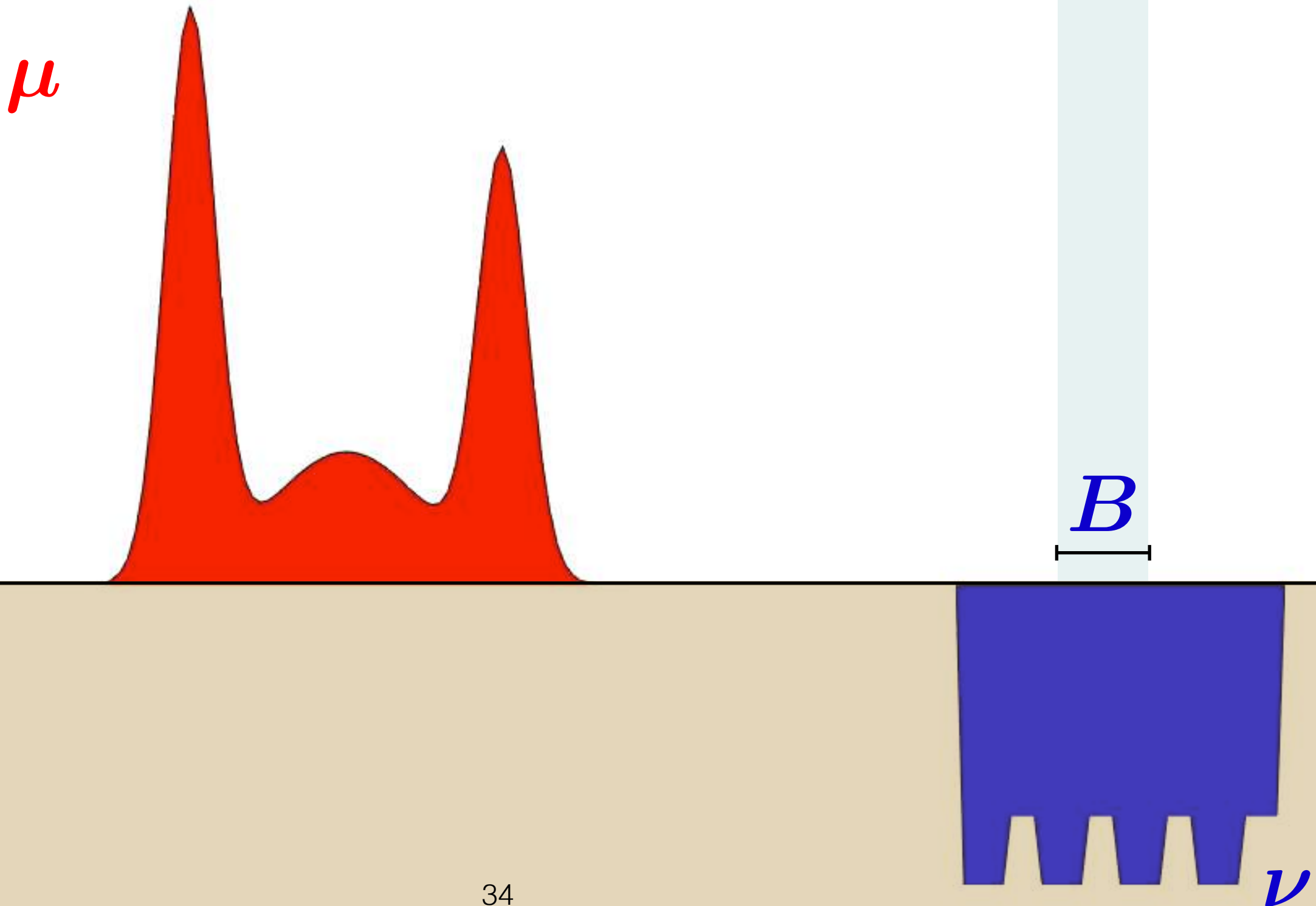
μ



ν

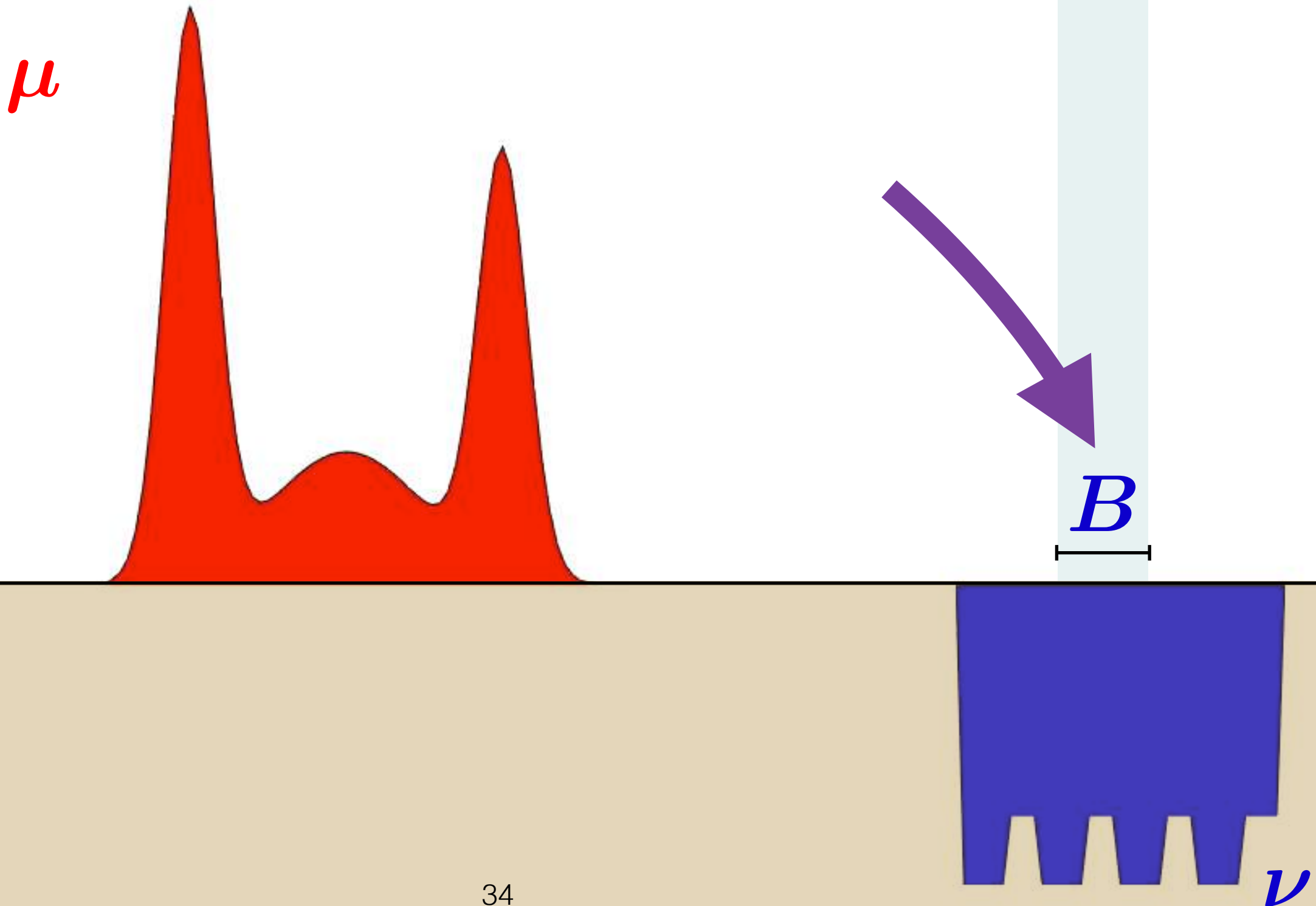
Origins: Monge's Problem

T must map red to blue.



Origins: Monge's Problem

T must map red to blue.



Origins: Monge's Problem

T must map red to blue.

μ

$$T^{-1}(B) = \{x | T(x) \in B\}$$

B

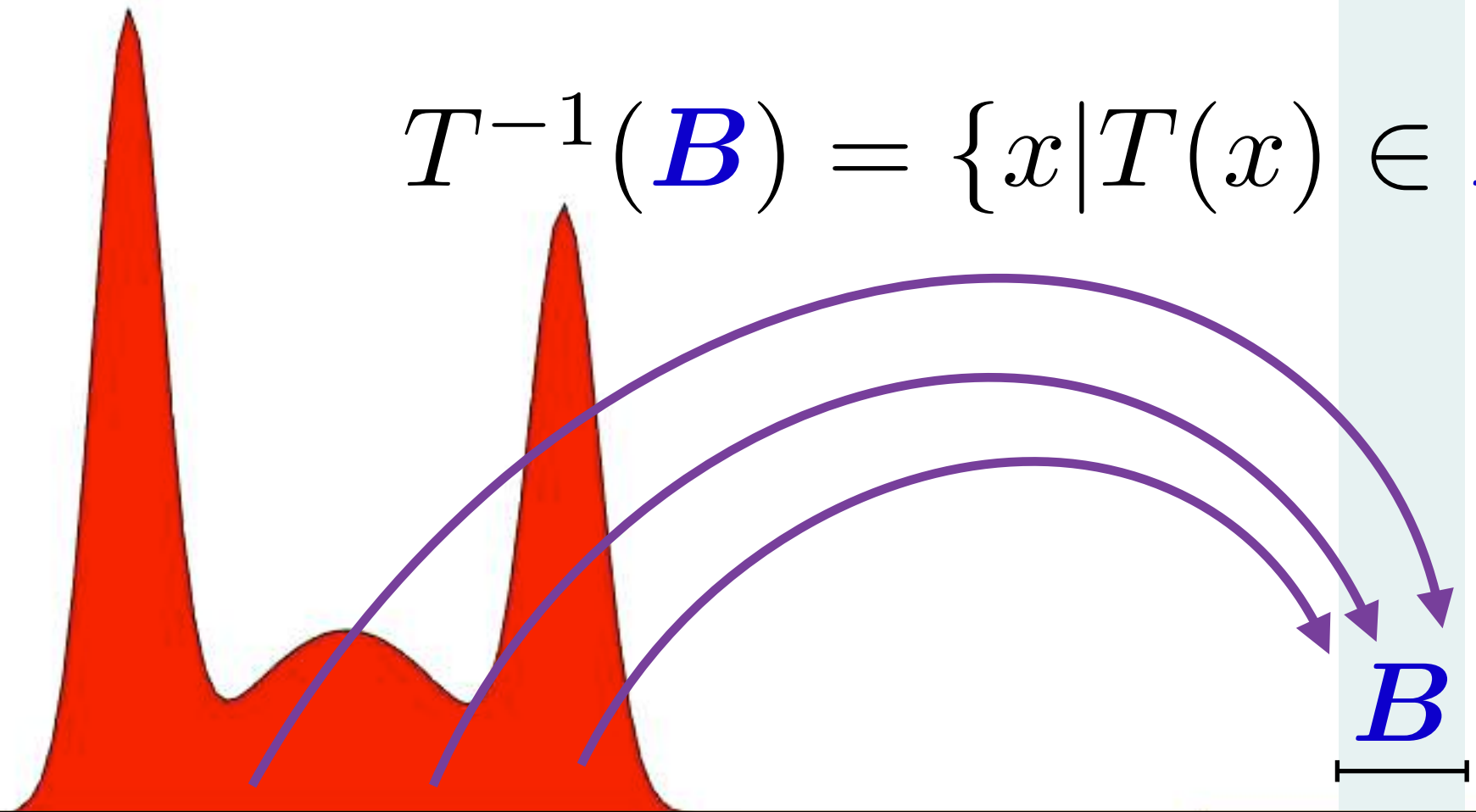
ν

Origins: Monge's Problem

T must map red to blue.

μ

$$T^{-1}(\mathbf{B}) = \{x | T(x) \in \mathbf{B}\}$$

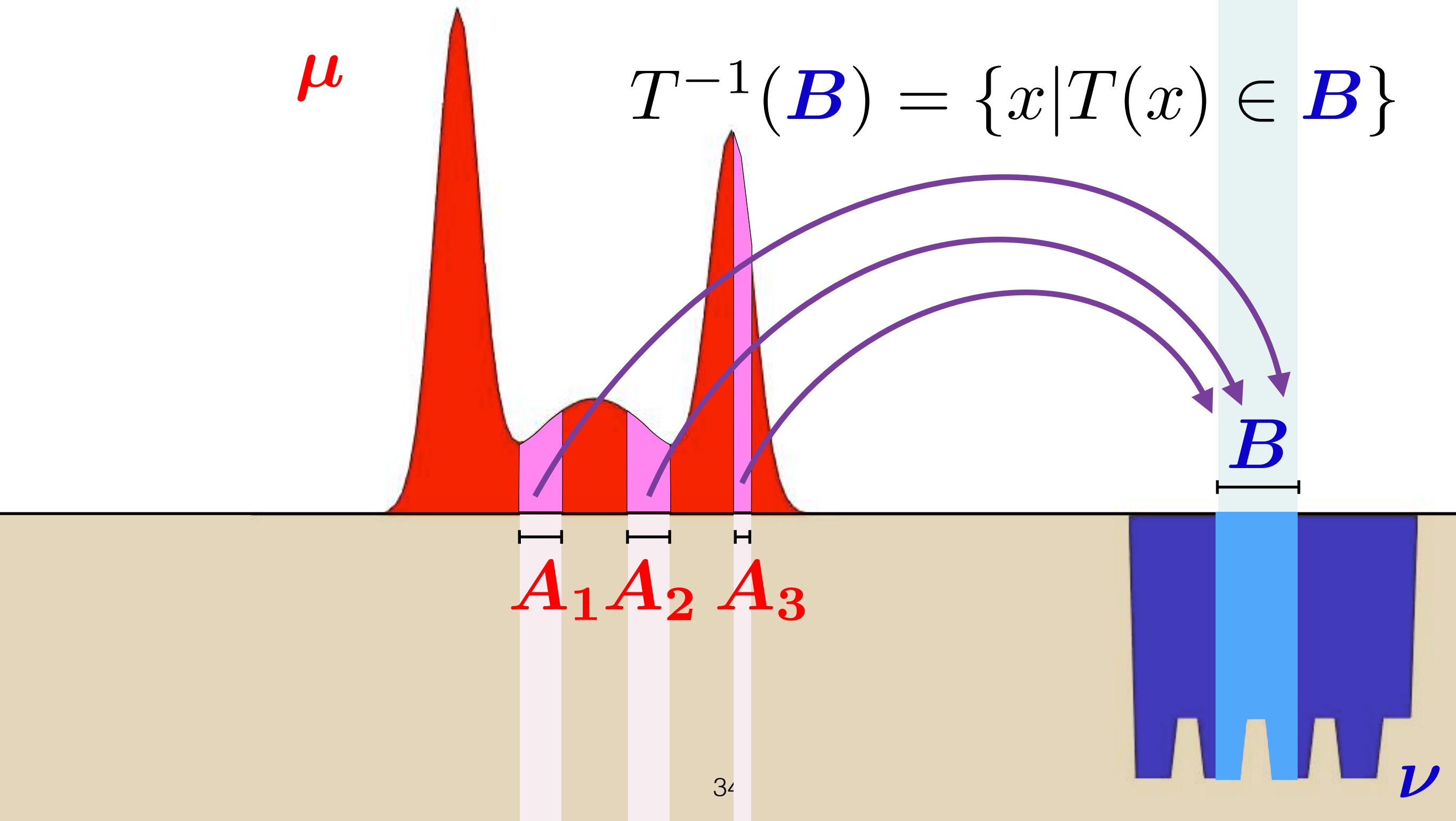


Origins: Monge's Problem

T must map red to blue.

μ

$$T^{-1}(\mathbf{B}) = \{x | T(x) \in \mathbf{B}\}$$

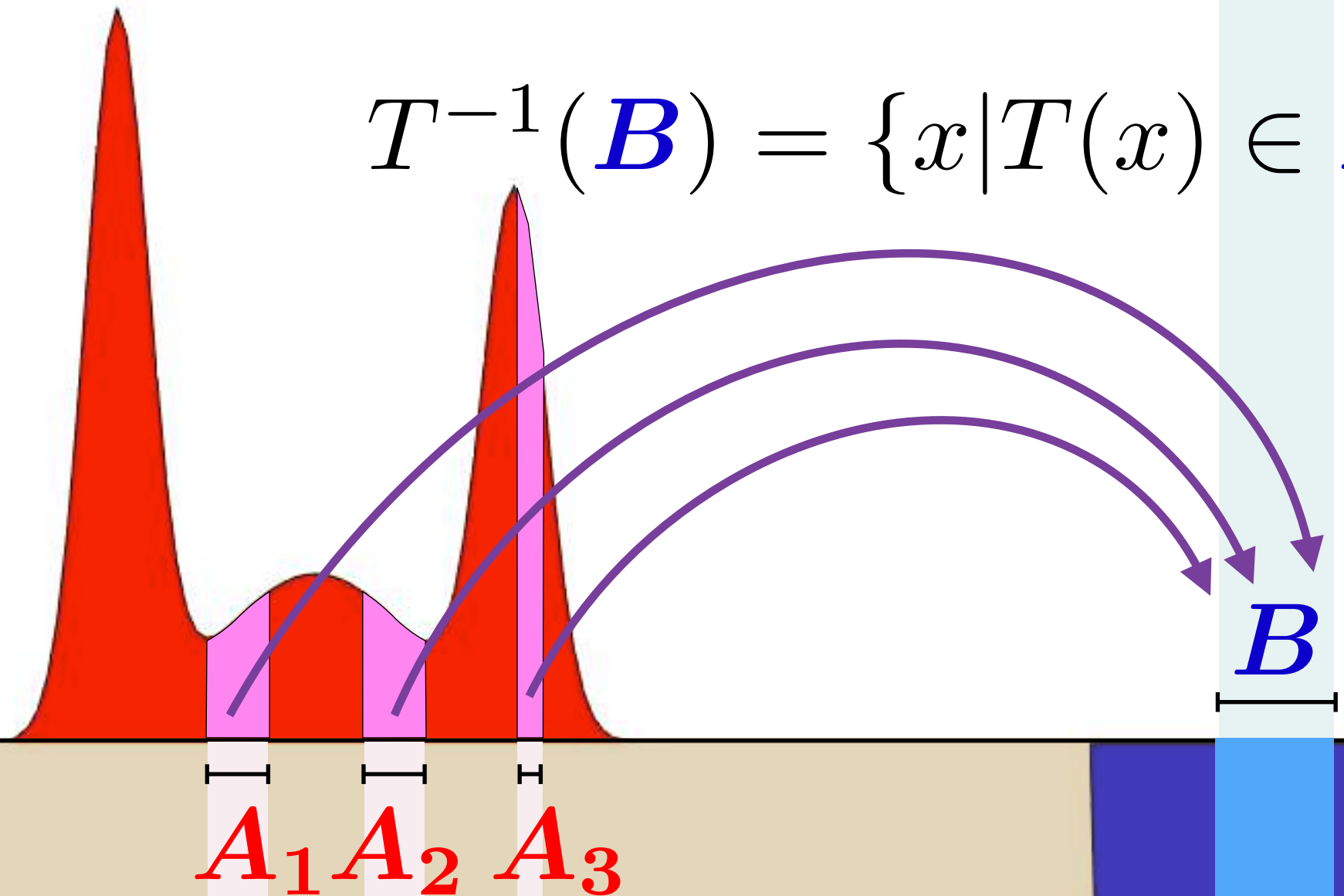


Origins: Monge's Problem

T must map red to blue.

μ

$$T^{-1}(\mathbf{B}) = \{x | T(x) \in \mathbf{B}\}$$



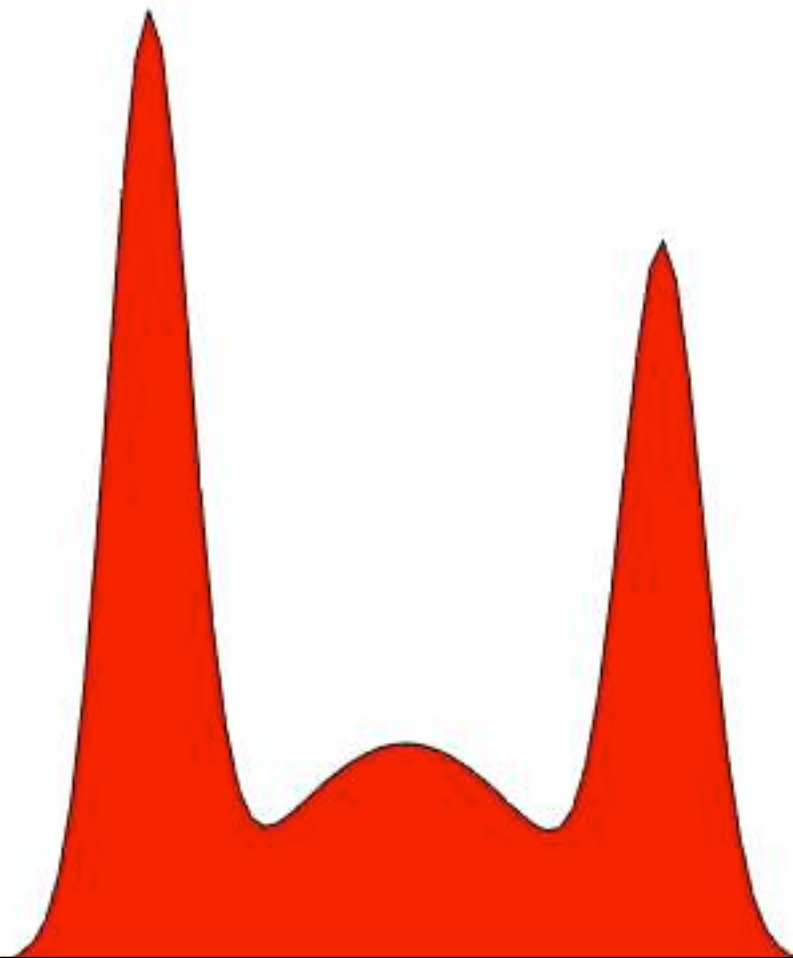
$$\mu(A_1) + \mu(A_2) + \mu(A_3) = \nu(B)$$

ν

Origins: Monge's Problem

T must map red to blue.

μ



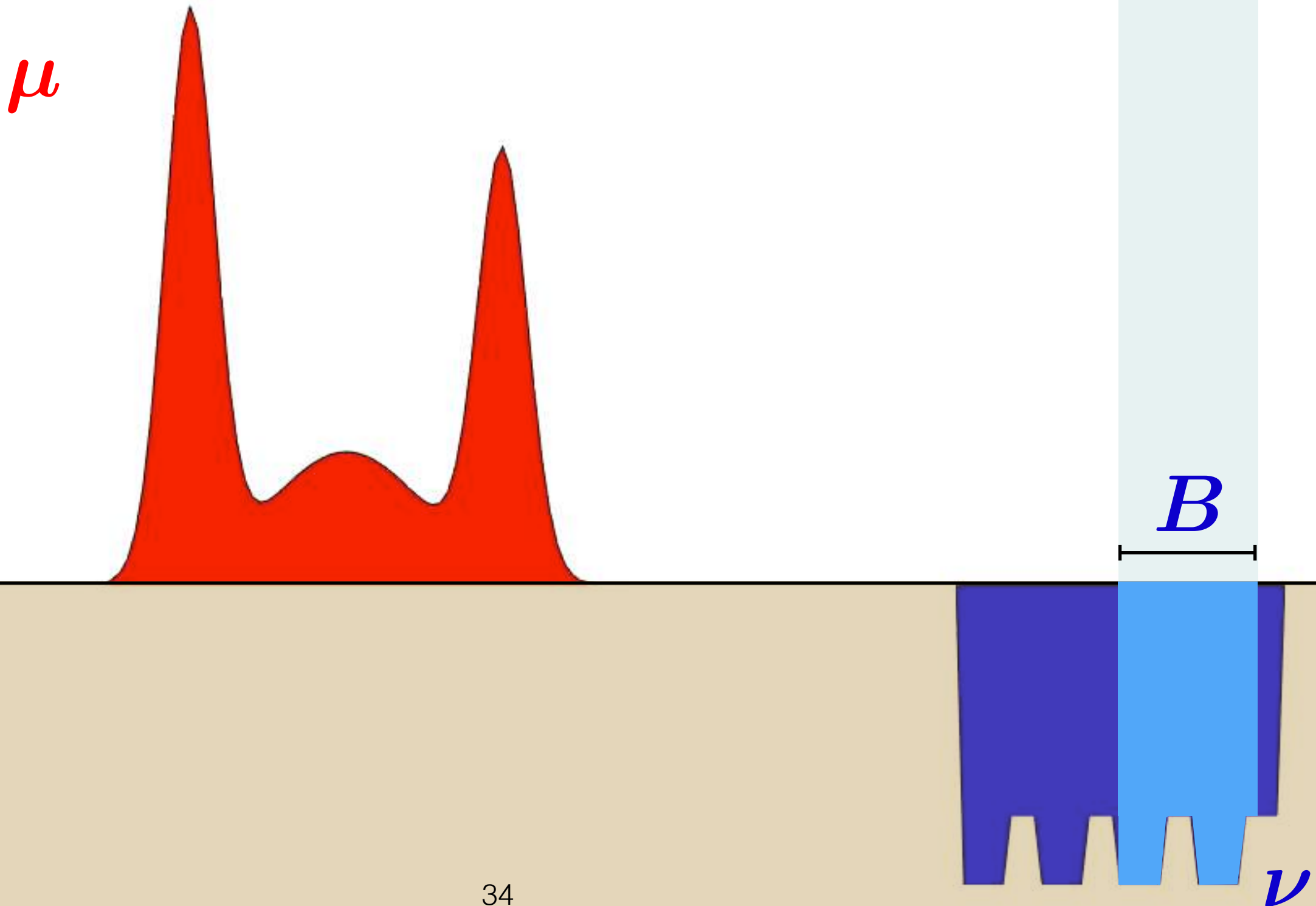
B



ν

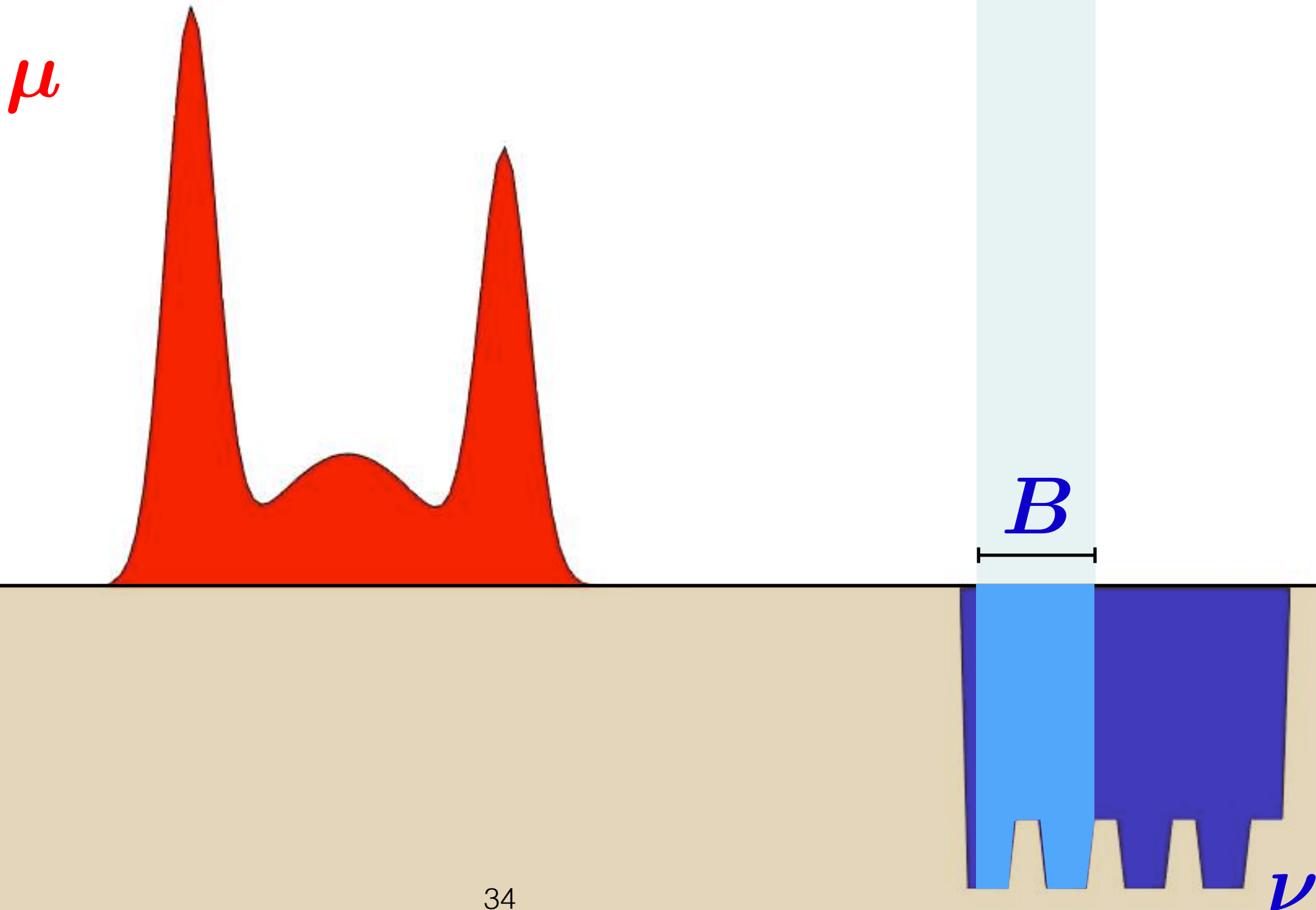
Origins: Monge's Problem

T must map red to blue.



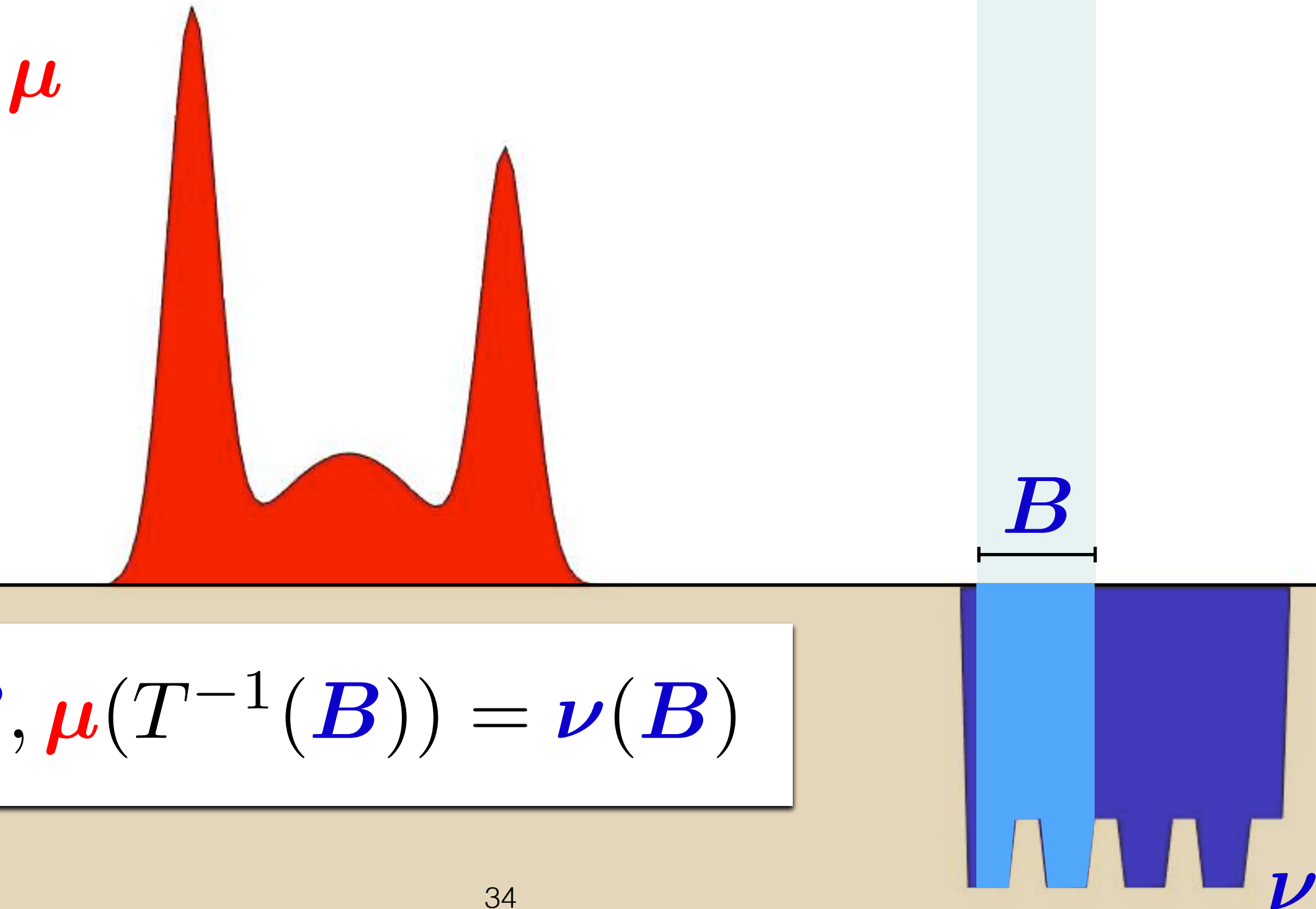
Origins: Monge's Problem

T must map red to blue.



Origins: Monge's Problem

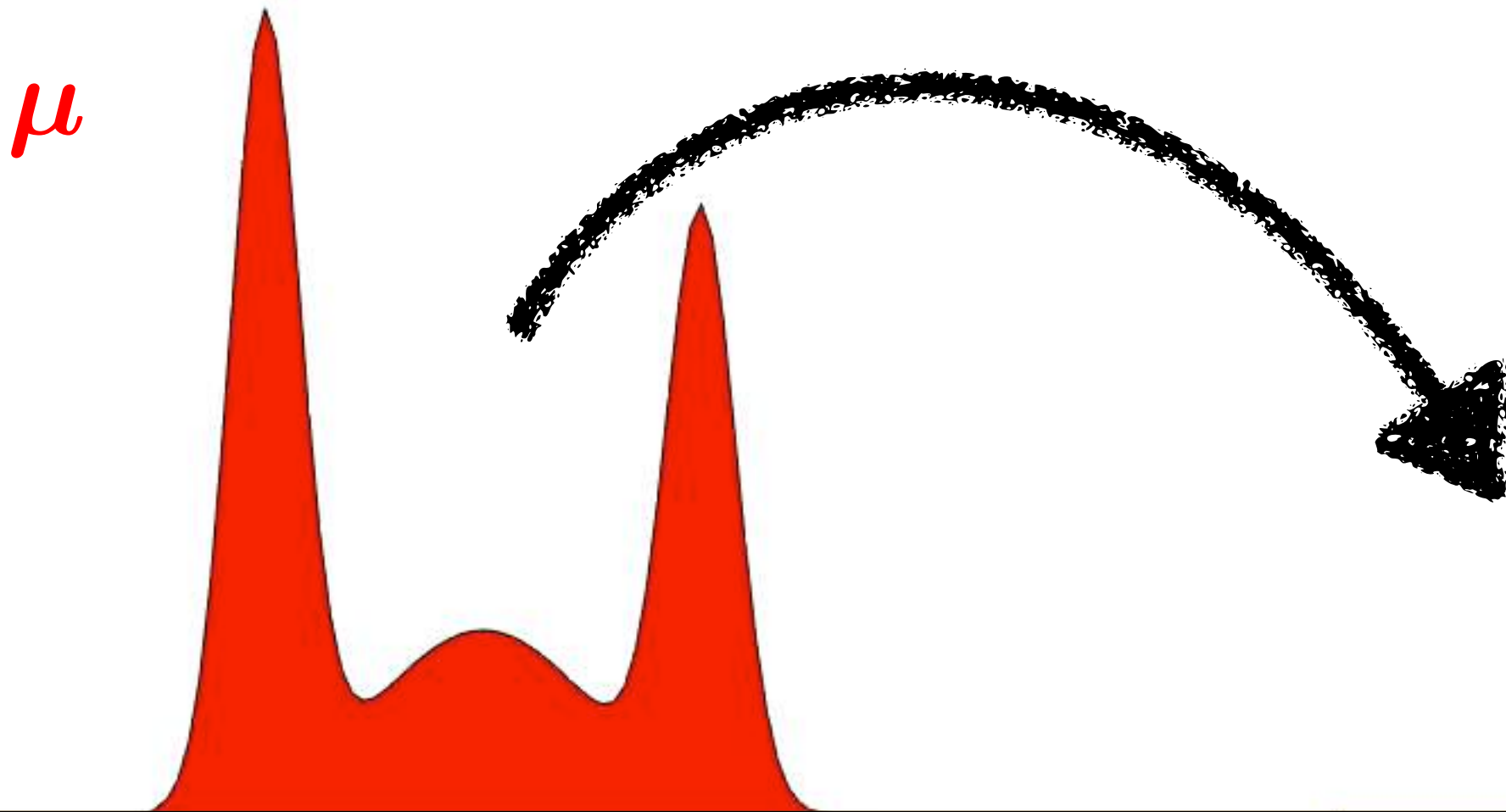
T must map red to blue.



$$\forall B, \mu(T^{-1}(B)) = \nu(B)$$

Origins: Monge's Problem

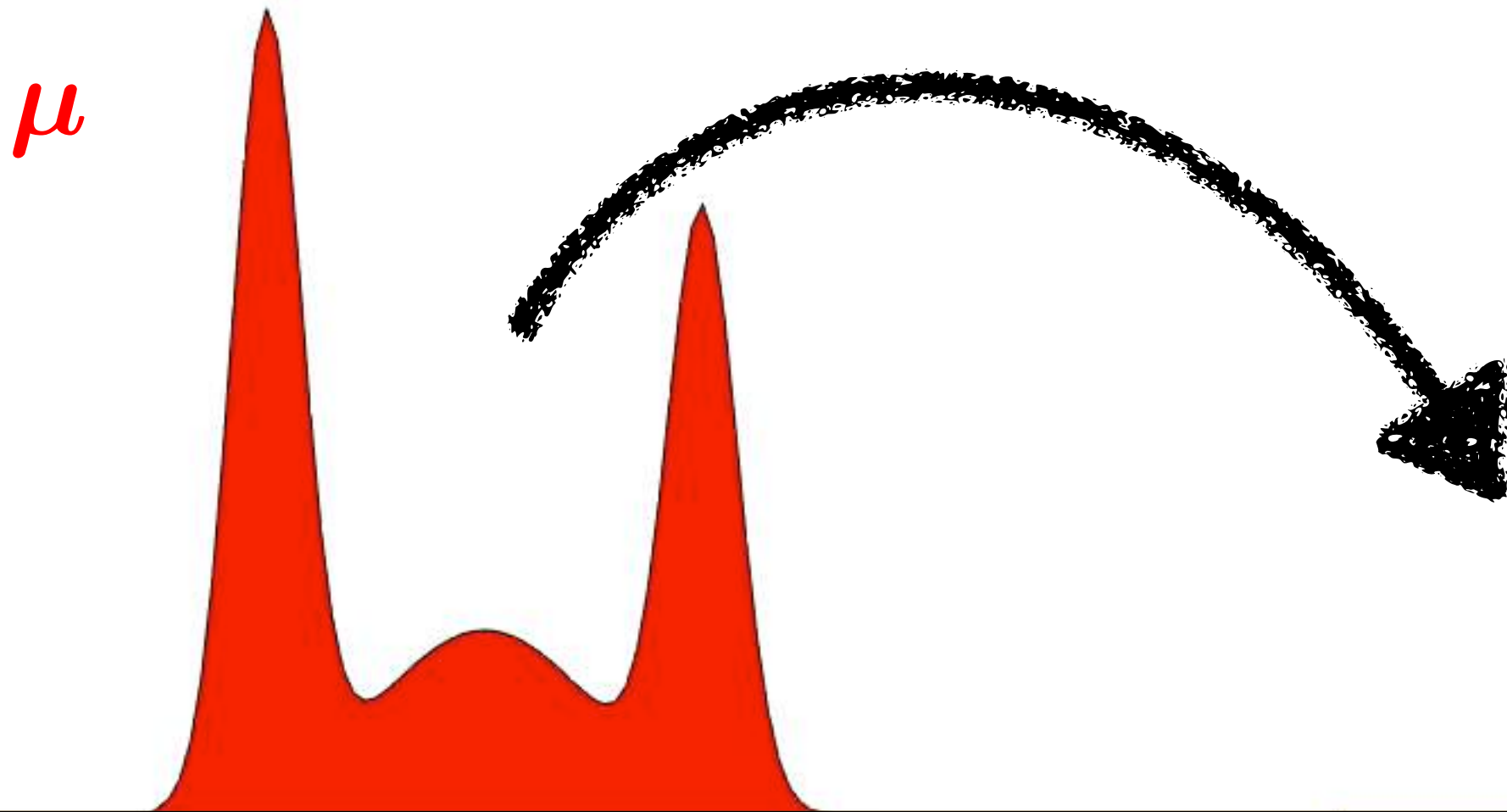
*T must **push-forward** the red measure towards the blue*



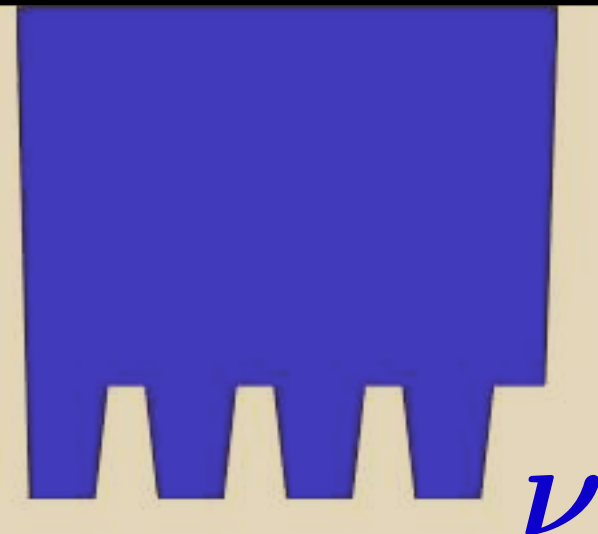
$$T_{\#} \mu = \nu$$

Origins: Monge's Problem

T must **push-forward** the red measure towards the blue



What T s.t. $T_{\#}\mu = \nu$
minimizes $\int D(x, T(x)) \mu(dx)$?



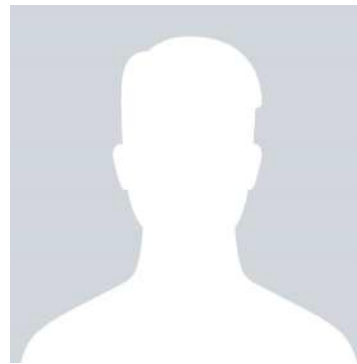
Kantorovich Problem



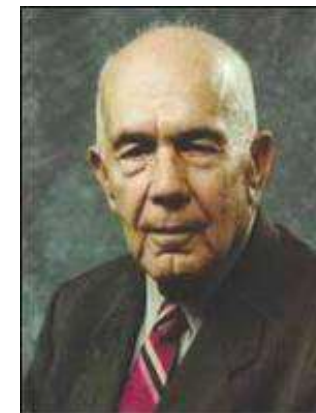
Kantorovich



1939



Tolstoi
1930



Hitchcock

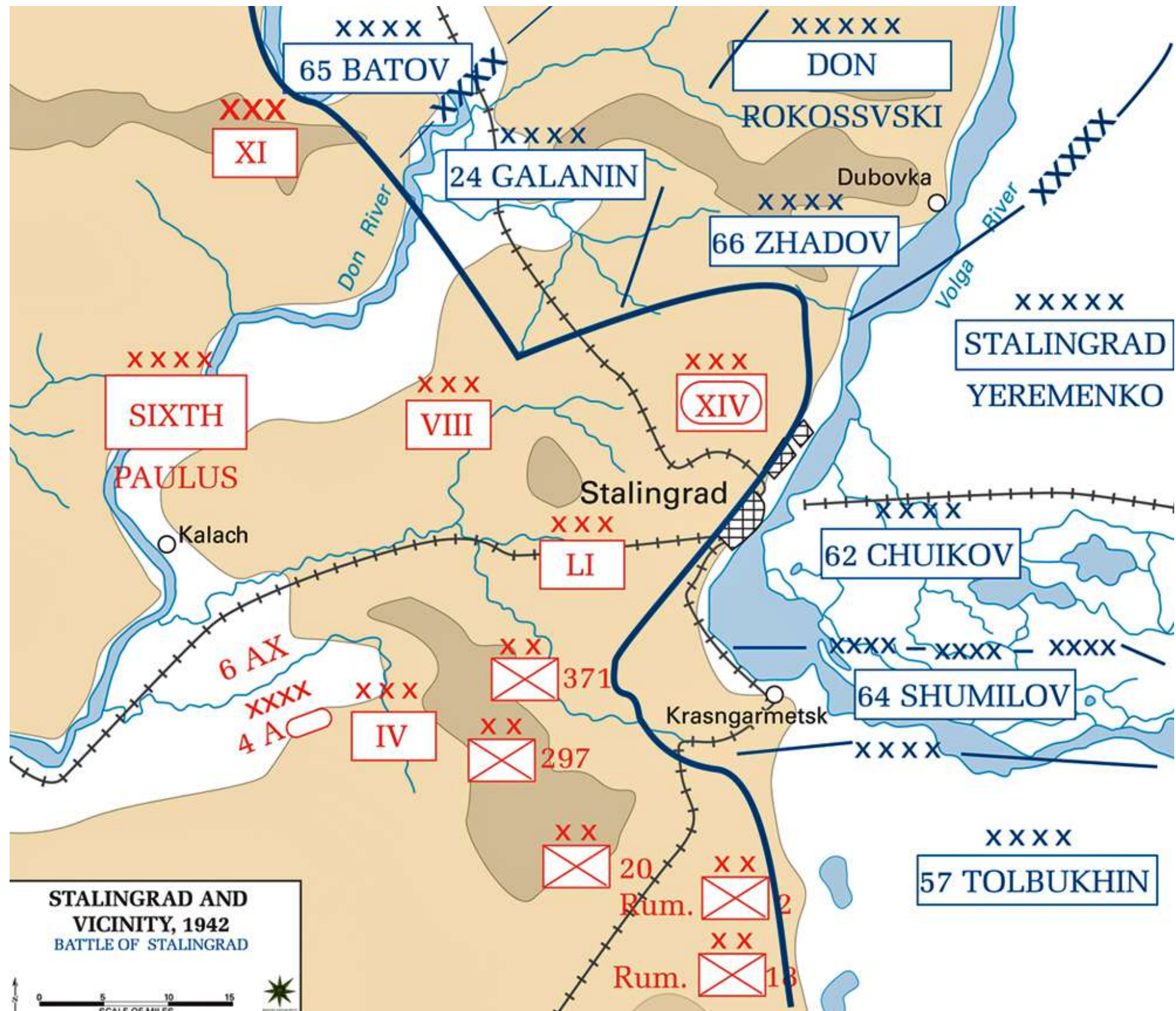
THE DISTRIBUTION OF A PRODUCT FROM SEVERAL SOURCES TO NUMEROUS LOCALITIES

BY FRANK L. HITCHCOCK

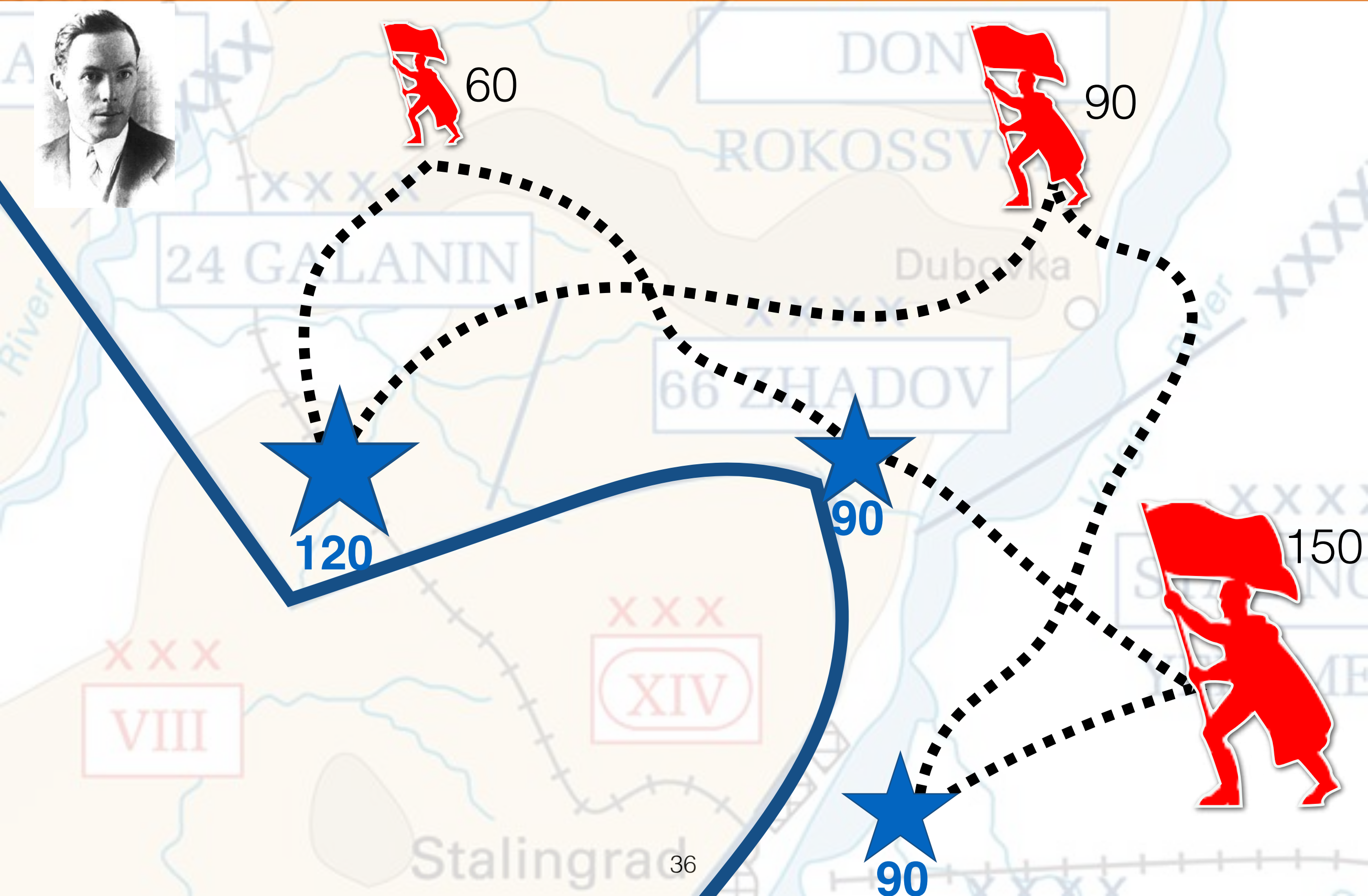
1. Statement of the problem. When several factories supply a product to a number of cities we desire the least costly manner of distribution. Due to freight rates and other matters the cost of a ton of product to a particular city will vary according to which factory supplies it, and will also vary from city to city.

1941

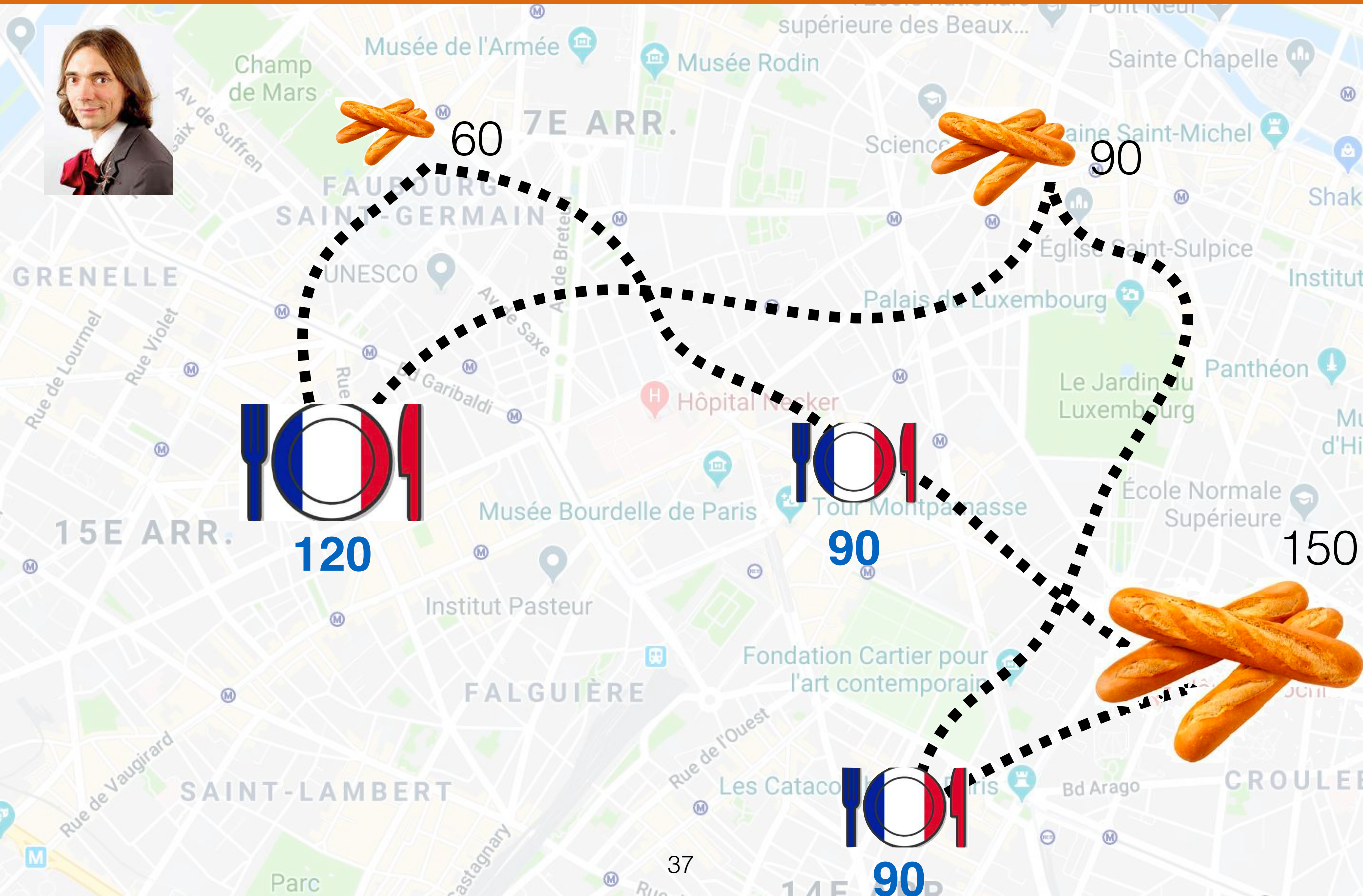
Kantorovich Problem



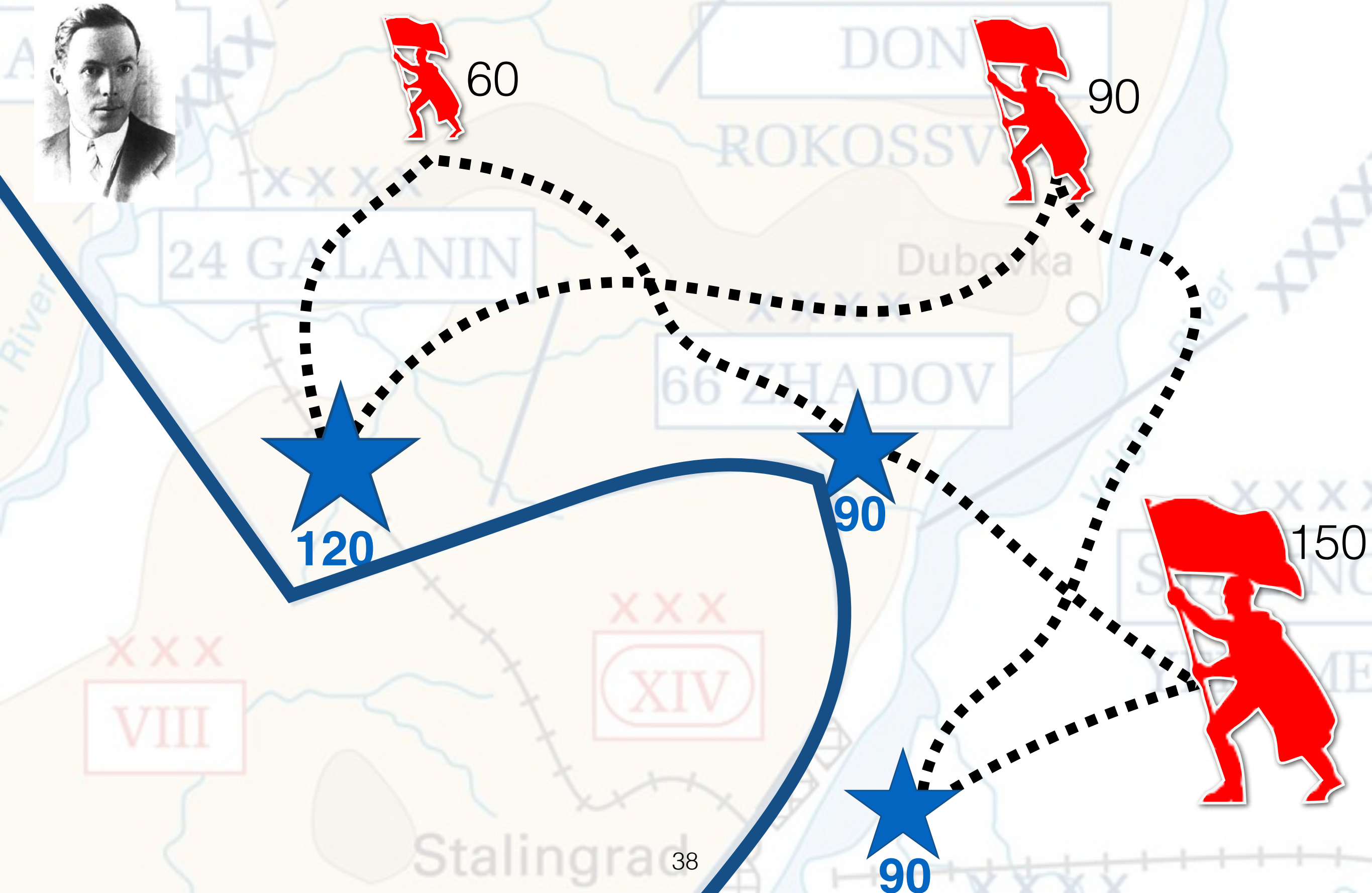
Kantorovich Problem



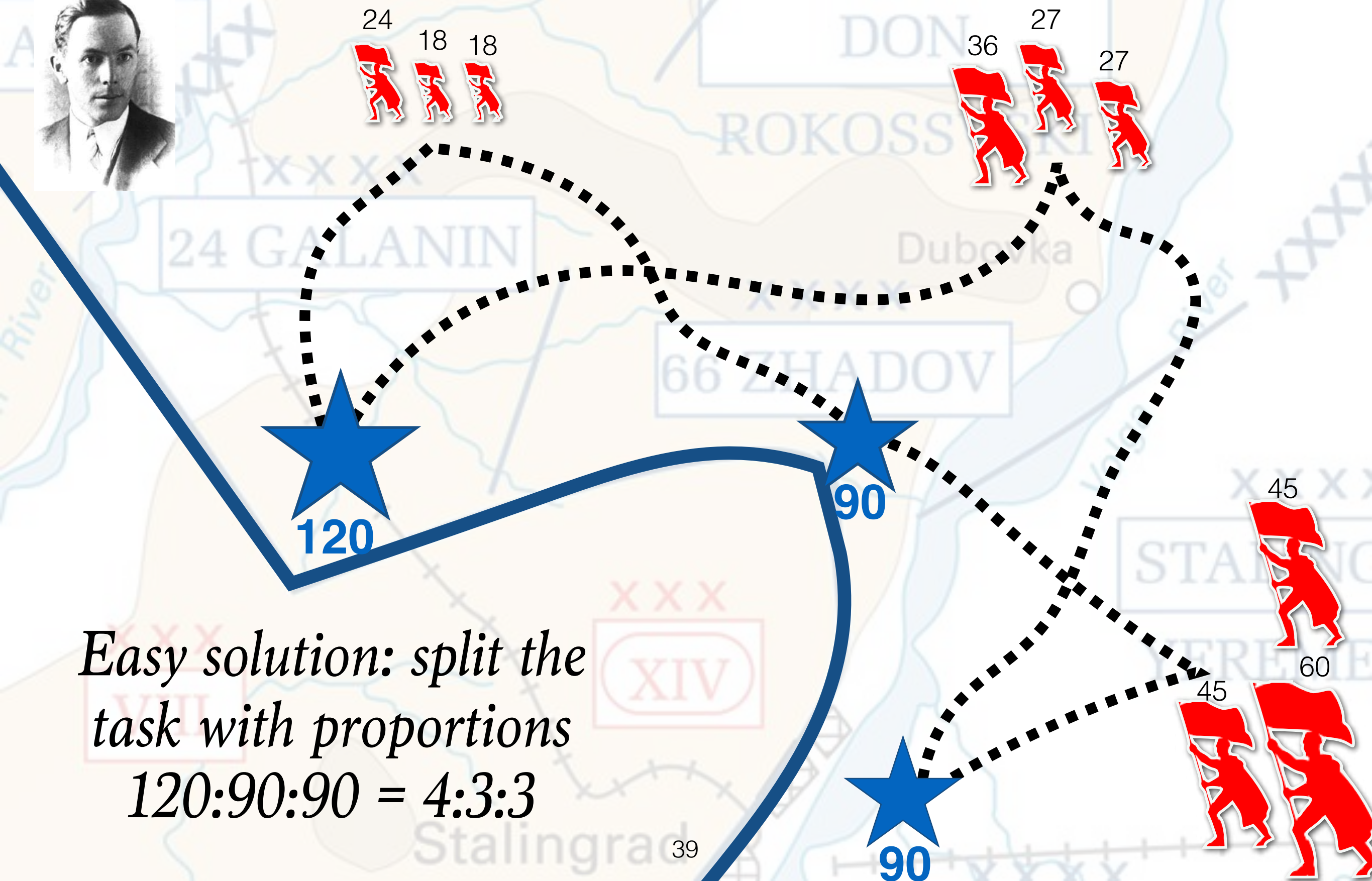
Kantorovich Problem *à la française*



Kantorovich Problem

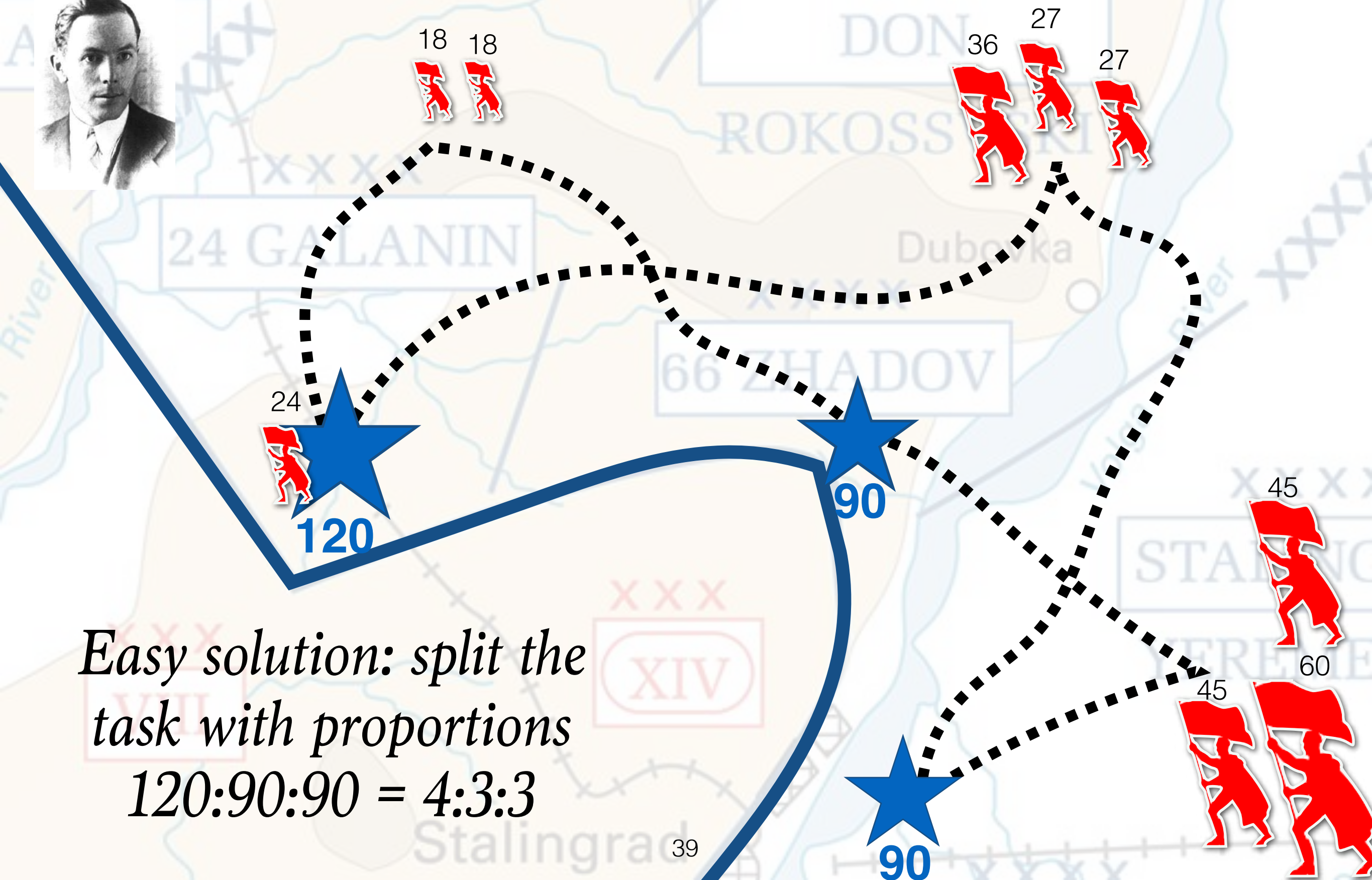


Kantorovich Problem



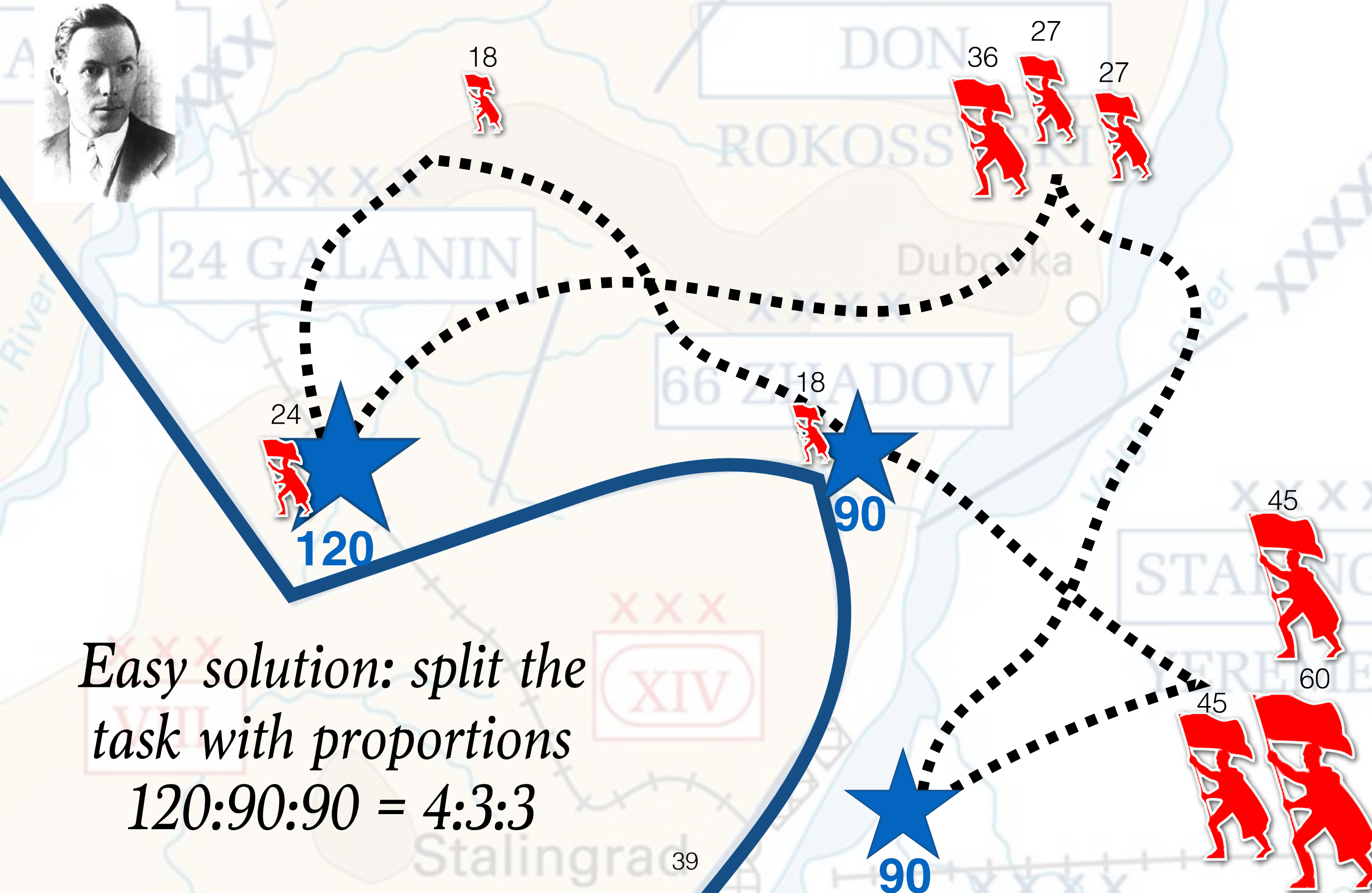
Easy solution: split the task with proportions
 $120:90:90 = 4:3:3$

Kantorovich Problem



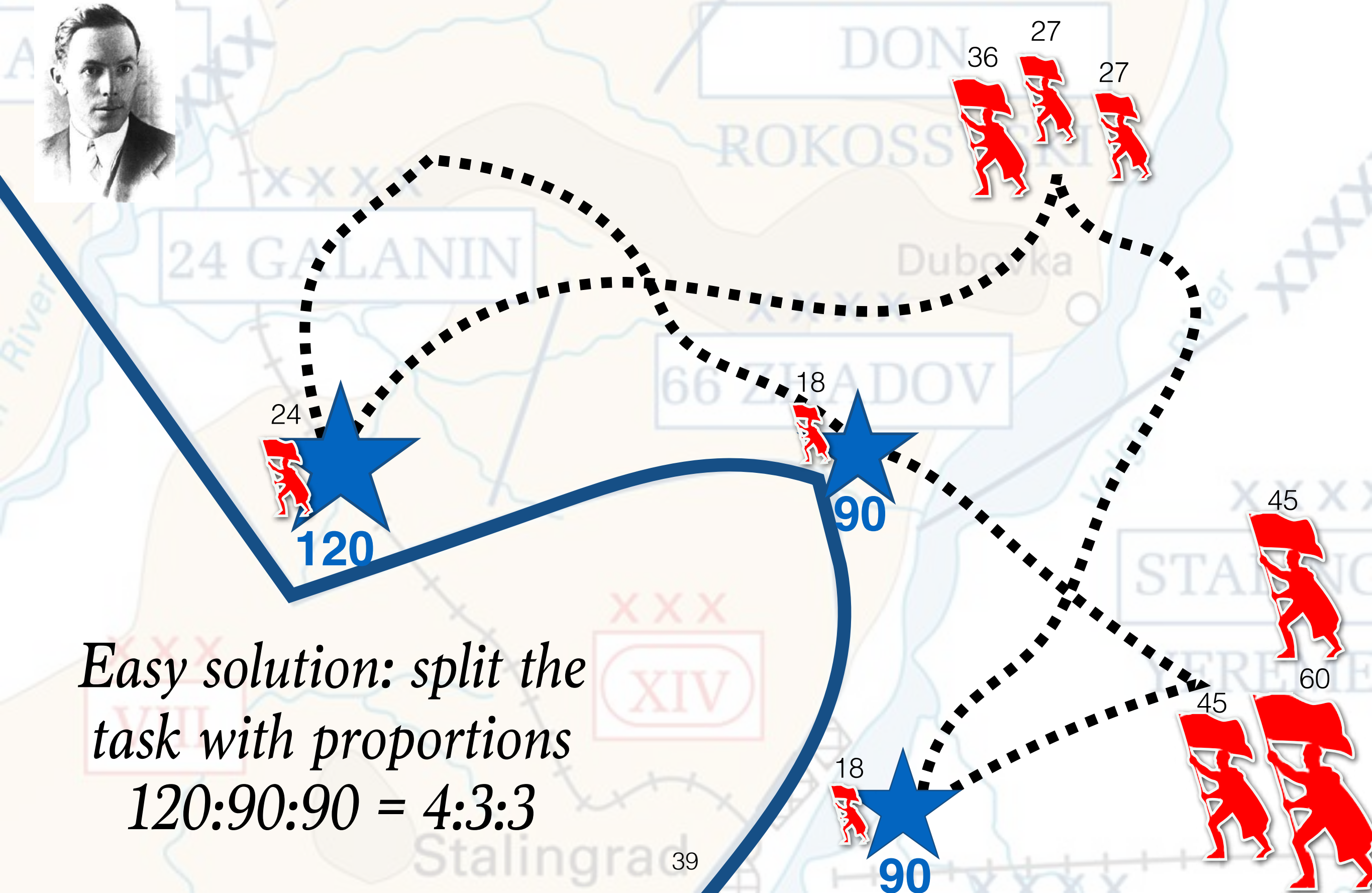
Easy solution: split the task with proportions
 $120:90:90 = 4:3:3$

Kantorovich Problem



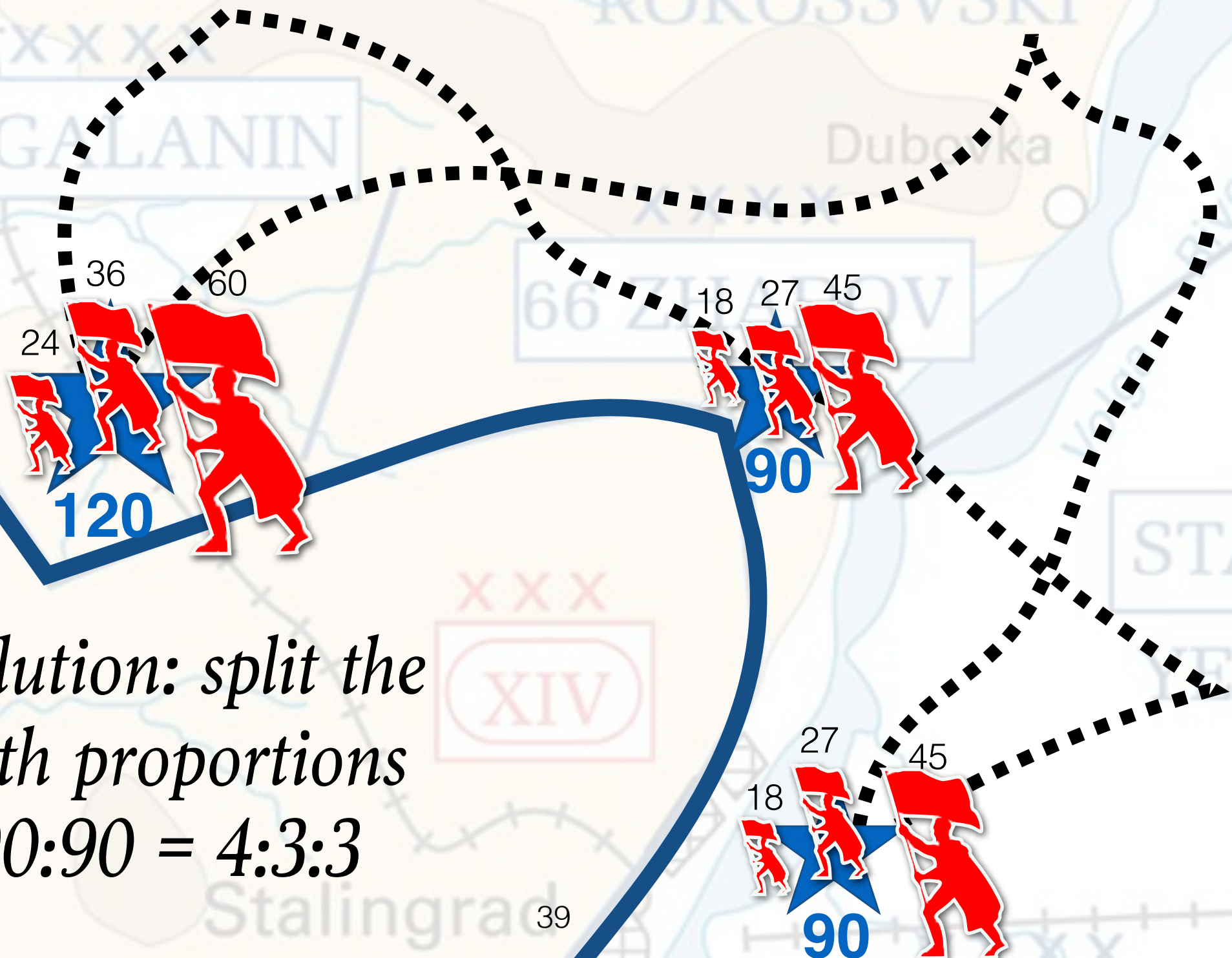
Easy solution: split the task with proportions
 $120:90:90 = 4:3:3$

Kantorovich Problem



Easy solution: split the task with proportions
 $120:90:90 = 4:3:3$

Kantorovich Problem



Easy solution: split the task with proportions
 $120:90:90 = 4:3:3$

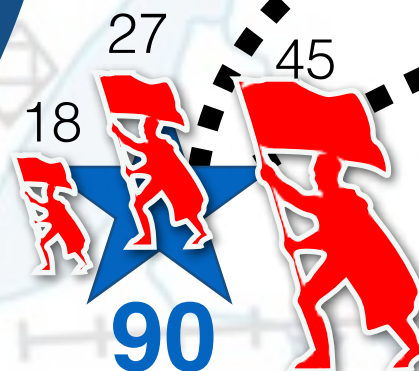
Kantorovich Problem



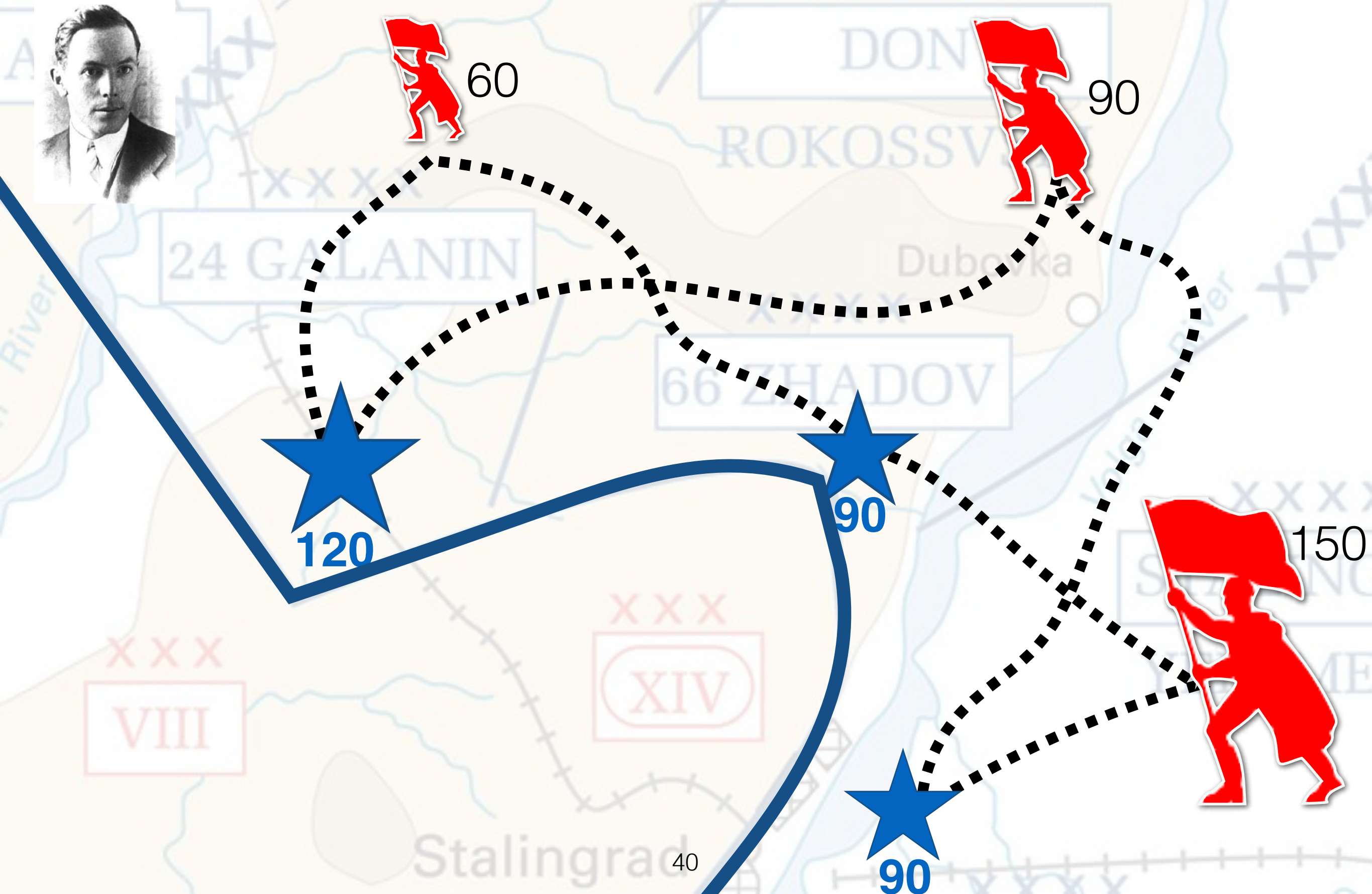
**Naive approach results in
many displacements...**

**Can we find a cheaper
alternative?**

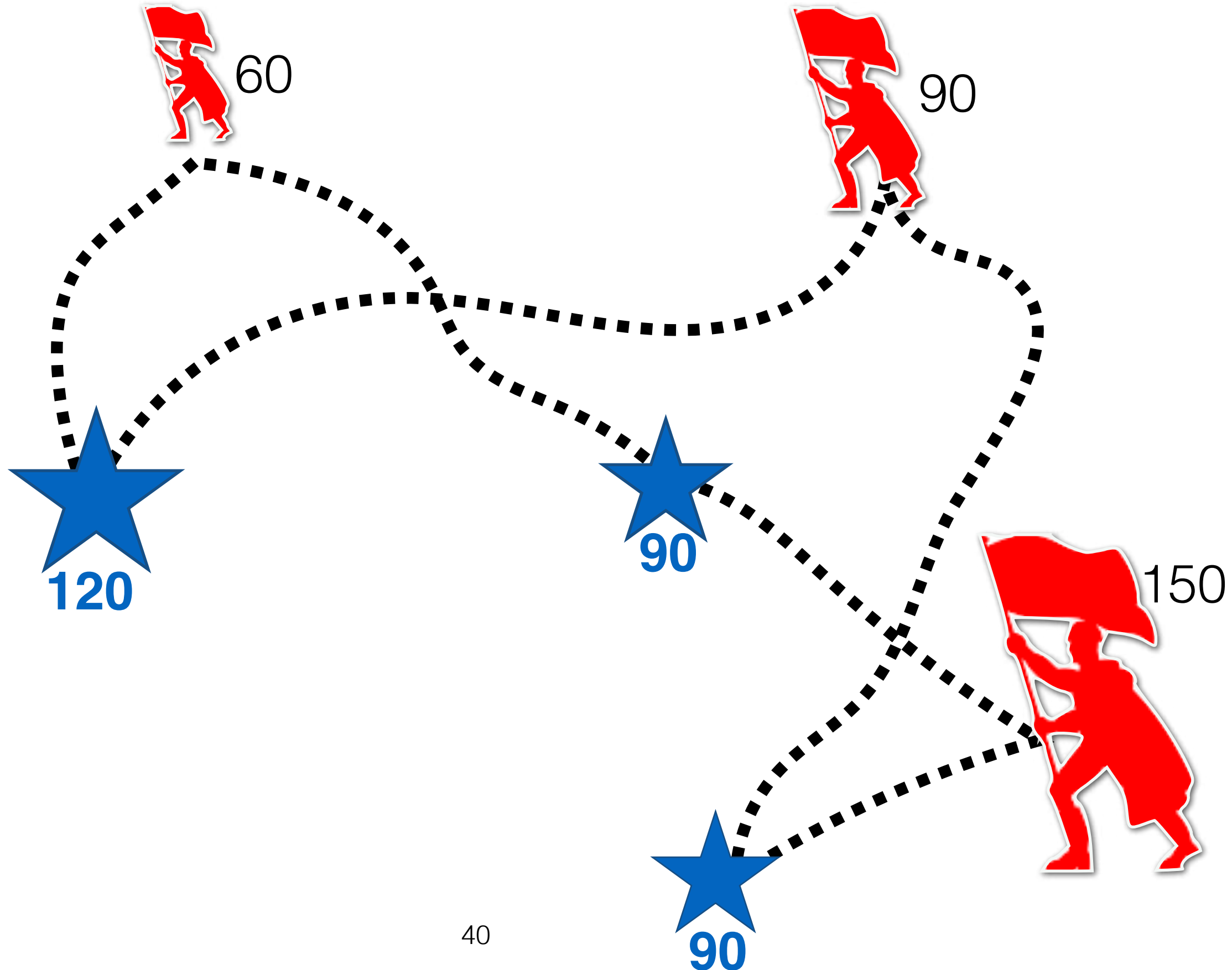
*Easy solution: split the
task with proportions
 $120:90:90 = 4:3:3$*



Kantorovich Problem



Kantorovich Problem



Kantorovich Problem



60

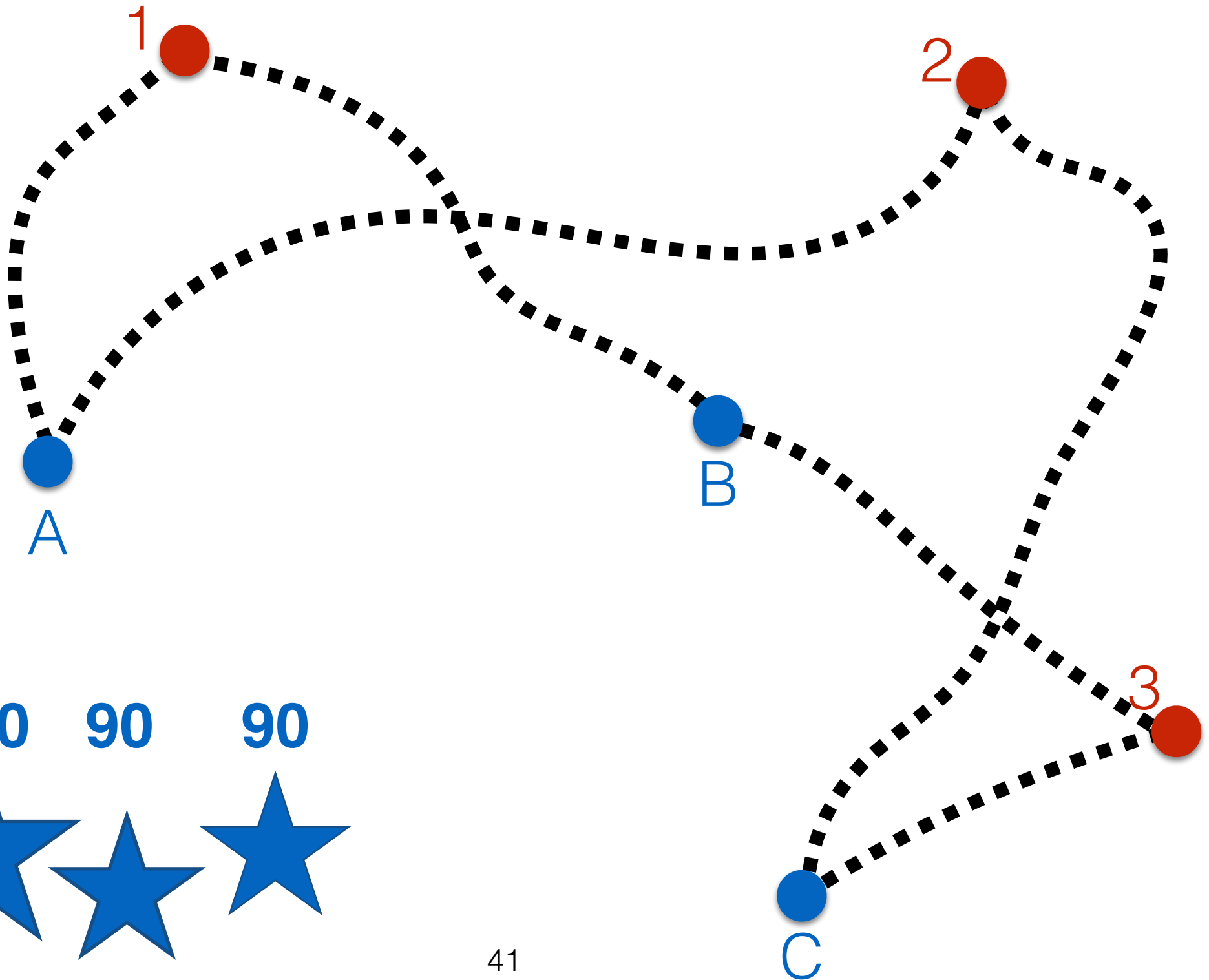
90

150

120

90

90



Kantorovich Problem



60

90

150

120

90

90



Kantorovich Problem



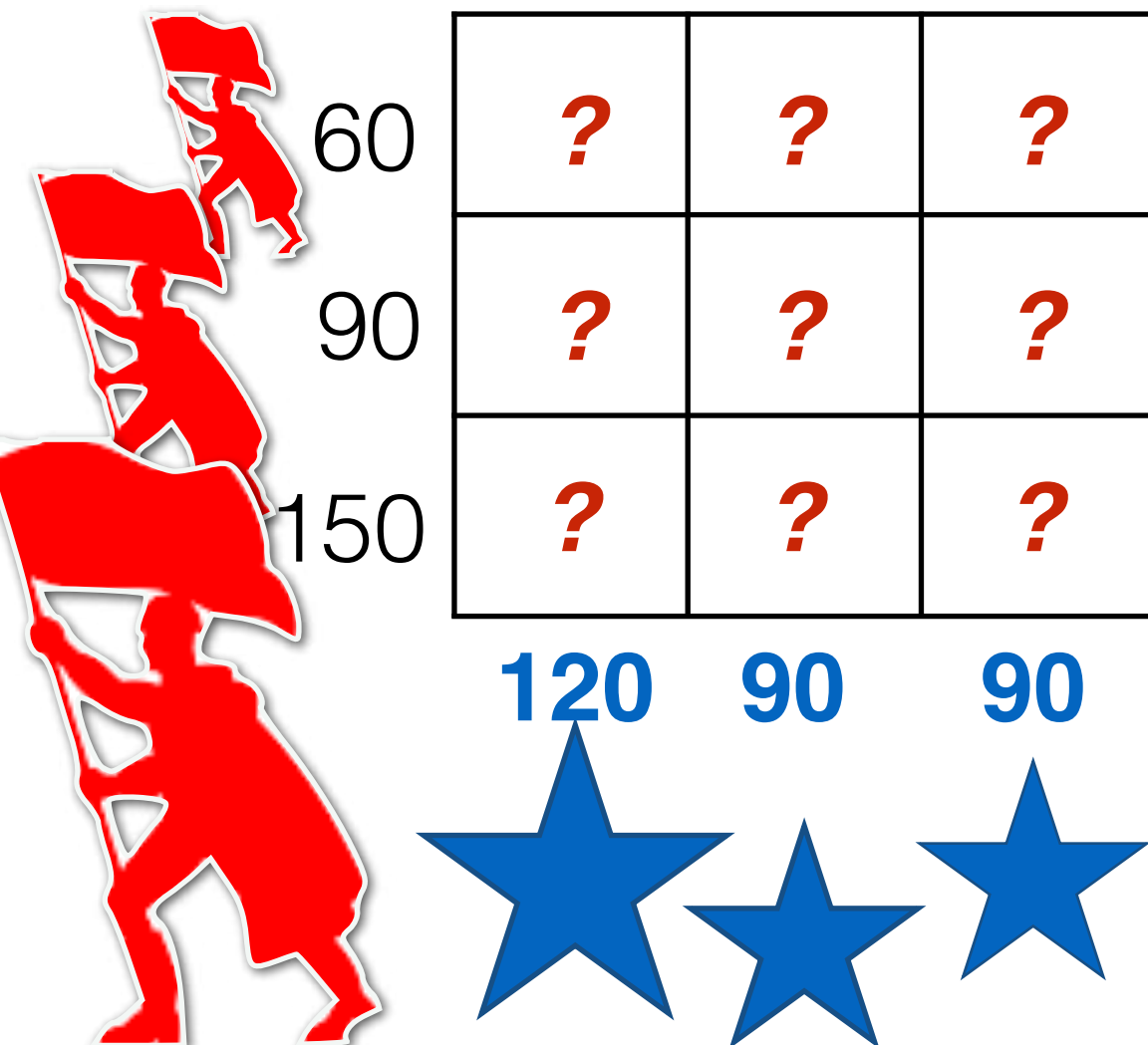
60	?	?	?
90	?	?	?
150	?	?	?
	120	90	90
	★	★	★

1 ●	d_{1A}	d_{1B}	d_{1C}
2 ●	d_{2A}	d_{2B}	d_{2C}
3 ●	d_{3A}	d_{3B}	d_{3C}
	● A	● B	● C

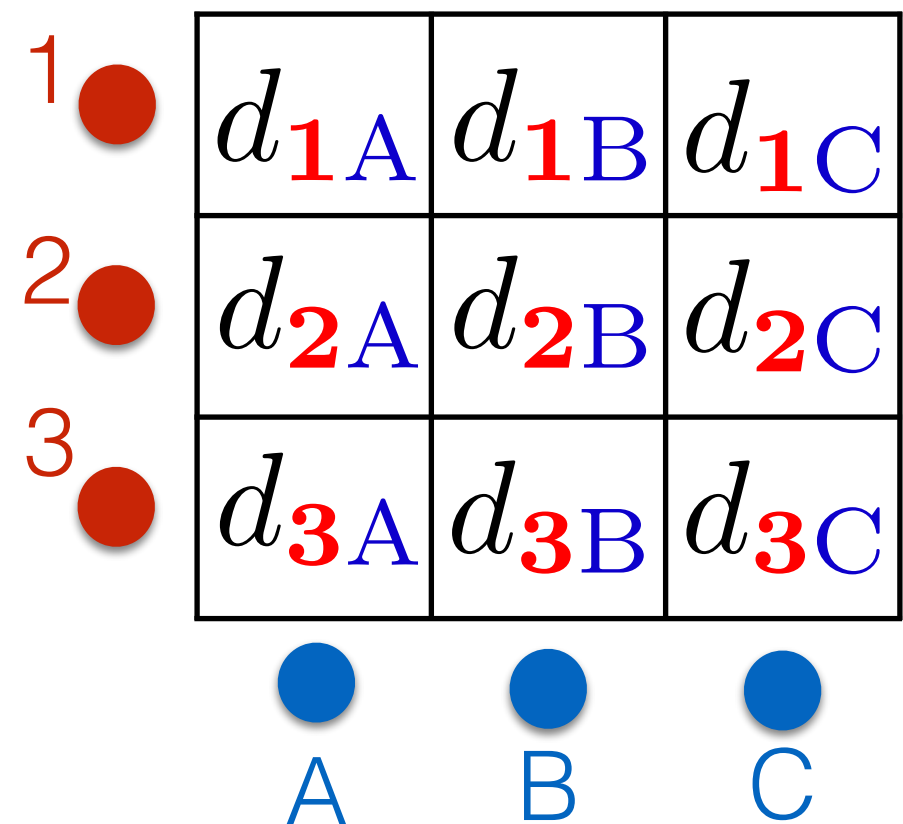
Kantorovich Problem



Transportation matrix



Distance matrix



Kantorovich Problem



*The problem is entirely described by
counts and a **cost/distance matrix***

Transportation matrix



60	?	?	?
90	?	?	?
150	?	?	?
	120	90	90

Three blue stars are positioned below the bottom row of the table.

Distance matrix

1 ●	d_{1A}	d_{1B}	d_{1C}
2 ●	d_{2A}	d_{2B}	d_{2C}
3 ●	d_{3A}	d_{3B}	d_{3C}
	● A	● B	● C

Kantorovich Problem

Transportation matrix

60	?	?	?
90	?	?	?
150	?	?	?
	120	90	90

Distance matrix

1	d_{1A}	d_{1B}	d_{1C}
2	d_{2A}	d_{2B}	d_{2C}
3	d_{3A}	d_{3B}	d_{3C}
	A	B	C

Kantorovich Problem

Transportation matrix

60	p_{1A}	p_{1B}	p_{1C}
90	p_{2A}	p_{2B}	p_{2C}
150	p_{3A}	p_{3B}	p_{3C}
	120	90	90

Distance matrix

1	d_{1A}	d_{1B}	d_{1C}
2	d_{2A}	d_{2B}	d_{2C}
3	d_{3A}	d_{3B}	d_{3C}
	A	B	C

Kantorovich Problem

Transportation matrix

a_1	p_{1A}	p_{1B}	p_{1C}
a_2	p_{2A}	p_{2B}	p_{2C}
a_3	p_{3A}	p_{3B}	p_{3C}
	b_A	b_B	b_C

Distance matrix

1	d_{1A}	d_{1B}	d_{1C}
2	d_{2A}	d_{2B}	d_{2C}
3	d_{3A}	d_{3B}	d_{3C}
	A	B	C

Kantorovich Problem

Transportation matrix

a_1	p_{1A}	p_{1B}	p_{1C}
a_2	p_{2A}	p_{2B}	p_{2C}
a_3	p_{3A}	p_{3B}	p_{3C}
	b_A	b_B	b_C

Distance matrix

1	d_{1A}	d_{1B}	d_{1C}
2	d_{2A}	d_{2B}	d_{2C}
3	d_{3A}	d_{3B}	d_{3C}
	A	B	C

Constraints

$$\forall i \in \{1, 2, 3\}, \quad \sum_{j \in \{A, B, C\}} p_{ij} = a_i$$

$$\forall j \in \{A, B, C\}, \quad \sum_{i \in \{1, 2, 3\}} p_{ij} = b_j$$

$$p_{ij} \geq 0$$

Kantorovich Problem

Transportation matrix

a_1	p_{1A}	p_{1B}	p_{1C}
a_2	p_{2A}	p_{2B}	p_{2C}
a_3	p_{3A}	p_{3B}	p_{3C}
	b_A	b_B	b_C

Distance matrix

1	d_{1A}	d_{1B}	d_{1C}
2	d_{2A}	d_{2B}	d_{2C}
3	d_{3A}	d_{3B}	d_{3C}
	A	B	C

Constraints

$$\forall i \in \{1, 2, 3\}, \quad \sum_{j \in \{A, B, C\}} p_{ij} = a_i$$

$$\forall j \in \{A, B, C\}, \quad \sum_{i \in \{1, 2, 3\}} p_{ij} = b_j$$

$$p_{ij} \geq 0$$

Cost function

$$C(P) = \sum_{j \in \{A, B, C\}} \sum_{i \in \{1, 2, 3\}} p_{ij} d_{ij}$$

Kantorovich Problem

Transportation matrix

a_1	p_{1A}	p_{1B}	p_{1C}
a_2	p_{2A}	p_{2B}	p_{2C}
a_3	p_{3A}	p_{3B}	p_{3C}
	b_A	b_B	b_C

Distance matrix

1	d_{1A}	d_{1B}	d_{1C}
2	d_{2A}	d_{2B}	d_{2C}
3	d_{3A}	d_{3B}	d_{3C}
	A	B	C

Constraints

$$\forall i \in \{1, 2, 3\}, \quad \sum_{j \in \{A, B, C\}} p_{ij} = a_i$$

$$\forall j \in \{A, B, C\}, \quad \sum_{i \in \{1, 2, 3\}} p_{ij} = b_j$$

$$p_{ij} \geq 0$$

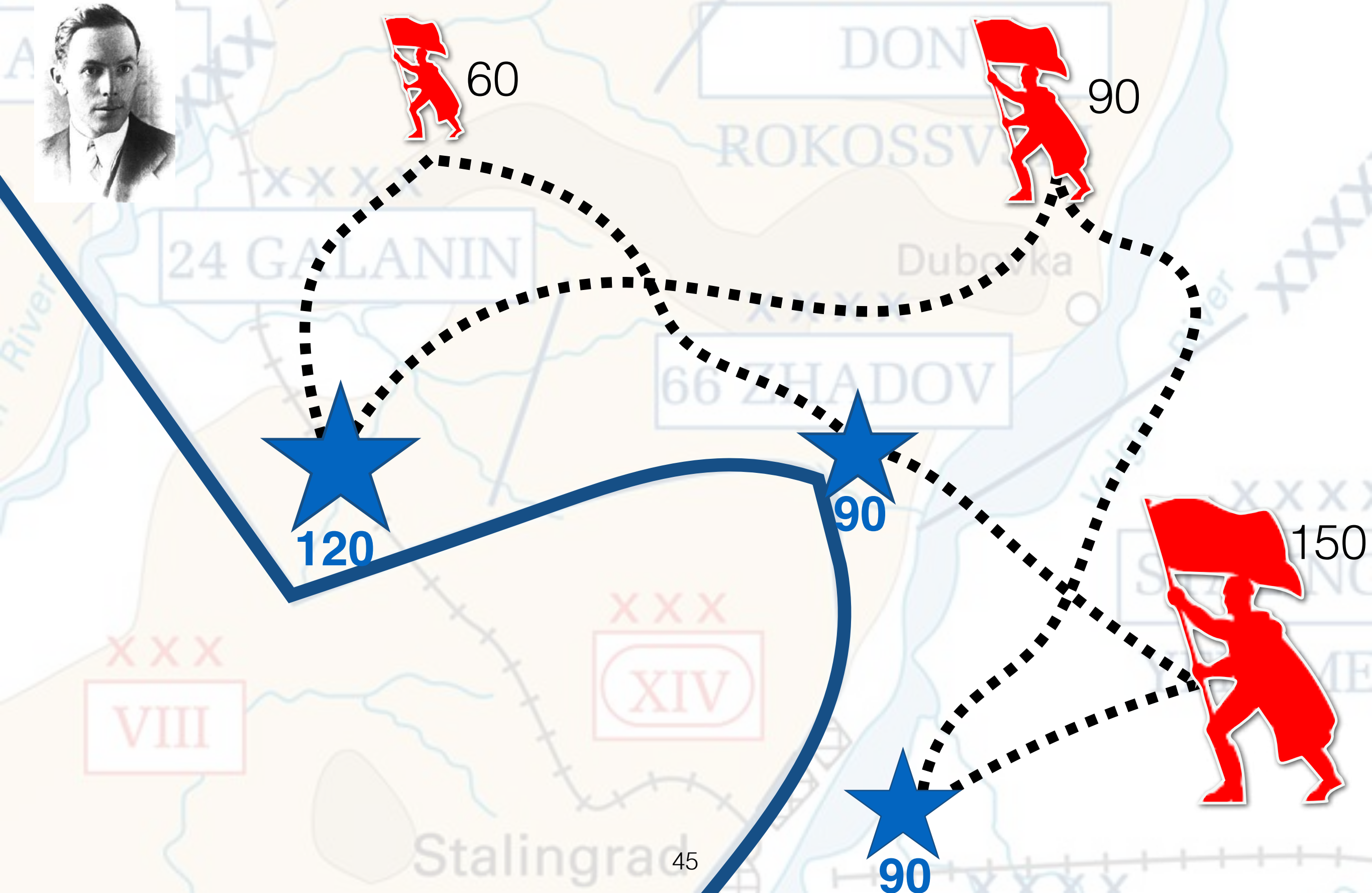
Cost function

$$C(\mathbf{P}) = \sum_{j \in \{A, B, C\}} \sum_{i \in \{1, 2, 3\}} p_{ij} d_{ij}$$

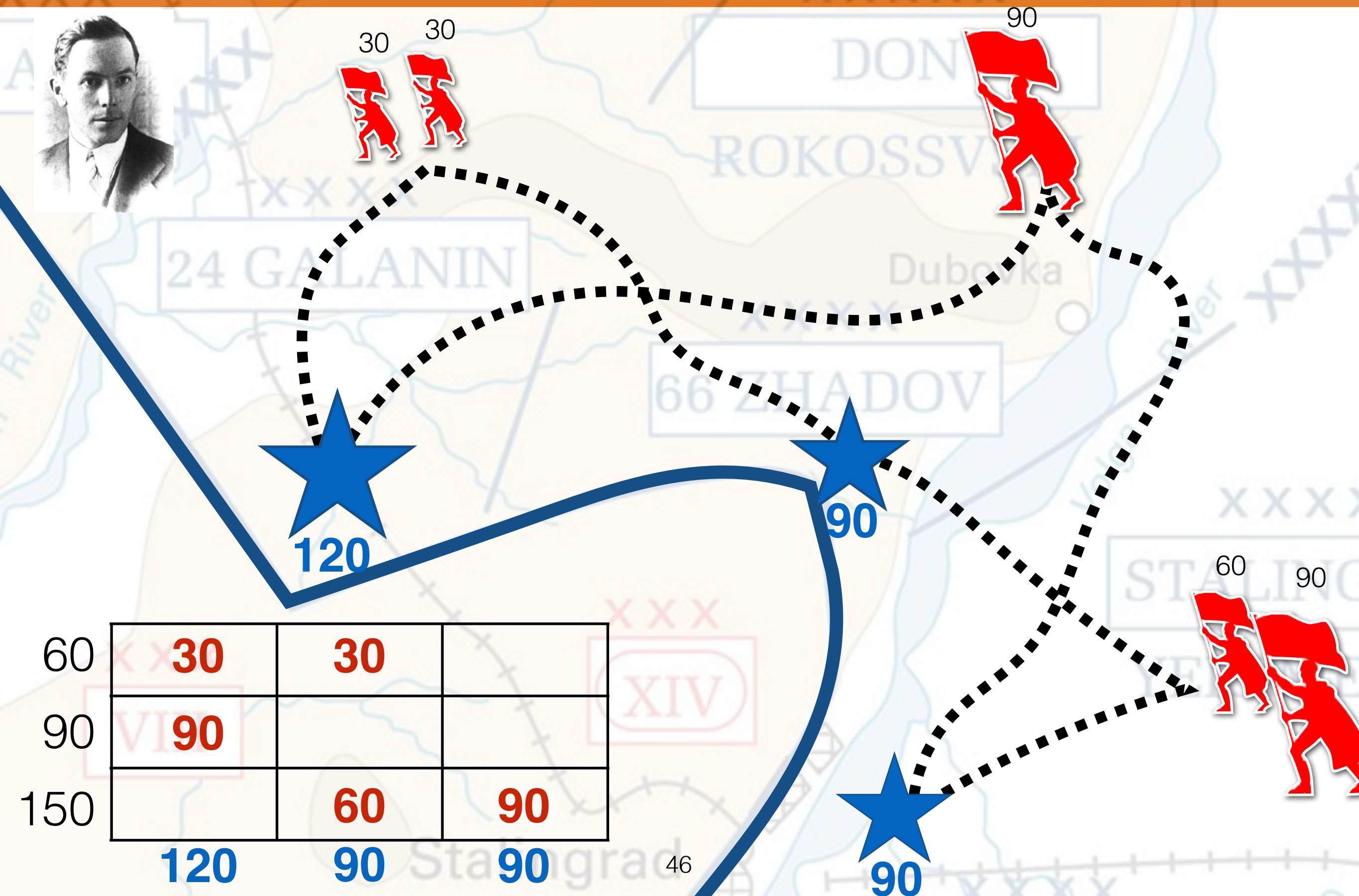
Problem

$$\min_{\text{all valid } \mathbf{P}} C(\mathbf{P})$$

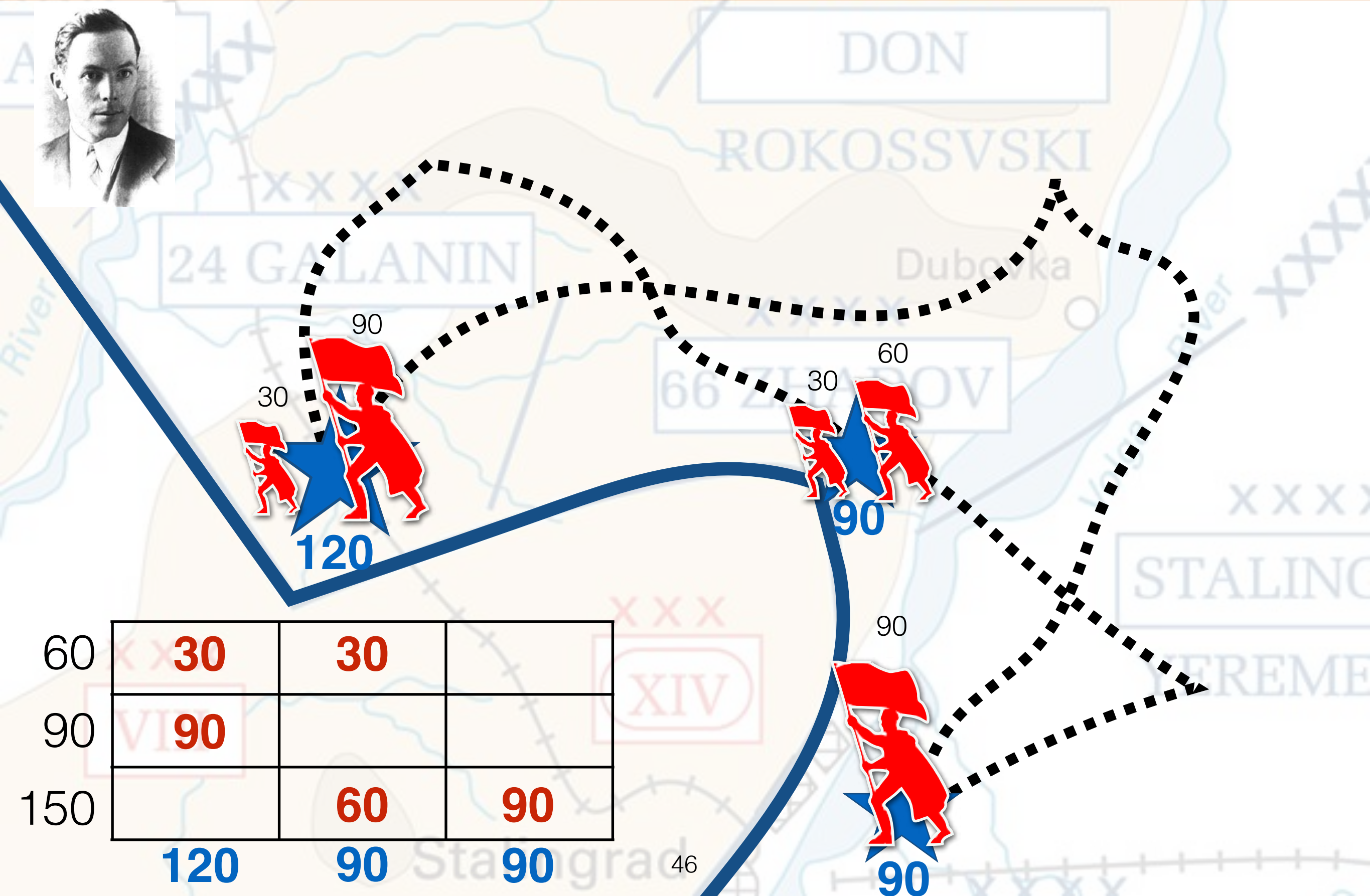
Kantorovich Problem



Kantorovich Problem

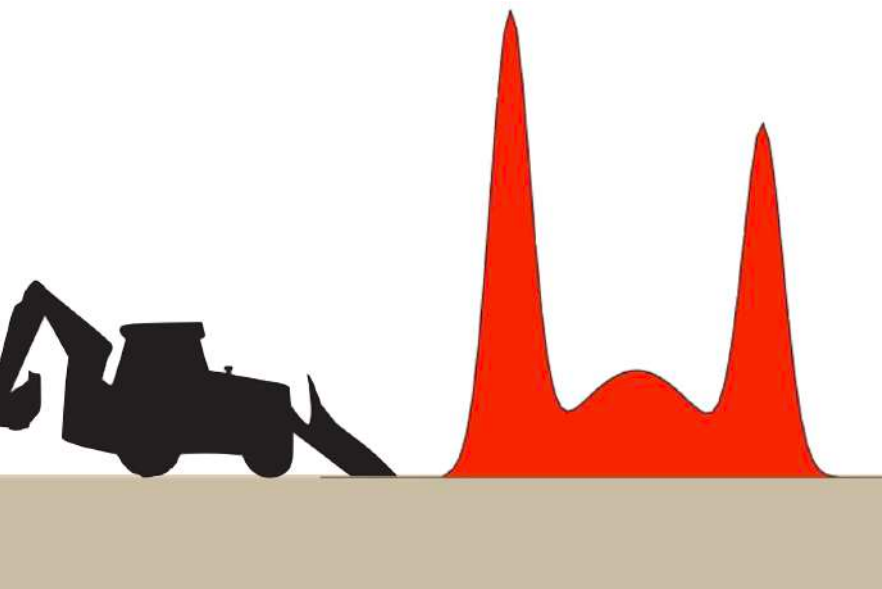


Kantorovich Problem



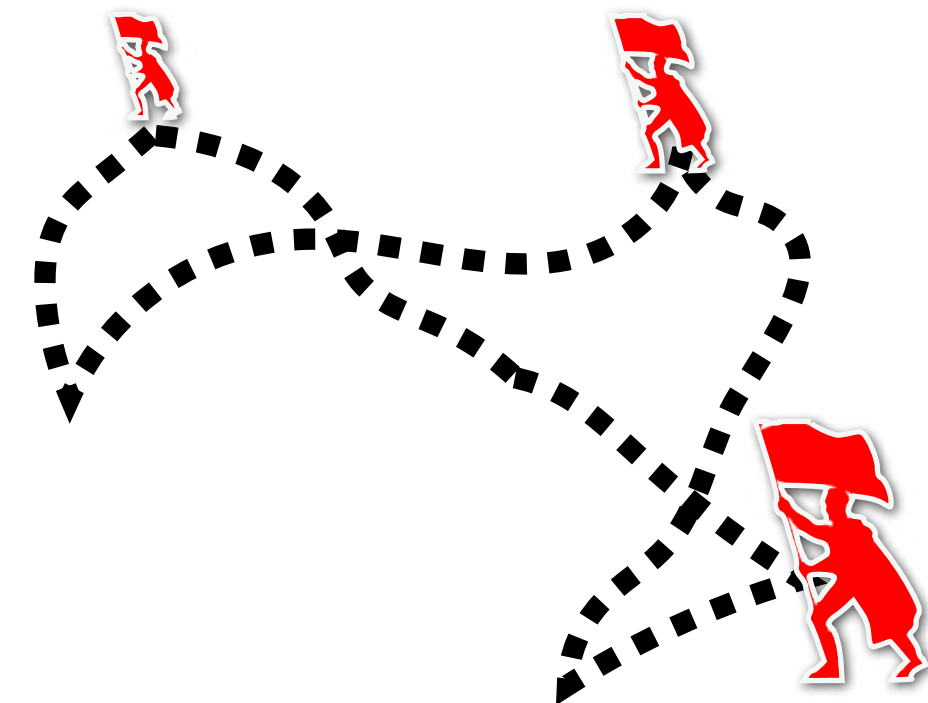
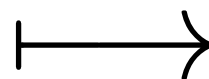
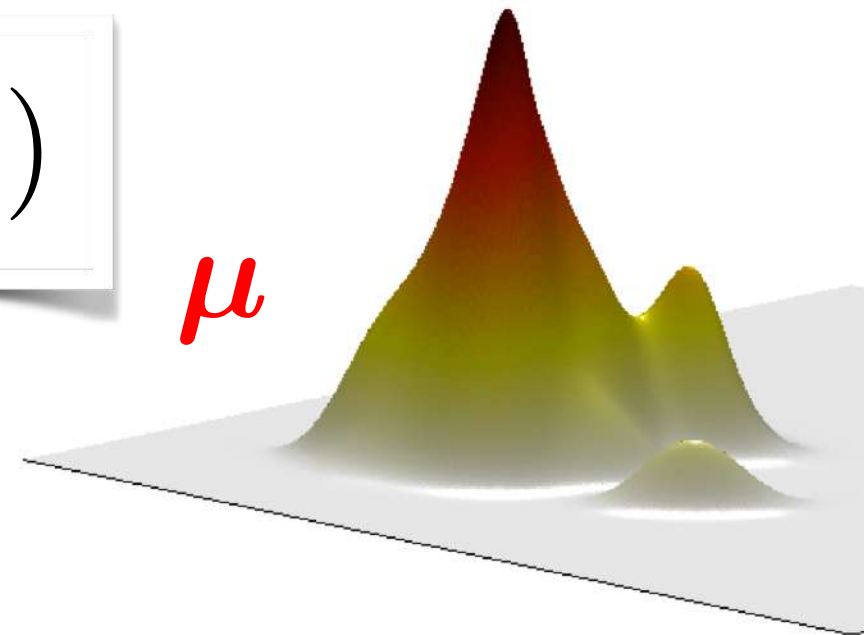
Mathematical Formalism

These problems involve discrete and continuous **probability measures** on a geometric space Ω

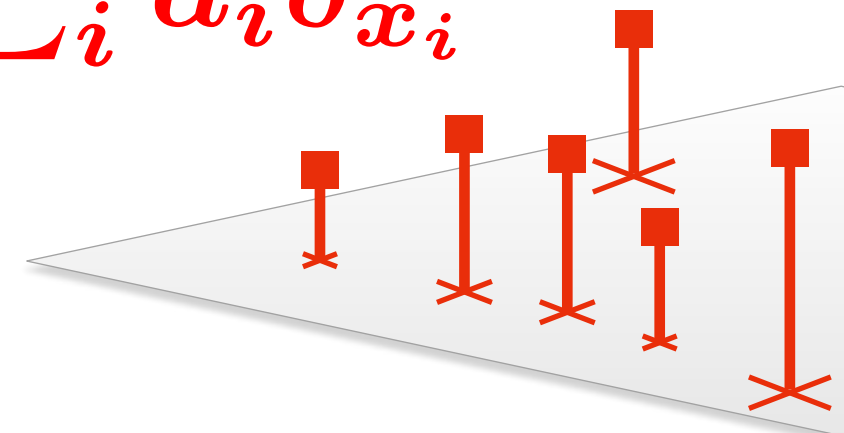


$$\mathcal{P}(\Omega)$$

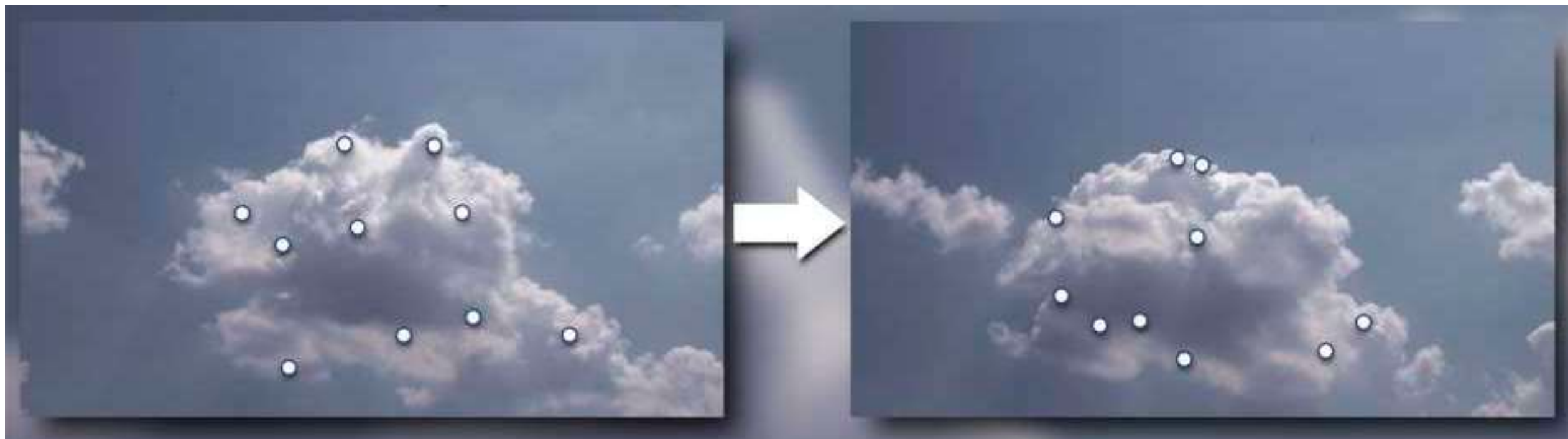
μ



$$\sum_i a_i \delta_{x_i}$$



OT: Nature's way to move particles



screenshot video: ETH Zurich / Michael Steiner, vistory GmbH



Figalli

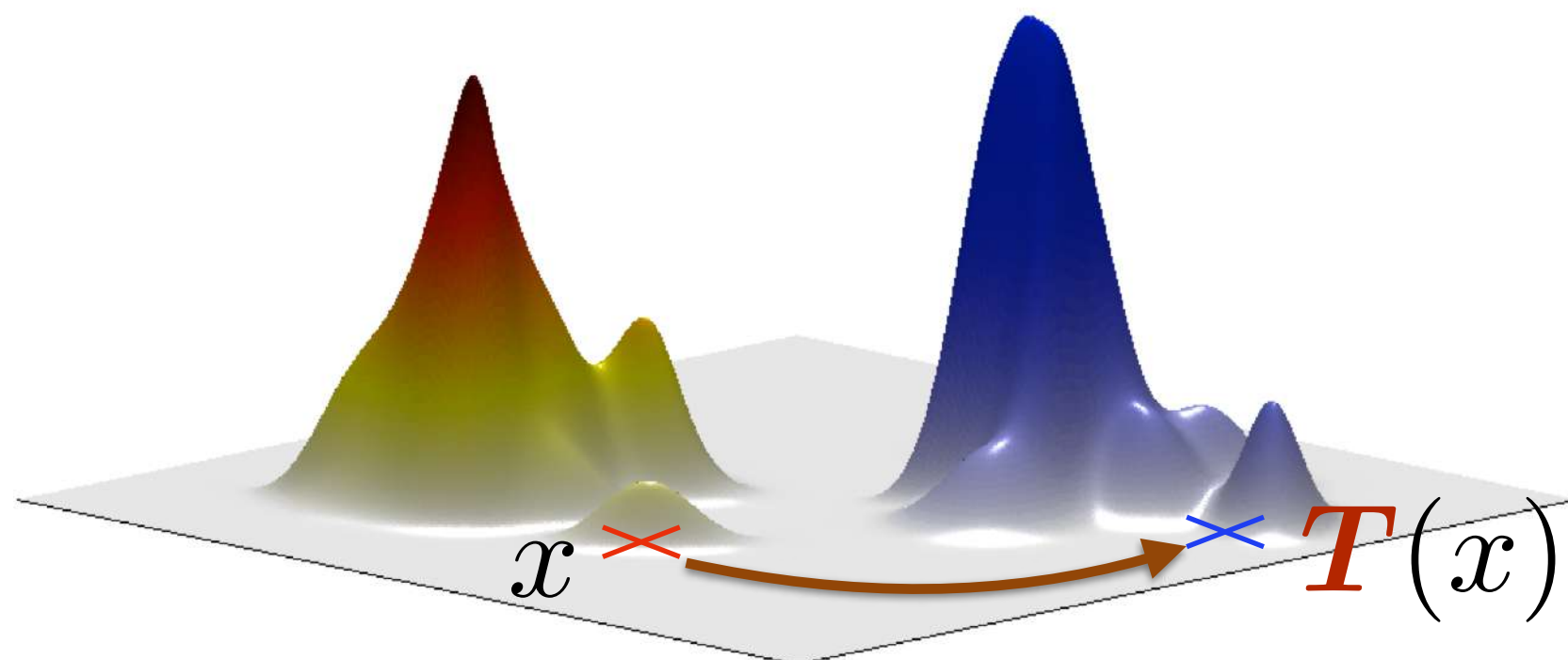
Fields'18

Monge Problem

Ω a measurable space, $\mathbf{c} : \Omega \times \Omega \rightarrow \mathbb{R}$.
 μ, ν two probability measures in $\mathcal{P}(\Omega)$.

[Monge'81] problem: find a map $T : \Omega \rightarrow \Omega$

$$\inf_{T_{\#}\mu=\nu} \int_{\Omega} \mathbf{c}(x, T(x)) \mu(dx)$$

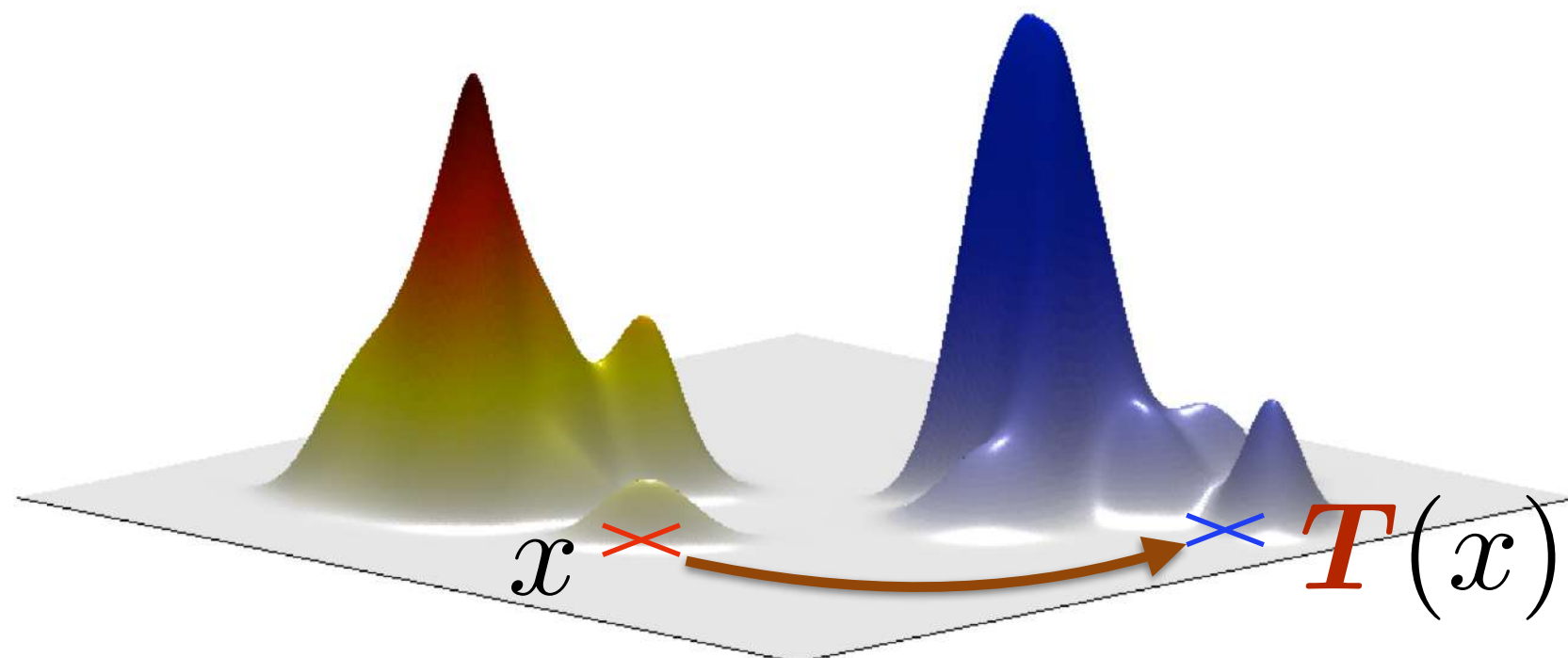


Monge Problem

Ω a measurable space, $\mathbf{c} : \Omega \times \Omega \rightarrow \mathbb{R}$.
 μ, ν two probability measures in $\mathcal{P}(\Omega)$.

[Monge'81] problem: find a map $\mathbf{T} : \Omega \rightarrow \Omega$

[Brenier'87] If $\Omega = \mathbb{R}^d$, $\mathbf{c} = \|\cdot - \cdot\|^2$,
 μ, ν a.c., then $\mathbf{T} = \nabla u$, u convex.



Monge Problem

Ω a measurable space, $\mathbf{c} : \Omega \times \Omega \rightarrow \mathbb{R}$.
 μ, ν two probability measures in $\mathcal{P}(\Omega)$.

[Monge'81] problem: find a map $\mathbf{T} : \Omega \rightarrow \Omega$

[Brenier'87] If $\Omega = \mathbb{R}^d$, $\mathbf{c} = \|\cdot - \cdot\|^2$,
 μ, ν a.c., then $\mathbf{T} = \nabla \mathbf{u}$, $\mathbf{u}$ convex.

[Brenier'87]: For any $\mathbf{u}$ convex, $\nabla \mathbf{u}$ is the OT
Monge map between μ and $\nabla \mathbf{u}_\# \mu$.

Links: Monge-Ampère Equation

If $\Omega = \mathbb{R}^d$, $\mathbf{c} = \|\cdot - \cdot\|^2$, μ, ν have densities p, q , then $T_{\#}\mu = \nu$ is equivalent to

$$p(x) = q(T(x)) |\det J_T(x)|$$

Monge-Ampère: find convex f such that

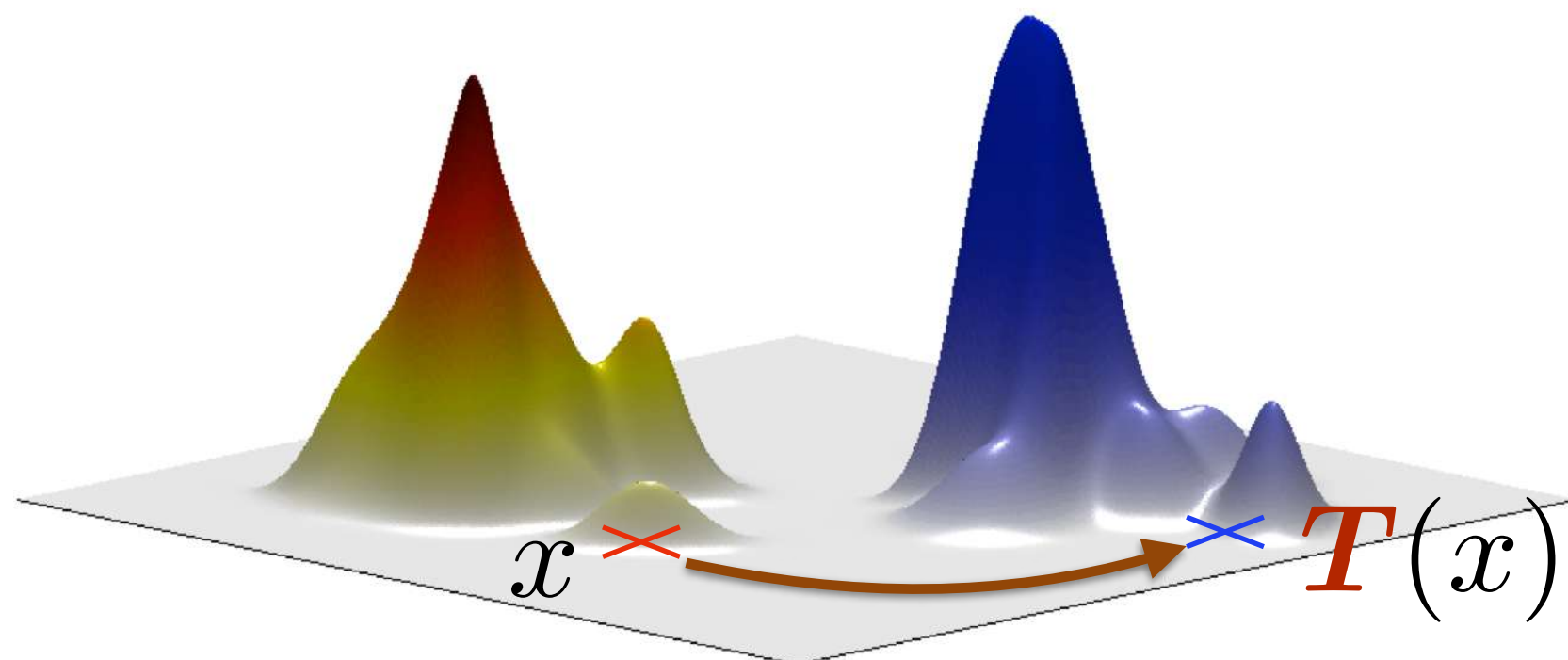
$$|\nabla^2 f(x)| = \frac{p(x)}{q(\nabla f(x))}$$

Monge Problem

Ω a measurable space, $\mathbf{c} : \Omega \times \Omega \rightarrow \mathbb{R}$.
 μ, ν two probability measures in $\mathcal{P}(\Omega)$.

[Monge'81] problem: find a map $T : \Omega \rightarrow \Omega$

$$\inf_{T_{\#}\mu=\nu} \int_{\Omega} \mathbf{c}(x, T(x)) \mu(dx)$$

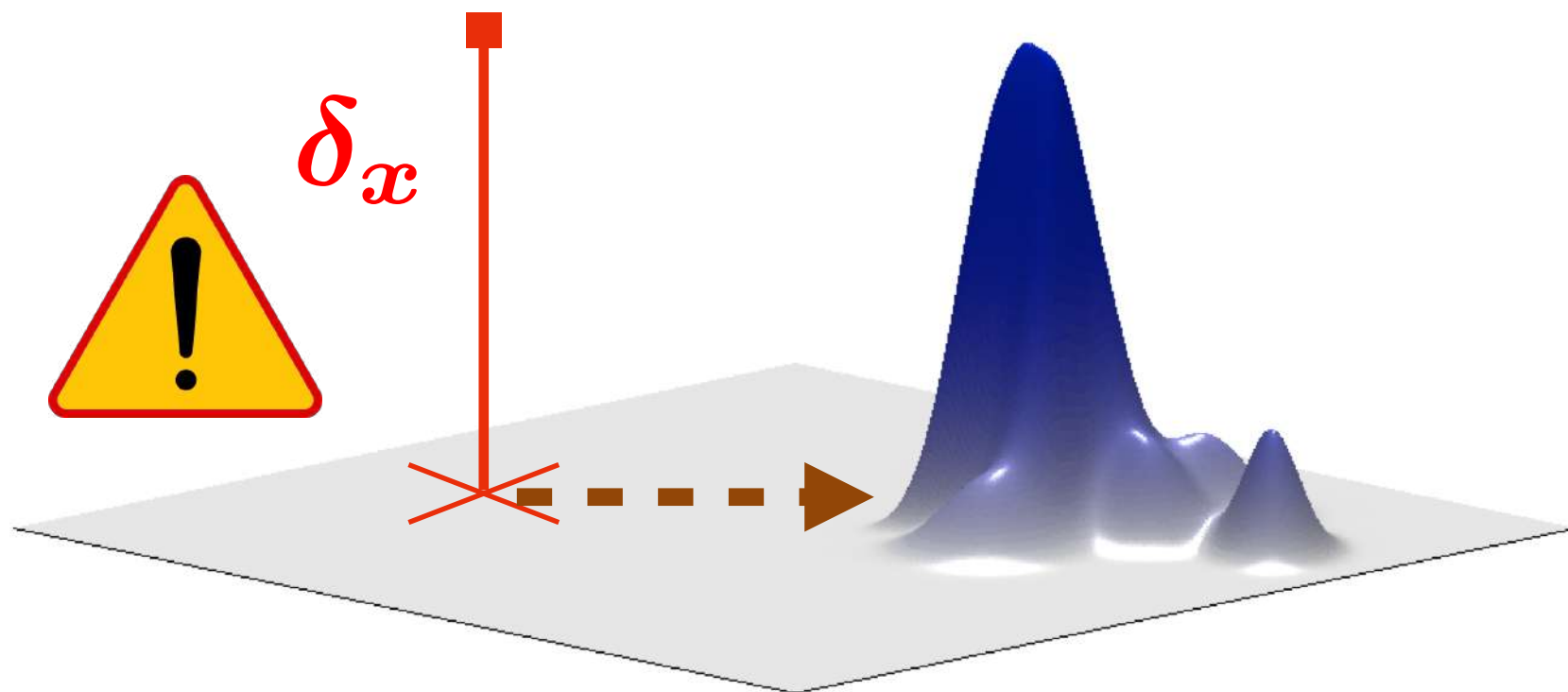


Monge Problem

Ω a measurable space, $\mathbf{c} : \Omega \times \Omega \rightarrow \mathbb{R}$.
 μ, ν two probability measures in $\mathcal{P}(\Omega)$.

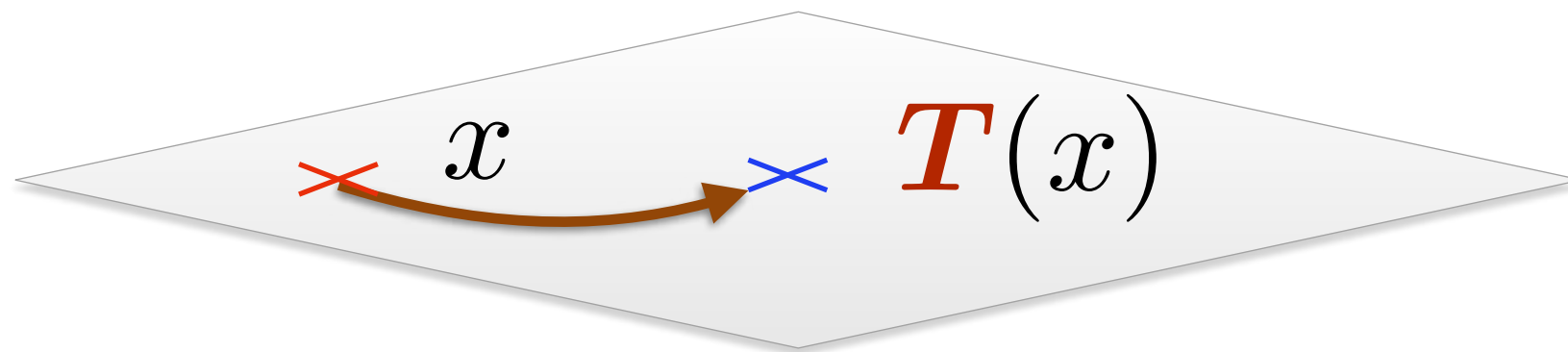
[Monge'81] problem: find a map $T : \Omega \rightarrow \Omega$

$$\inf_{T_{\#}\mu=\nu} \int_{\Omega} \mathbf{c}(x, T(x)) \mu(dx)$$



Kantorovich Relaxation

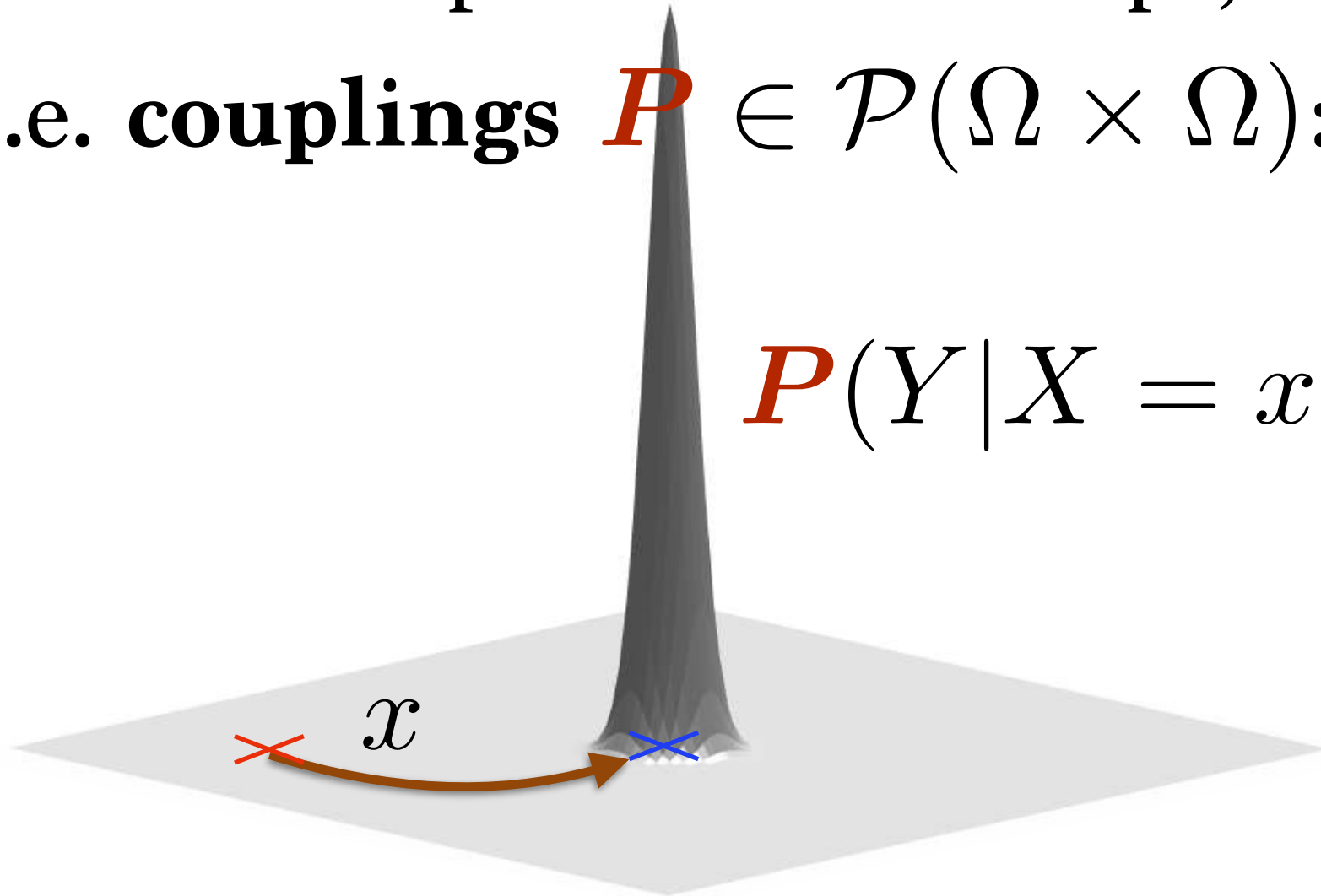
Instead of maps $T : \Omega \rightarrow \Omega$,
consider probabilistic maps,
i.e. couplings $P \in \mathcal{P}(\Omega \times \Omega)$:



Kantorovich Relaxation

Instead of maps $T : \Omega \rightarrow \Omega$,
consider probabilistic maps,
i.e. couplings $P \in \mathcal{P}(\Omega \times \Omega)$:

$$P(Y|X = x)$$



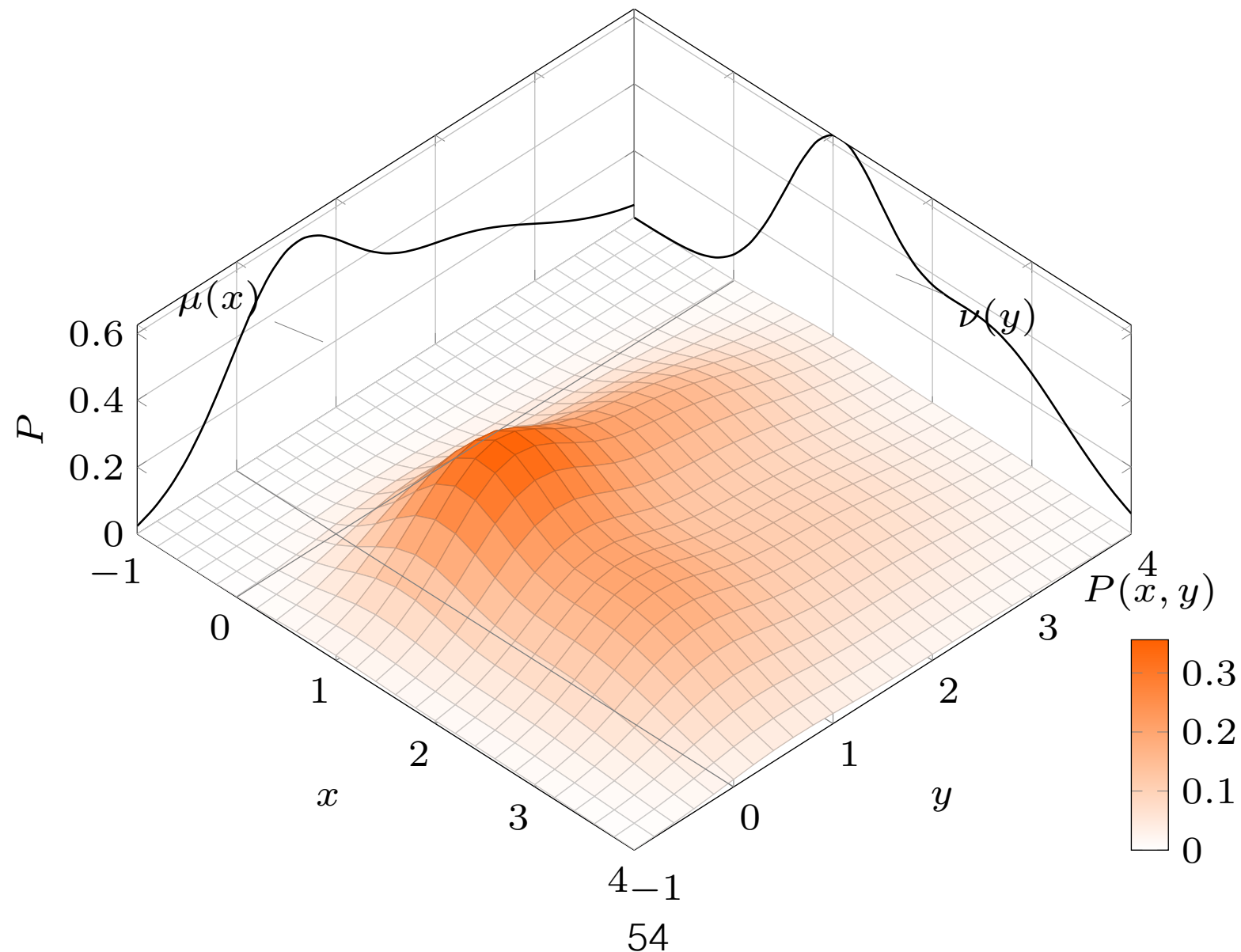
Kantorovich Relaxation

Instead of maps $T : \Omega \rightarrow \Omega$,
consider probabilistic maps,
i.e. **couplings** $P \in \mathcal{P}(\Omega \times \Omega)$:

$$\Pi(\mu, \nu) \stackrel{\text{def}}{=} \{P \in \mathcal{P}(\Omega \times \Omega) \mid \forall A, B \subset \Omega, \\ P(A \times \Omega) = \mu(A), \\ P(\Omega \times B) = \nu(B)\}$$

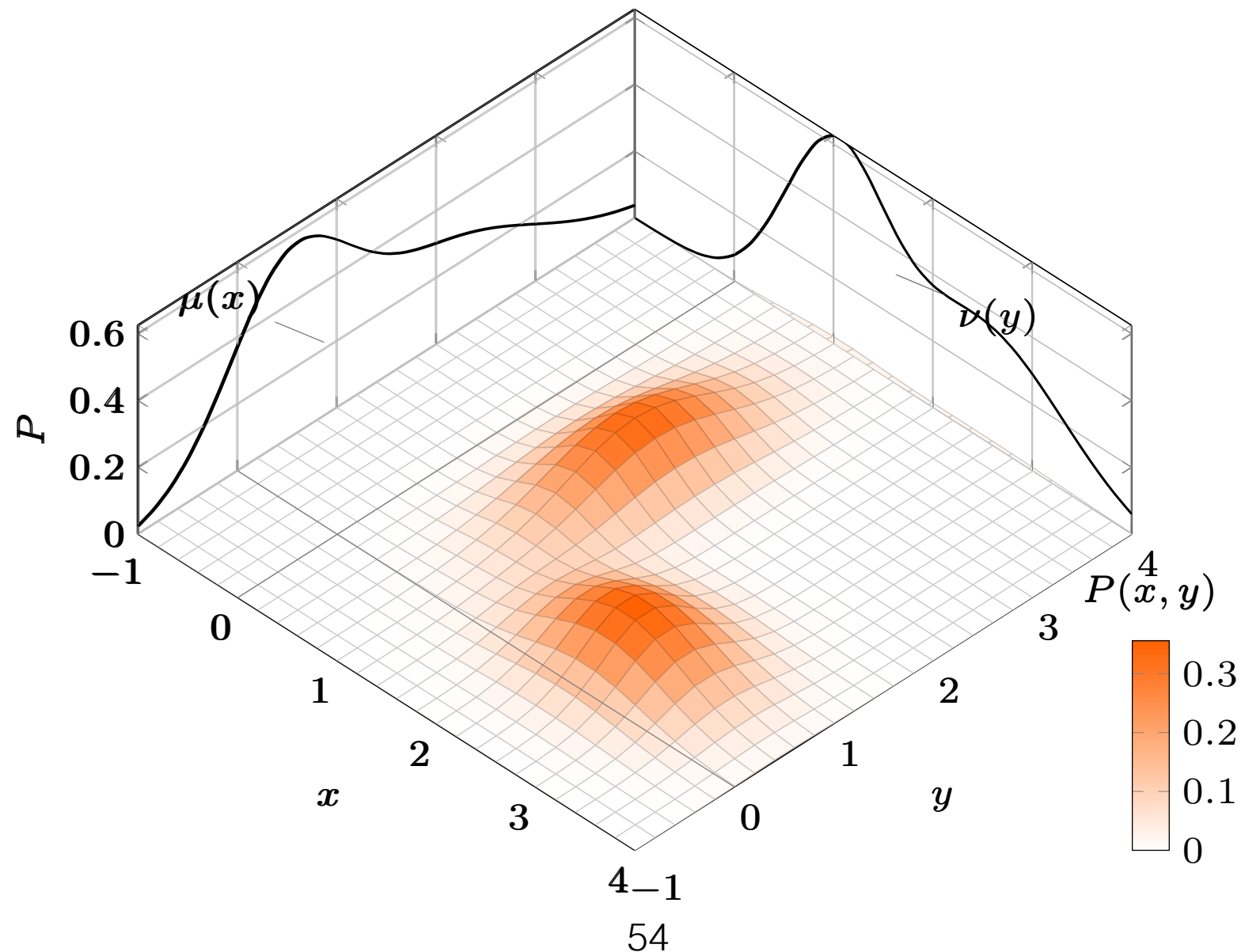
Kantorovich Relaxation

$$\Pi(\mu, \nu) \stackrel{\text{def}}{=} \{P \in \mathcal{P}(\Omega \times \Omega) \mid \forall A, B \subset \Omega, \\ P(A \times \Omega) = \mu(A), P(\Omega \times B) = \nu(B)\}$$



Kantorovich Relaxation

$$\Pi(\mu, \nu) \stackrel{\text{def}}{=} \{P \in \mathcal{P}(\Omega \times \Omega) \mid \forall A, B \subset \Omega, \\ P(A \times \Omega) = \mu(A), P(\Omega \times B) = \nu(B)\}$$



Kantorovich Problem

$$\inf_{T_{\#}\mu=\nu} \int_{\Omega} \mathbf{c}(x, T(x)) \mu(dx) \quad \text{MONGE}$$

Def. Given μ, ν in $\mathcal{P}(\Omega)$; a cost function $\mathbf{c}$ on $\Omega \times \Omega$, the Kantorovich problem is

$$\inf_{P \in \Pi(\mu, \nu)} \iint \mathbf{c}(x, y) P(dx, dy).$$

PRIMAL

Kantorovich Problem

Def. Given μ, ν in $\mathcal{P}(\Omega)$; a cost function c on $\Omega \times \Omega$, the Kantorovich problem is

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

PRIMAL

Kantorovich Problem

Def. Given μ, ν in $\mathcal{P}(\Omega)$; a cost function c on $\Omega \times \Omega$, the Kantorovich problem is

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

PRIMAL

$$\sup_{\substack{\varphi \in L_1(\mu), \psi \in L_1(\nu) \\ \varphi(x) + \psi(y) \leq c(x, y)}} \int \varphi d\mu + \int \psi d\nu.$$

DUAL

Kantorovich Problem

Def. Given μ, ν in $\mathcal{P}(\Omega)$; a cost function c on $\Omega \times \Omega$, the Kantorovich problem is

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

PRIMAL

For two real-valued functions φ, ψ on Ω ,

$$(\varphi \oplus \psi)(x, y) \stackrel{\text{def}}{=} \varphi(x) + \psi(y)$$

Kantorovich Problem

Def. Given μ, ν in $\mathcal{P}(\Omega)$; a cost function c on $\Omega \times \Omega$, the Kantorovich problem is

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

PRIMAL

$$\sup_{\substack{\varphi \in L_1(\mu), \psi \in L_1(\nu) \\ \varphi \oplus \psi \leq c}} \int \varphi d\mu + \int \psi d\nu.$$

DUAL

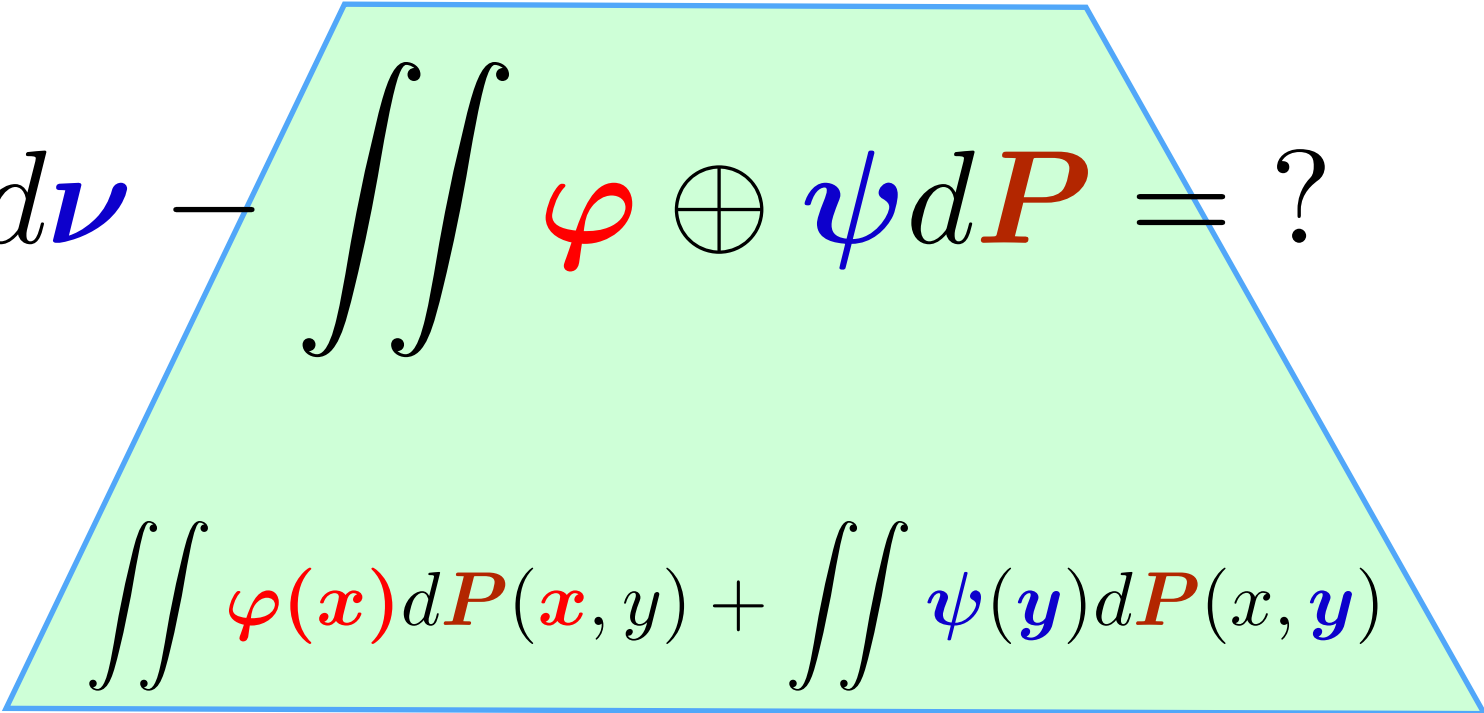
Deriving Kantorovich Duality

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \Pi(\mu, \nu)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP = ?$$

Deriving Kantorovich Duality

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \Pi(\mu, \nu)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP = ?$$

$$\iint \varphi(x) dP(x, y) + \iint \psi(y) dP(x, y)$$

Deriving Kantorovich Duality

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \Pi(\mu, \nu)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP = ?$$
$$\iint \varphi(x) dP(x, y) + \iint \psi(y) dP(x, y)$$

Deriving Kantorovich Duality

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \Pi(\mu, \nu)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP = 0$$

Deriving Kantorovich Duality

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \Pi(\mu, \nu)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP = 0$$

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \mathcal{P}_+(\Omega^2)$.

Deriving Kantorovich Duality

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \Pi(\mu, \nu)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP = 0$$

Let $\varphi, \psi : \Omega \rightarrow \mathbb{R}$, and $P \in \mathcal{P}_+(\Omega^2)$.

$$\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi dP =$$

$$\int \varphi d \underbrace{(\mu - P_X)}_{\neq 0} + \int \psi d \underbrace{(\nu - P_Y)}_{\neq 0} \quad \text{and / or}$$

Deriving Kantorovich Duality

$$\begin{aligned}\iota_{\Pi}(\boldsymbol{P}) &= \sup_{\varphi, \psi} \left[\int \varphi d\boldsymbol{\mu} + \int \psi d\boldsymbol{\nu} - \iint \varphi \oplus \psi d\boldsymbol{P} \right] \\ &= \begin{cases} 0 & \text{if } \boldsymbol{P} \in \Pi(\boldsymbol{\mu}, \boldsymbol{\nu}), \\ +\infty & \text{otherwise.} \end{cases}\end{aligned}$$

Deriving Kantorovich Duality

$$\begin{aligned}\iota_{\Pi}(\boldsymbol{P}) &= \sup_{\varphi, \psi} \left[\int \varphi d\boldsymbol{\mu} + \int \psi d\boldsymbol{\nu} - \iint \varphi \oplus \psi d\boldsymbol{P} \right] \\ &= \begin{cases} 0 & \text{if } \boldsymbol{P} \in \Pi(\boldsymbol{\mu}, \boldsymbol{\nu}), \\ +\infty & \text{otherwise.} \end{cases}\end{aligned}$$

$$\inf_{\boldsymbol{P} \in \Pi(\boldsymbol{\mu}, \boldsymbol{\nu})} \iint \boldsymbol{c} d\boldsymbol{P}$$

Deriving Kantorovich Duality

$$\begin{aligned} \iota_{\Pi}(\mathbf{P}) &= \sup_{\varphi, \psi} \left[\int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi d\mathbf{P} \right] \\ &= \begin{cases} 0 & \text{if } \mathbf{P} \in \Pi(\mu, \nu), \\ +\infty & \text{otherwise.} \end{cases} \end{aligned}$$

$$\inf_{\mathbf{P} \in \Pi(\mu, \nu)} \iint \mathbf{c} d\mathbf{P}$$

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \boxed{\iota_{\Pi}(\mathbf{P})}$$

Deriving Kantorovich Duality

$$\begin{aligned}\iota_{\Pi}(\boldsymbol{P}) &= \sup_{\varphi, \psi} \left[\int \varphi d\boldsymbol{\mu} + \int \psi d\boldsymbol{\nu} - \iint \varphi \oplus \psi d\boldsymbol{P} \right] \\ &= \begin{cases} 0 & \text{if } \boldsymbol{P} \in \Pi(\boldsymbol{\mu}, \boldsymbol{\nu}), \\ +\infty & \text{otherwise.} \end{cases}\end{aligned}$$

$$\inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \iint \boldsymbol{c} d\boldsymbol{P} + \iota_{\Pi}(\boldsymbol{P})$$

Deriving Kantorovich Duality

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \iota_\Pi(\mathbf{P})$$

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \sup_{\varphi, \psi} \int \varphi d\boldsymbol{\mu} + \int \psi d\boldsymbol{\nu} - \iint \varphi \oplus \psi d\mathbf{P}$$

Deriving Kantorovich Duality

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \iota_\Pi(\mathbf{P})$$

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \sup_{\varphi, \psi} \iint \mathbf{c} d\mathbf{P} + \int \varphi d\mu + \int \psi d\nu - \iint \varphi \oplus \psi d\mathbf{P}$$

Deriving Kantorovich Duality

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \iota_\Pi(\mathbf{P})$$

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \sup_{\varphi, \psi} \iint \mathbf{c} d\mathbf{P} - \iint \varphi \oplus \psi d\mathbf{P} + \int \varphi d\mu + \int \psi d\nu$$

Deriving Kantorovich Duality

$$\inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \iint \boldsymbol{c} \, d\boldsymbol{P} + \iota_\Pi(\boldsymbol{P})$$

$$\inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \sup_{\boldsymbol{\varphi}, \boldsymbol{\psi}} \iint (\boldsymbol{c} - \boldsymbol{\varphi} \oplus \boldsymbol{\psi}) \, d\boldsymbol{P} \quad + \int \boldsymbol{\varphi} \, d\boldsymbol{\mu} + \int \boldsymbol{\psi} \, d\boldsymbol{\nu}$$

Deriving Kantorovich Duality

$$\inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \iint \boldsymbol{c} \, d\boldsymbol{P} + \iota_\Pi(\boldsymbol{P})$$

$$\sup_{\boldsymbol{\varphi}, \boldsymbol{\psi}} \inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \iint (\boldsymbol{c} - \boldsymbol{\varphi} \oplus \boldsymbol{\psi}) \, d\boldsymbol{P} \quad + \int \boldsymbol{\varphi} \, d\boldsymbol{\mu} + \int \boldsymbol{\psi} \, d\boldsymbol{\nu}$$

Deriving Kantorovich Duality

$$\inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \iint \boldsymbol{c} \, d\boldsymbol{P} + \iota_\Pi(\boldsymbol{P})$$

$$\sup_{\boldsymbol{\varphi}, \boldsymbol{\psi}} \inf_{\boldsymbol{P} \in \mathcal{P}_+(\Omega^2)} \iint (\boldsymbol{c} - \boldsymbol{\varphi} \oplus \boldsymbol{\psi}) \, d\boldsymbol{P} \quad + \int \boldsymbol{\varphi} \, d\boldsymbol{\mu} + \int \boldsymbol{\psi} \, d\boldsymbol{\nu}$$

Deriving Kantorovich Duality

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \iota_\Pi(\mathbf{P})$$

$$\sup_{\varphi, \psi} \inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint (\mathbf{c} - \varphi \oplus \psi) d\mathbf{P} \quad + \int \varphi d\mu + \int \psi d\nu$$

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega)} \iint (\mathbf{c} - \varphi \oplus \psi) d\mathbf{P} = \begin{cases} 0 & \text{if } \mathbf{c} - \varphi \oplus \psi \geq 0. \\ -\infty & \text{otherwise} \end{cases}$$

Deriving Kantorovich Duality

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint \mathbf{c} d\mathbf{P} + \iota_\Pi(\mathbf{P})$$

$$\sup_{\varphi, \psi} \inf_{\mathbf{P} \in \mathcal{P}_+(\Omega^2)} \iint (\mathbf{c} - \varphi \oplus \psi) d\mathbf{P} \quad + \int \varphi d\mu + \int \psi d\nu$$

$$\inf_{\mathbf{P} \in \mathcal{P}_+(\Omega)} \iint (\mathbf{c} - \varphi \oplus \psi) d\mathbf{P} = \begin{cases} 0 & \text{if } \mathbf{c} - \varphi \oplus \psi \geq 0. \\ -\infty & \text{otherwise} \end{cases}$$

$$\sup_{\varphi \oplus \psi \leq \mathbf{c}} \int \varphi d\mu + \int \psi d\nu.$$

DUAL

Wasserstein Distances

Let $p \geq 1$. Let $\mathbf{c}(x, y) := \mathbf{D}^p(x, y)$, a metric.

Def. The p -Wasserstein distance between μ, ν in $\mathcal{P}(\Omega)$ is

$$W_p(\mu, \nu) \stackrel{\text{def}}{=} \left(\inf_{P \in \Pi(\mu, \nu)} \iint \mathbf{D}(x, y)^p P(dx, dy) \right)^{1/p}.$$

Wasserstein Distances

Let $p \geq 1$. Let $\mathbf{c}(x, y) := \mathbf{D}^p(x, y)$, a metric.

Def. The p -Wasserstein distance between μ, ν in $\mathcal{P}(\Omega)$ is

$$W_p(\mu, \nu) \stackrel{\text{def}}{=} \left(\inf_{P \in \Pi(\mu, \nu)} \iint \mathbf{D}(x, y)^p P(dx, dy) \right)^{1/p}.$$

Kantorovich Duality

$$W_p^p(\mu, \nu) = \sup_{\substack{\varphi \in L_1(\mu), \psi \in L_1(\nu) \\ \varphi(x) + \psi(y) \leq D^p(x, y)}} \int \varphi d\mu + \int \psi d\nu.$$

DUAL

- Kantorovich duality is computationally easier: easier to store 2 functions than an entire coupling.
- D transforms: useful machinery to consider only **two** instead of **one** dual potential, notably when $p = 1$

$$W_1(\mu, \nu) = \sup_{\varphi \text{ 1-Lipschitz}} \int \varphi (d\mu - d\nu).$$

Links between Monge & Kantorovich

Prop. For “well behaved” costs c , if μ has a density then an *optimal* Monge map T^* between μ and ν must exist.

Prop. In that case

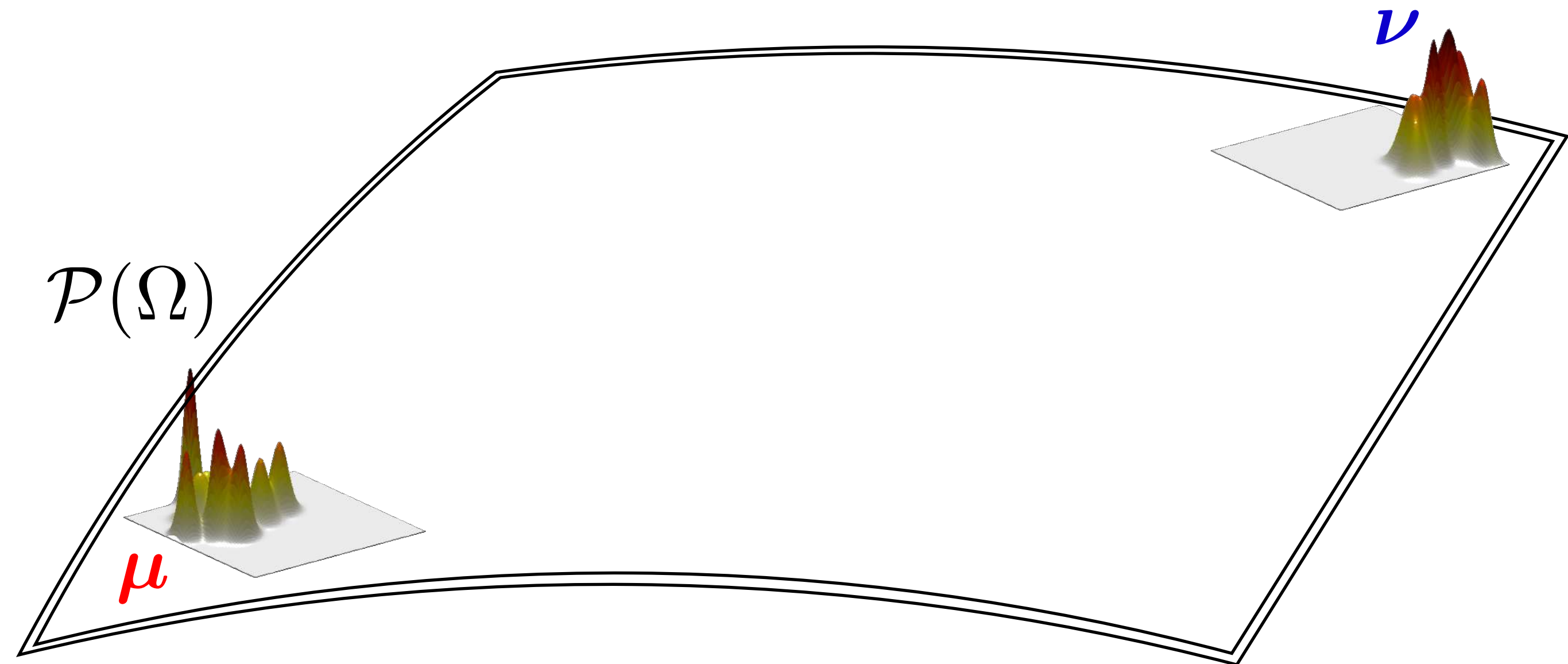
$$P^* := (\text{Id}, T^*)_{\#} \mu \in \Pi(\mu, \nu)$$

is also *optimal* for the Kantorovich problem.

[Brenier’91] [Smith&Knott’87] [McCann’01]

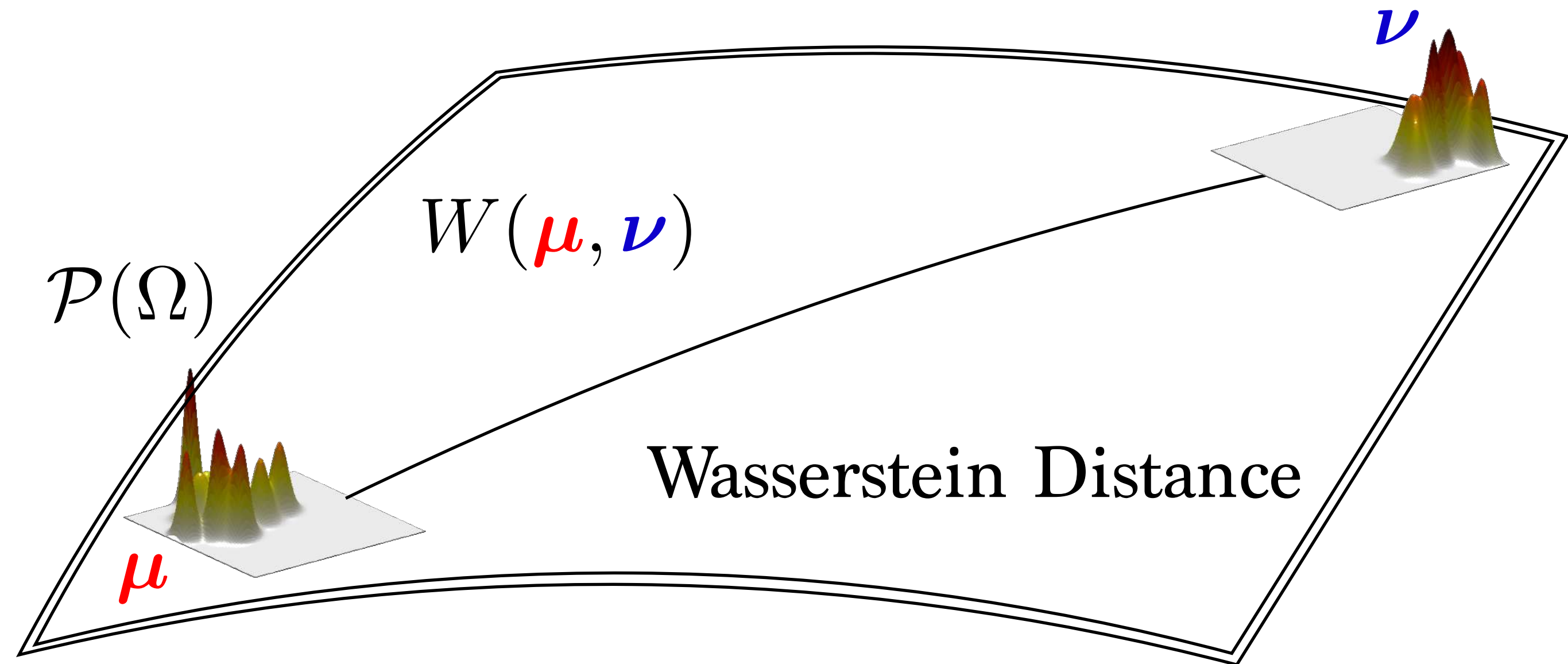
Optimal Transport Geometry

Very different geometry than standard information divergences (*KL*, Euclidean)



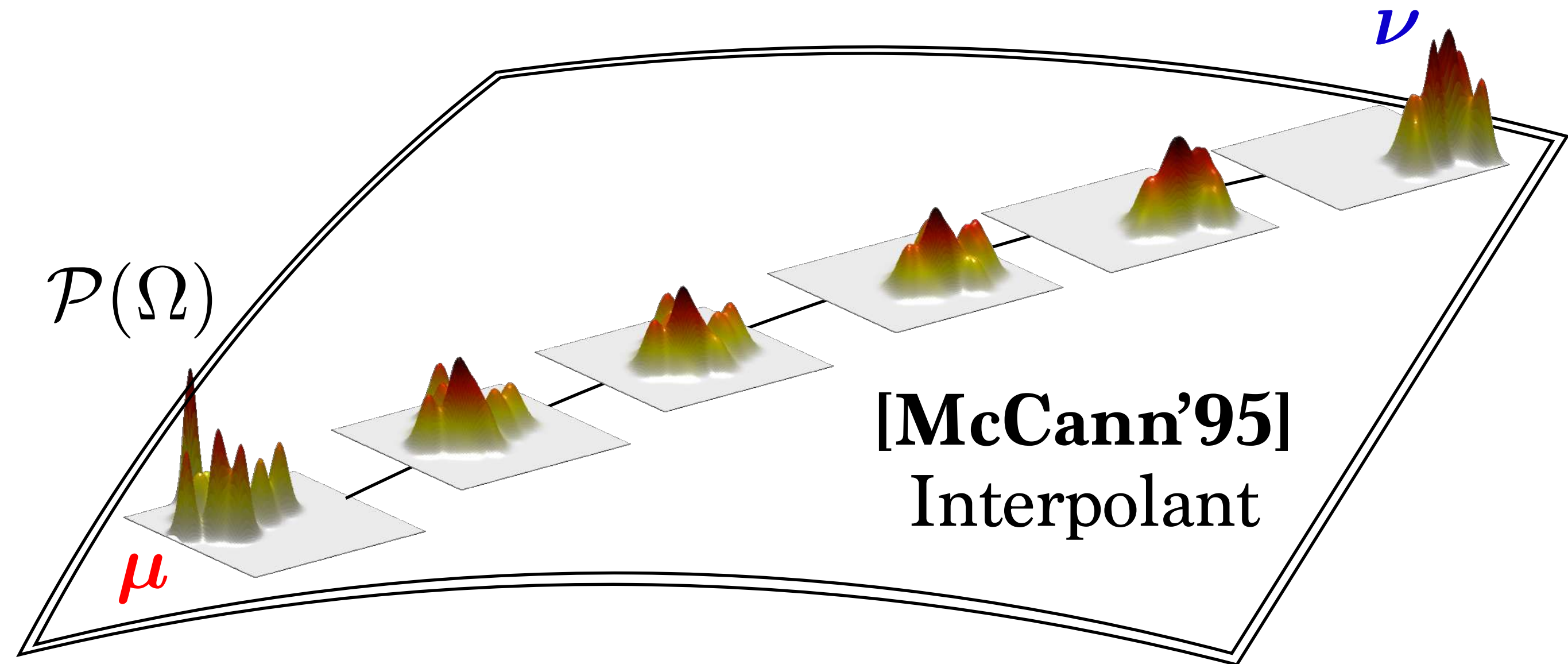
Optimal Transport Geometry

Very different geometry than standard information divergences (KL , Euclidean)



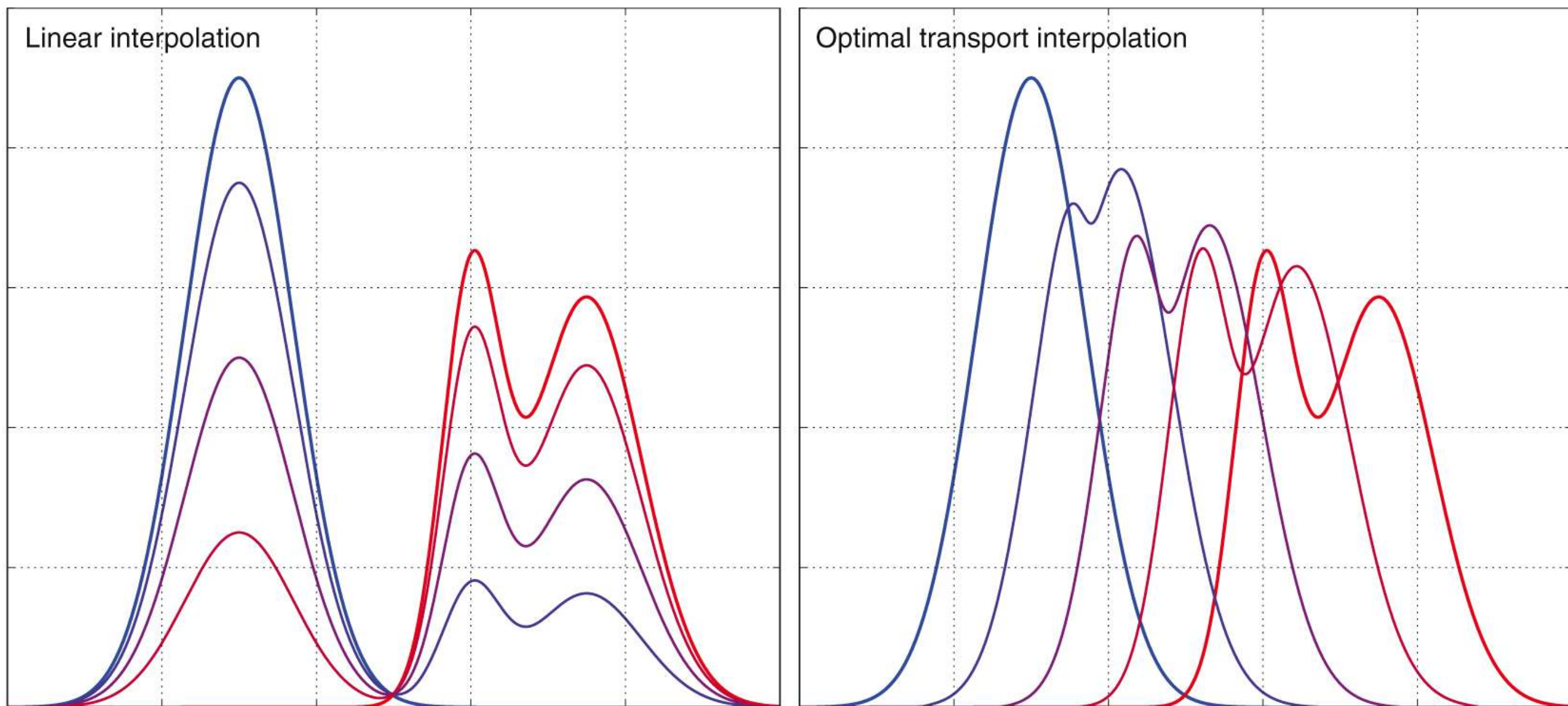
Optimal Transport Geometry

Very different geometry than standard information divergences (*KL*, Euclidean)

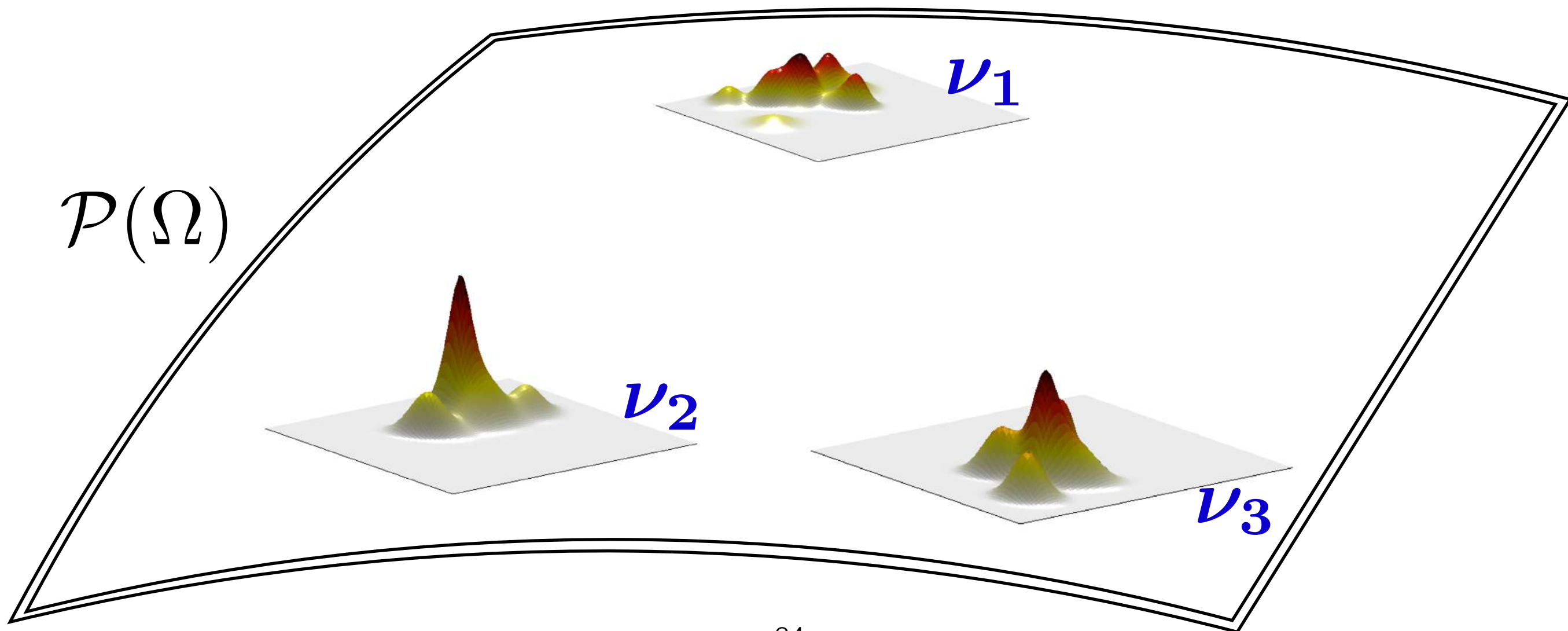


Optimal Transport Geometry

Very different geometry than standard information divergences (KL , Euclidean)

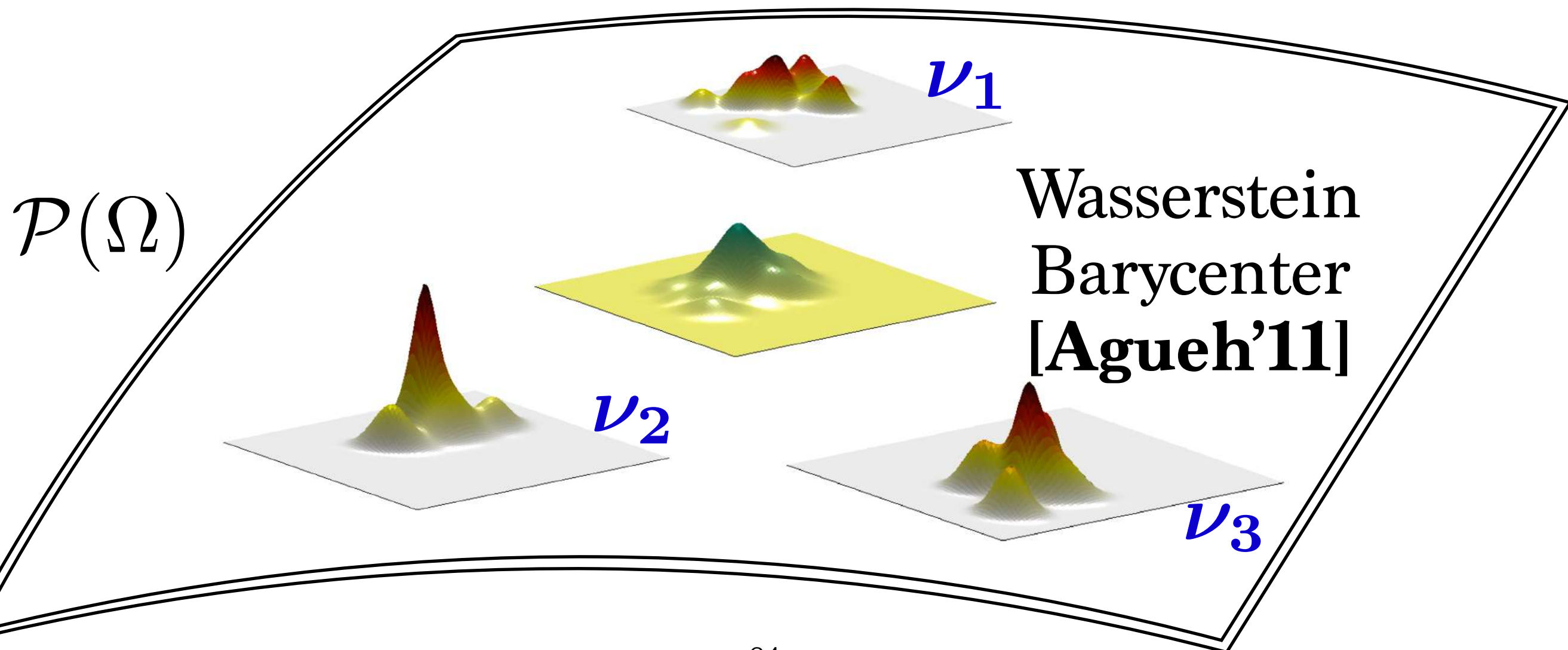


Optimal Transport Geometry



Optimal Transport Geometry

$$\min_{\mu \in \mathcal{P}(\Omega)} \sum_{i=1}^N \lambda_i W_p^p(\mu, \nu_i)$$



Optimal Transport Geometry

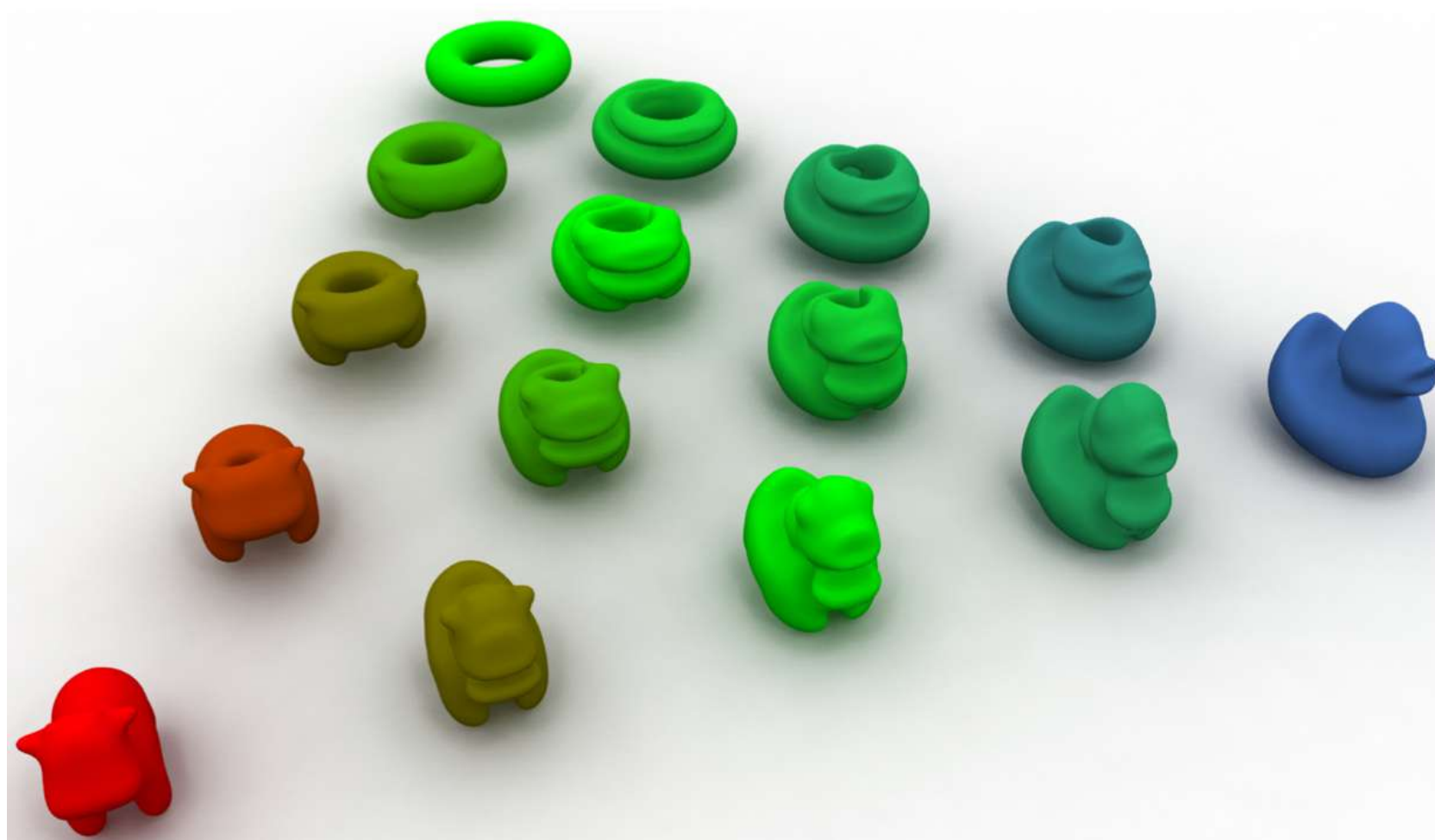
Very different geometry than standard information divergences (KL , Euclidean)



[SDPC.'15]

Optimal Transport Geometry

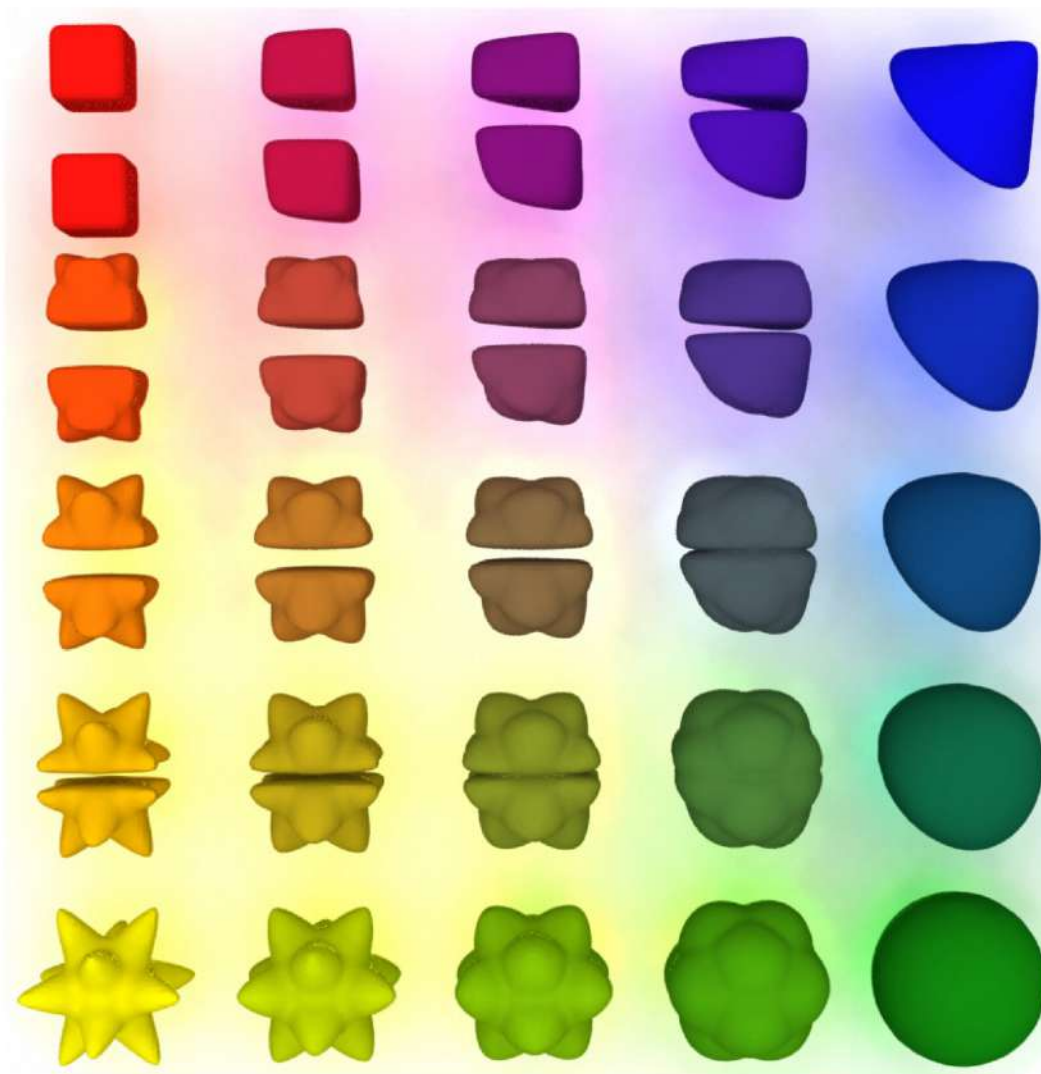
Very different geometry than standard information divergences (KL , Euclidean)



[SDPC.'15]

Optimal Transport Geometry

Very different geometry than standard information divergences (KL , Euclidean)



[SDPC.'15]

Variational OT Problems in ML

Up to 2010: OT solvers
used mostly for retrieval
in databases of histograms

$$W_p(\mu, \nu) = ?$$

$$W_p(\mu, \nu) \leq \dots ?$$

OT is now used as a **loss** or **fidelity** term:

$$\operatorname{argmin}_{\mu \in \mathcal{P}(\Omega)} F(W_p(\mu, \nu_1), W_p(\mu, \nu_2), \dots, \mu) = ?$$

[Jordan Kinderlehrer Otto'98]

$$“\nabla_{\mu}” W_p(\mu, \nu) = ?$$

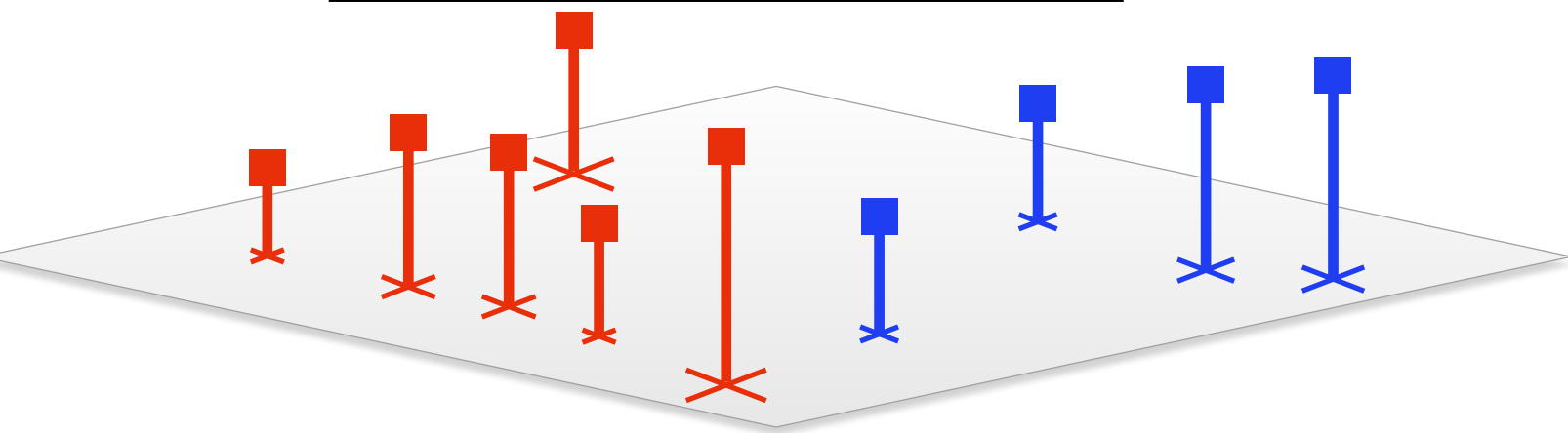
[Ambrosio Gigli Savaré'05]

2. Computing OT exactly

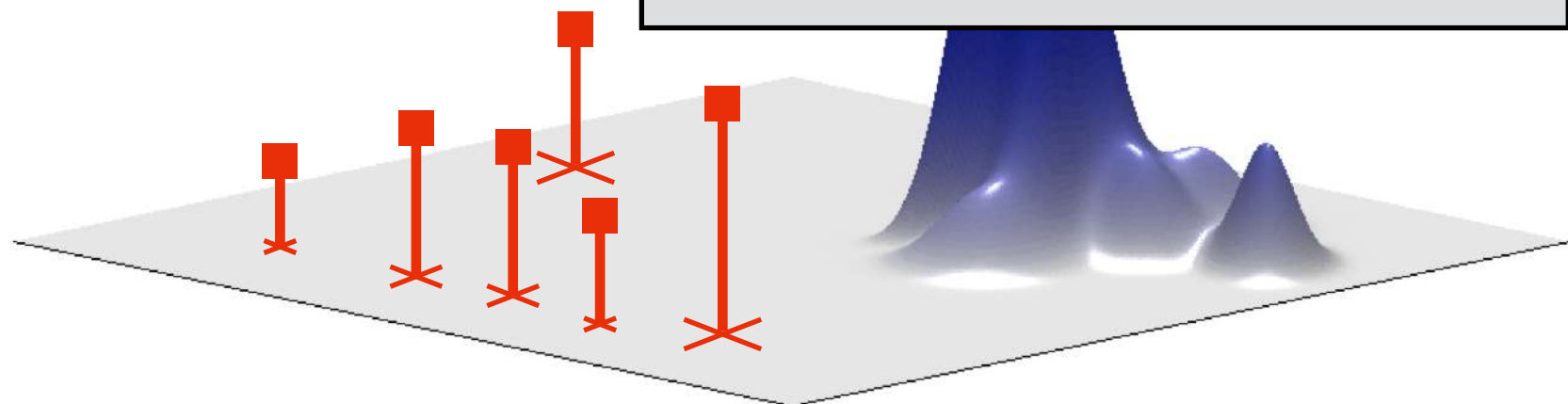
- Typology: discrete/continuous problems
- Easy cases and exact solvers for discrete measures.

When can we compute OT?

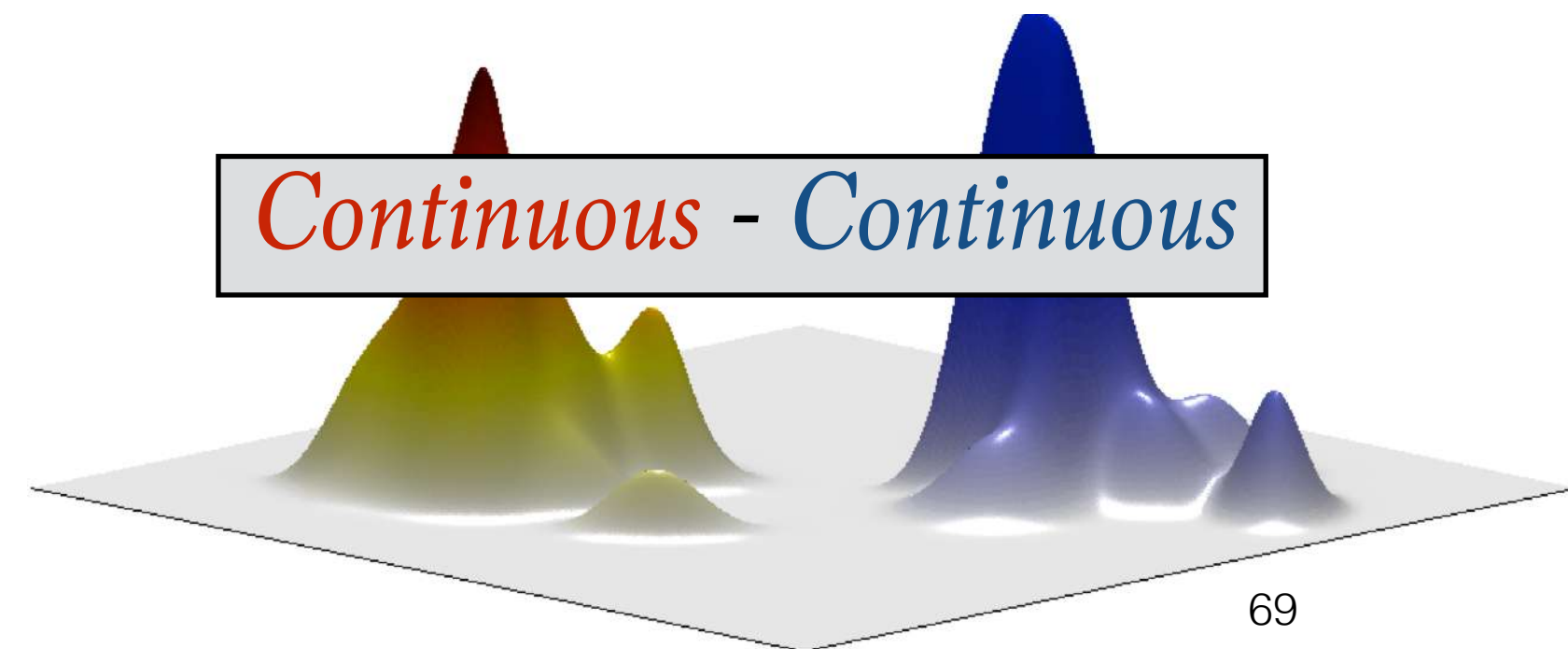
Discrete - Discrete



Discrete - Continuous



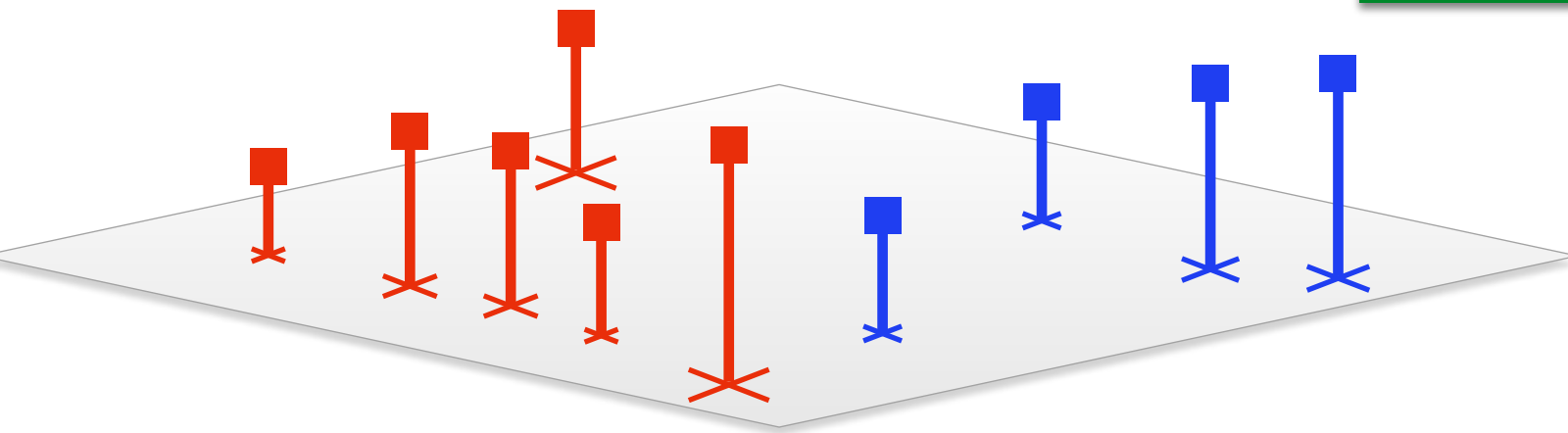
Continuous - Continuous



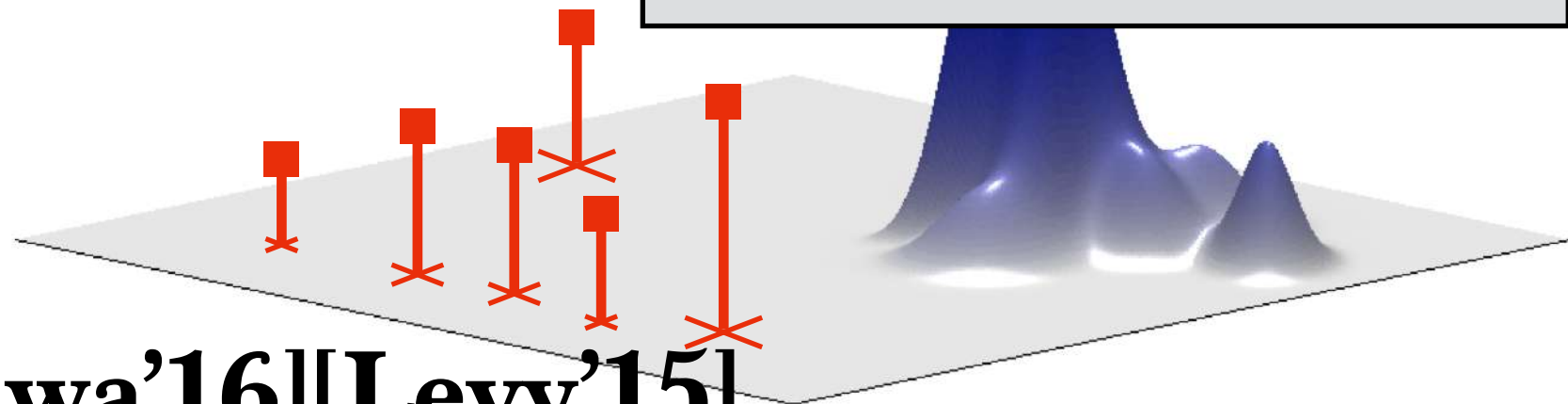
When can we compute OT?

Discrete - *Discrete*

Network flow solvers



Discrete - *Continuous*



low dim.

[Mérigot'11][Kitagawa'16][Levy'15]

Continuous - *Continuous*

Stochastic
Optimization

[Genevay'16]

[Arjovsky'17]

PDE's

[Benamou'98]

Easy (1): Univariate Measures

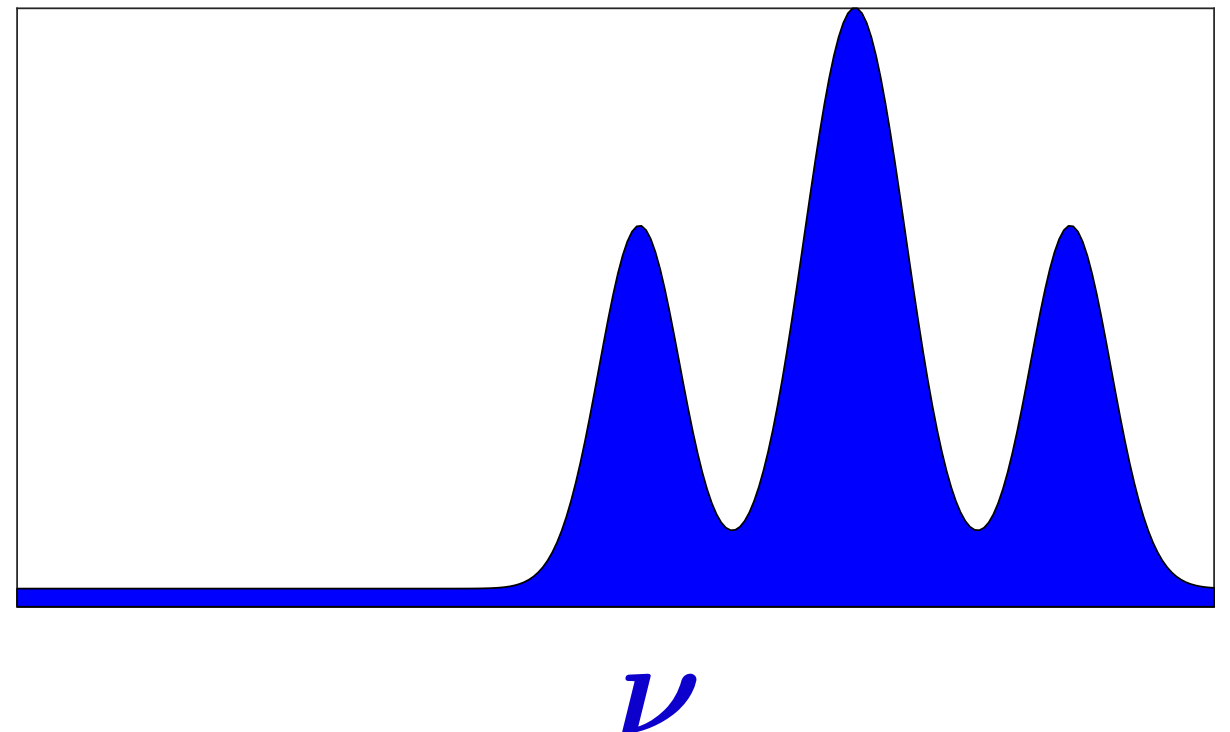
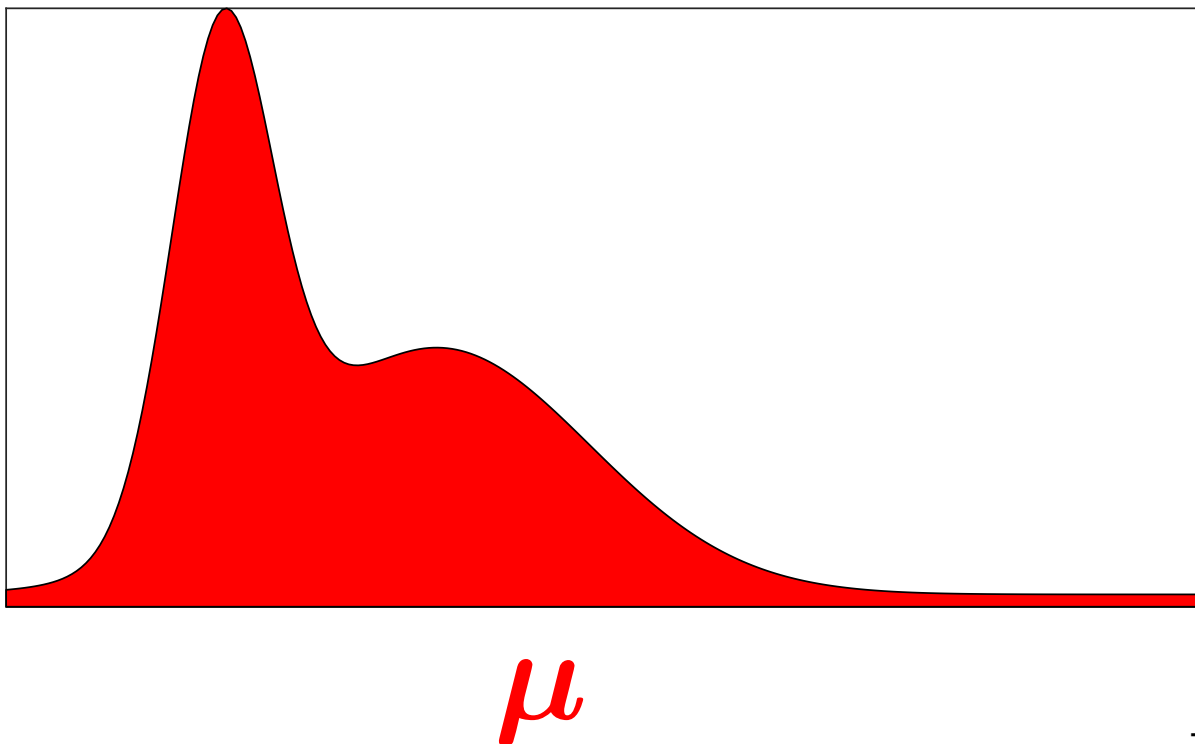
Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$

Easy (1): Univariate Measures

Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

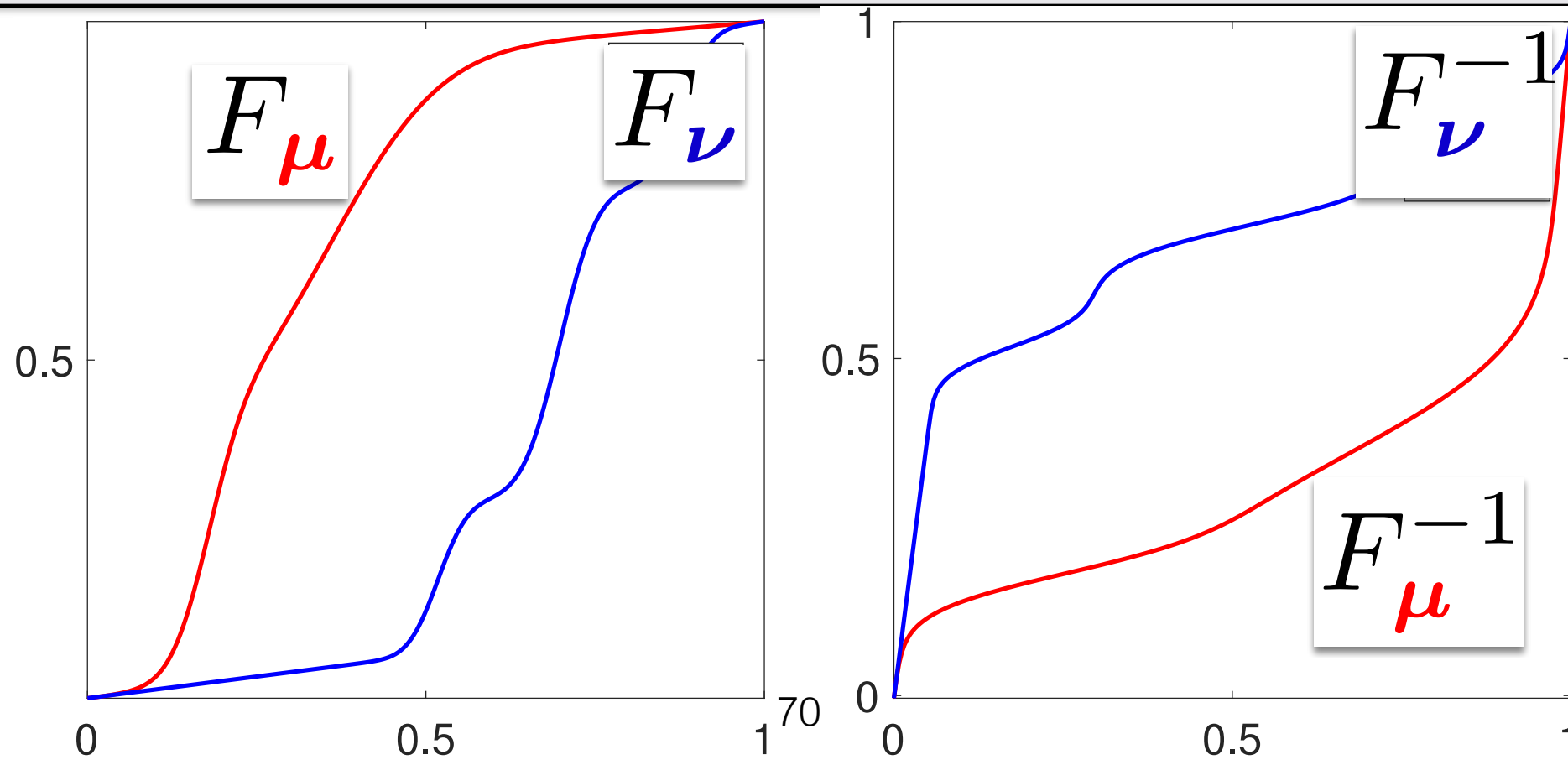
$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$



Easy (1): Univariate Measures

Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

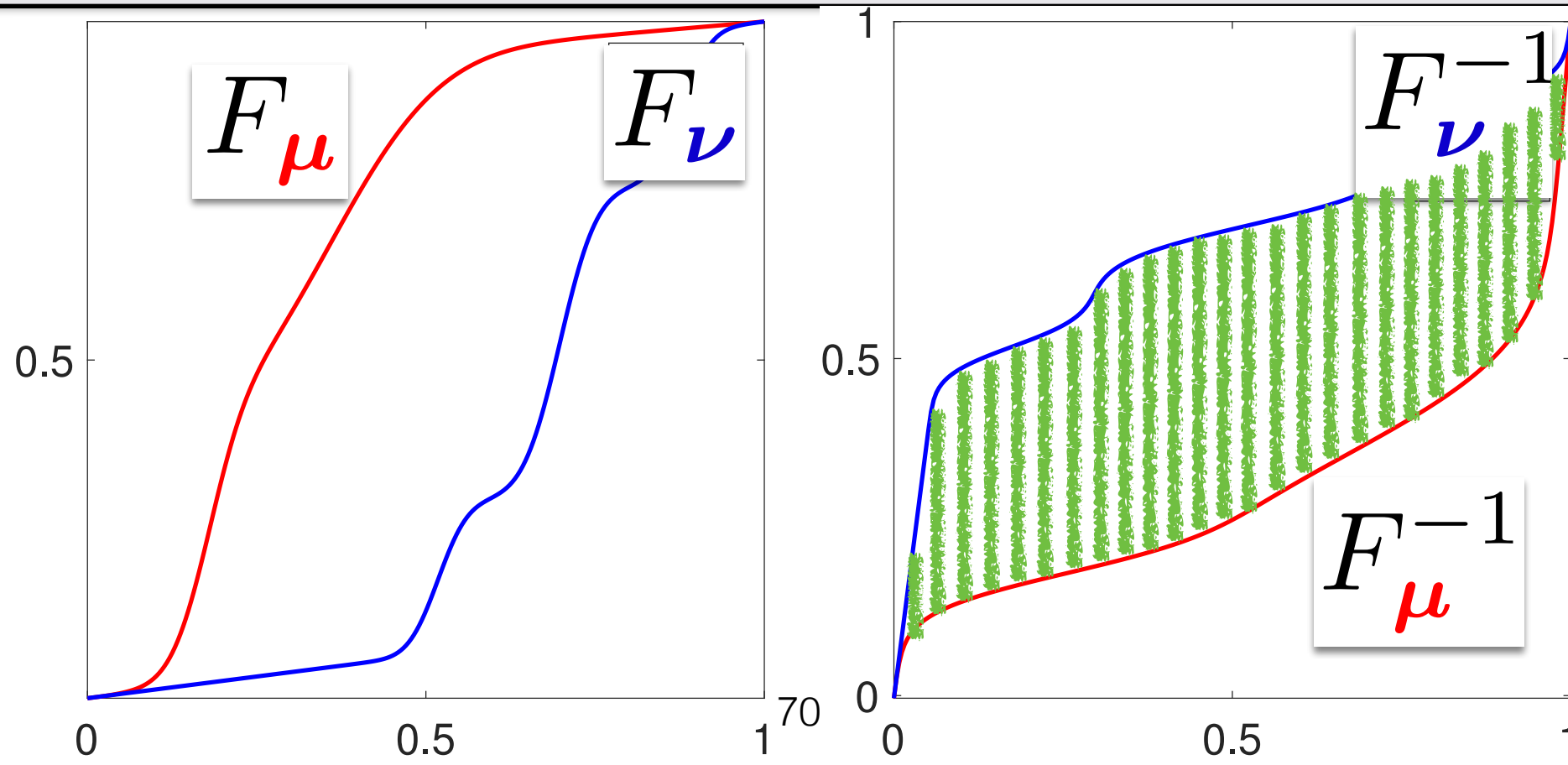
$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$



Easy (1): Univariate Measures

Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$



Easy (2): Gaussian Measures

Remark. If $\Omega = \mathbb{R}^d$, $\mathbf{c}(x, y) = \|x - y\|^2$, and $\mu = \mathcal{N}(\mathbf{m}_\mu, \Sigma_\mu)$, $\nu = \mathcal{N}(\mathbf{m}_\nu, \Sigma_\nu)$ then

$$W_2^2(\mu, \nu) = \|\mathbf{m}_\mu - \mathbf{m}_\nu\|^2 + B(\Sigma_\mu, \Sigma_\nu)^2$$

where B is the Bures metric

$$B(\Sigma_\mu, \Sigma_\nu)^2 = \text{trace}(\Sigma_\mu + \Sigma_\nu - 2(\Sigma_\mu^{1/2} \Sigma_\nu \Sigma_\mu^{1/2})^{1/2}).$$

Easy (2): Gaussian Measures

Remark. If $\Omega = \mathbb{R}^d$, $\mathbf{c}(x, y) = \|x - y\|^2$, and $\mu = \mathcal{N}(\mathbf{m}_\mu, \Sigma_\mu)$, $\nu = \mathcal{N}(\mathbf{m}_\nu, \Sigma_\nu)$ then

$$W_2^2(\mu, \nu) = \|\mathbf{m}_\mu - \mathbf{m}_\nu\|^2 + B(\Sigma_\mu, \Sigma_\nu)^2$$

where B is the Bures metric

$$B(\Sigma_\mu, \Sigma_\nu)^2 = \text{trace}(\Sigma_\mu + \Sigma_\nu - 2(\Sigma_\mu^{1/2} \Sigma_\nu \Sigma_\mu^{1/2})^{1/2}).$$

The map $T : x \mapsto \mathbf{m}_\nu + A(x - \mathbf{m}_\mu)$ is **optimal**,

$$\text{where } A = \Sigma_\mu^{-\frac{1}{2}} \left(\Sigma_\mu^{\frac{1}{2}} \Sigma_\nu \Sigma_\mu^{\frac{1}{2}} \right)^{\frac{1}{2}} \Sigma_\mu^{-\frac{1}{2}}.$$

Easy (2): Gaussian Measures

Remark. If $\Omega = \mathbb{R}^d$, $\mathbf{c}(x, y) = \|x - y\|^2$, and $\mu = \mathcal{N}(\mathbf{m}_\mu, \Sigma_\mu)$, $\nu = \mathcal{N}(\mathbf{m}_\nu, \Sigma_\nu)$ then

$$W_2^2(\mu, \nu) = \|\mathbf{m}_\mu - \mathbf{m}_\nu\|^2 + B(\Sigma_\mu, \Sigma_\nu)^2$$

where B is the Bures distance, defined by

$$B(\Sigma_\mu, \Sigma_\nu) = \sqrt{\frac{1}{2} \text{trace}(\Sigma_\mu + \Sigma_\nu - 2(\Sigma_\mu^{1/2} \Sigma_\nu \Sigma_\mu^{1/2})^{1/2})}.$$

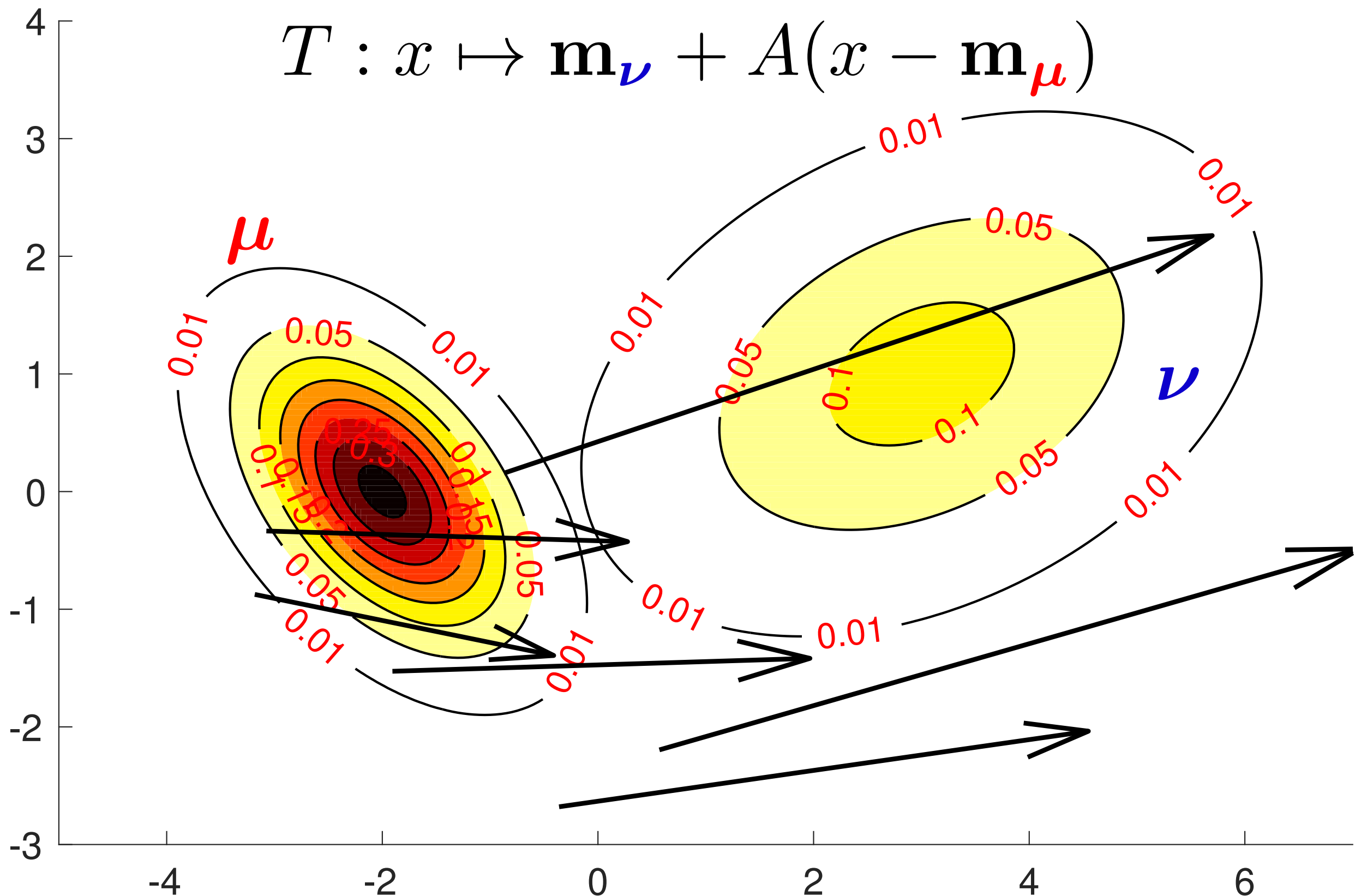
if $X \sim \mathcal{N}(m, \Sigma)$ then $Y := CX + b \sim \mathcal{N}(Cm + b, C\Sigma C^T)$.

The map $T : x \mapsto \mathbf{m}_\nu + A(x - \mathbf{m}_\mu)$ is optimal,

$$\text{where } A = \Sigma_\mu^{-\frac{1}{2}} \left(\Sigma_\mu^{\frac{1}{2}} \Sigma_\nu \Sigma_\mu^{\frac{1}{2}} \right)^{\frac{1}{2}} \Sigma_\mu^{-\frac{1}{2}}.$$

Easy (2): Gaussian Measures

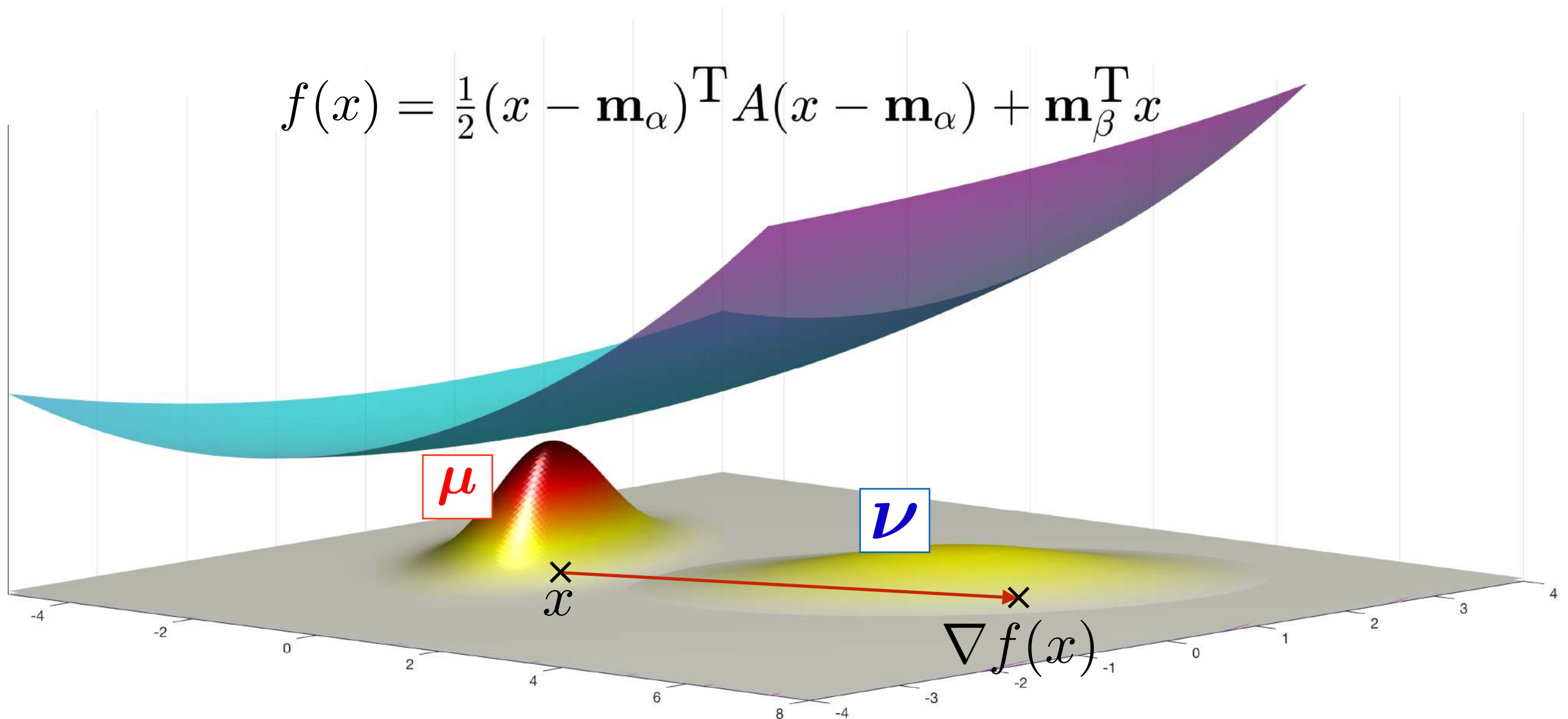
$$T : x \mapsto \mathbf{m}_\nu + A(x - \mathbf{m}_\mu)$$



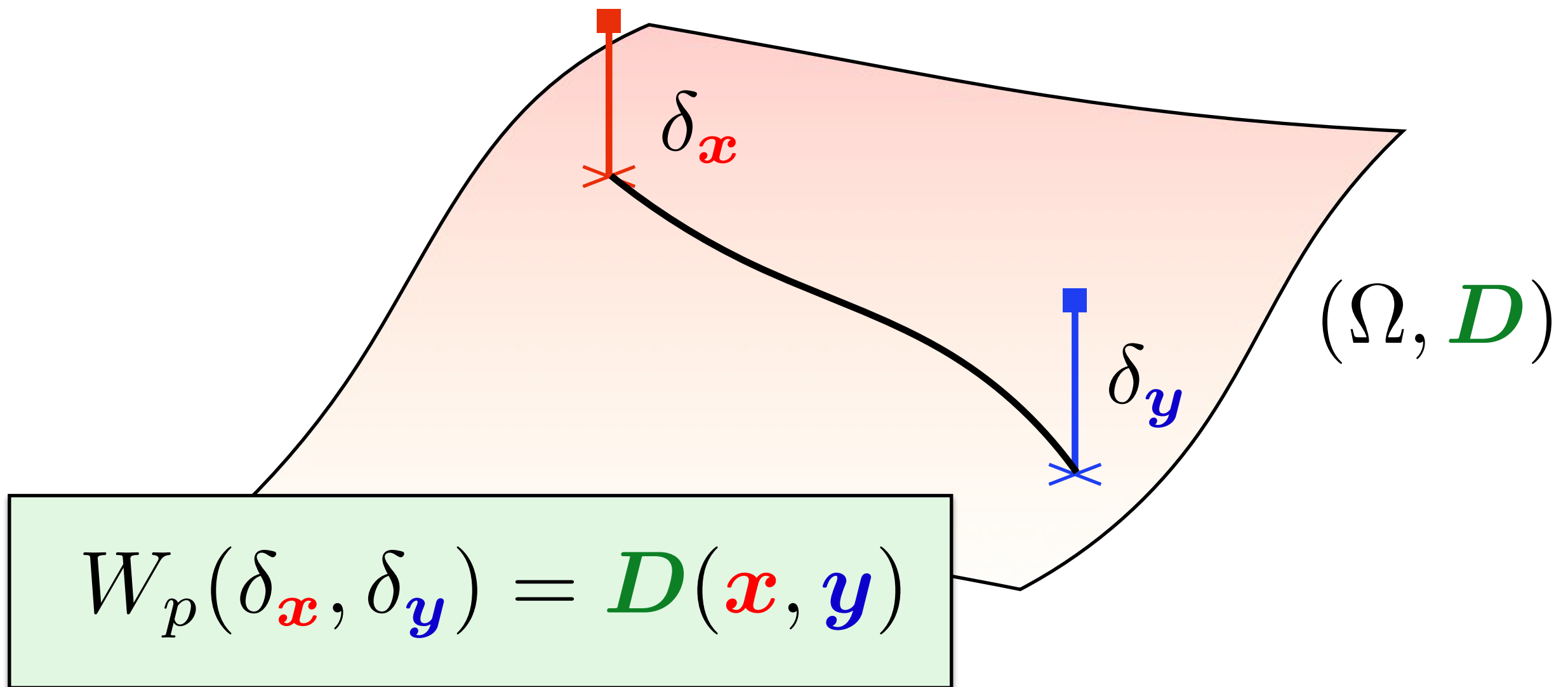
Easy (2): Gaussian Measures

$$T = \nabla \psi : x \mapsto \mathbf{m}_{\nu} + A(x - \mathbf{m}_{\mu})$$

$$f(x) = \frac{1}{2}(x - \mathbf{m}_{\alpha})^T A(x - \mathbf{m}_{\alpha}) + \mathbf{m}_{\beta}^T x$$



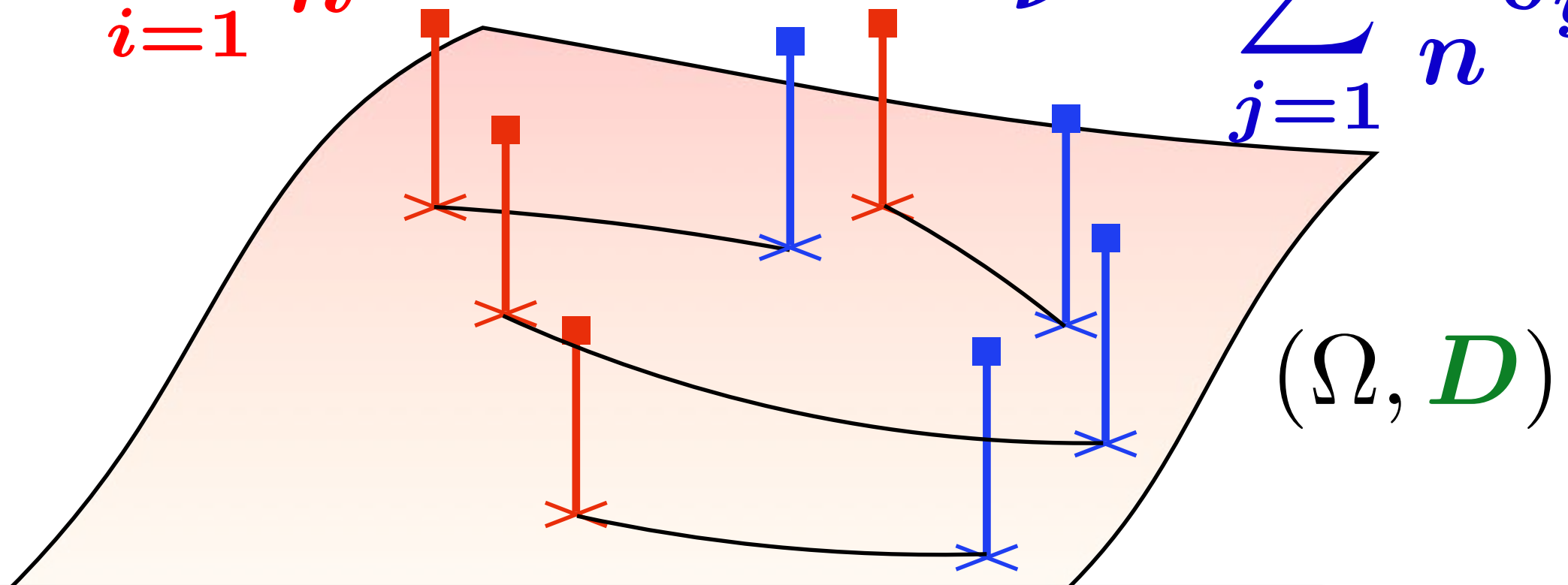
Wasserstein Between Two Diracs



Linear Assignment $\subset$ Wasserstein

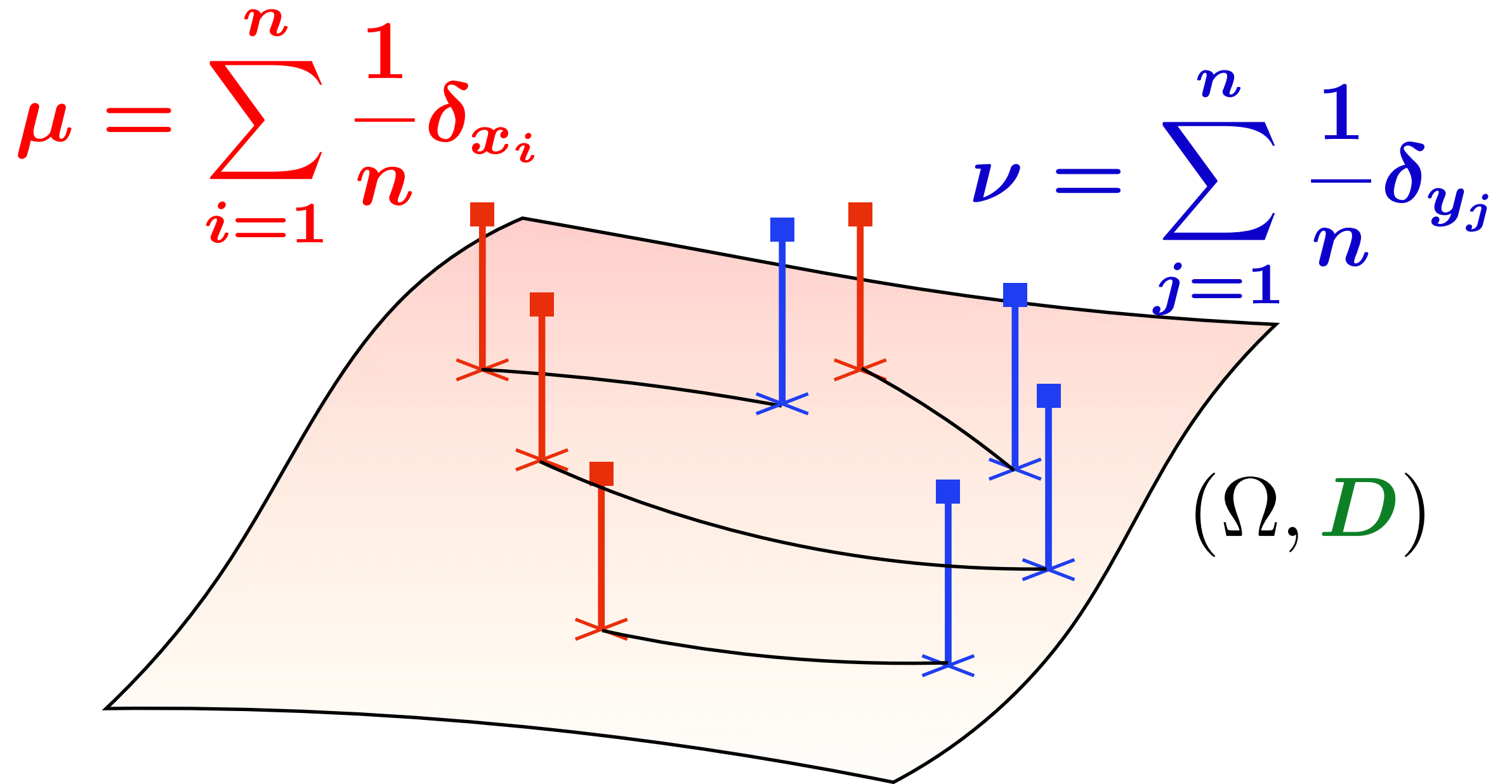
$$\mu = \sum_{i=1}^n \frac{1}{n} \delta_{x_i}$$

$$\nu = \sum_{j=1}^n \frac{1}{n} \delta_{y_j}$$



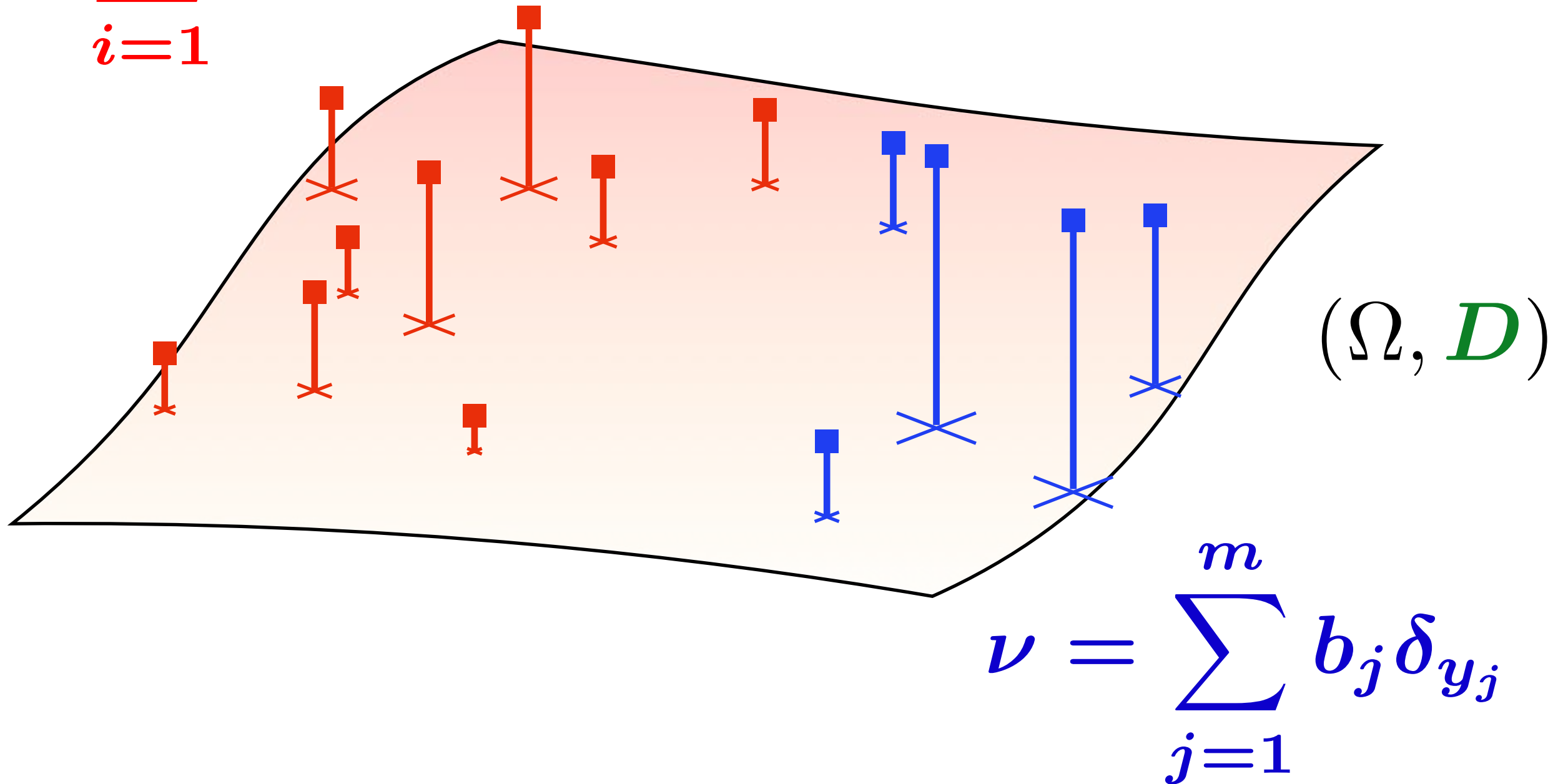
$$W_p^p(\mu, \nu) = \min_{\sigma \in S_n} \frac{1}{n} \sum_{i=1}^n D(x_i, y_{\sigma(i)})^p$$

Linear Assignment $\subset$ Wasserstein



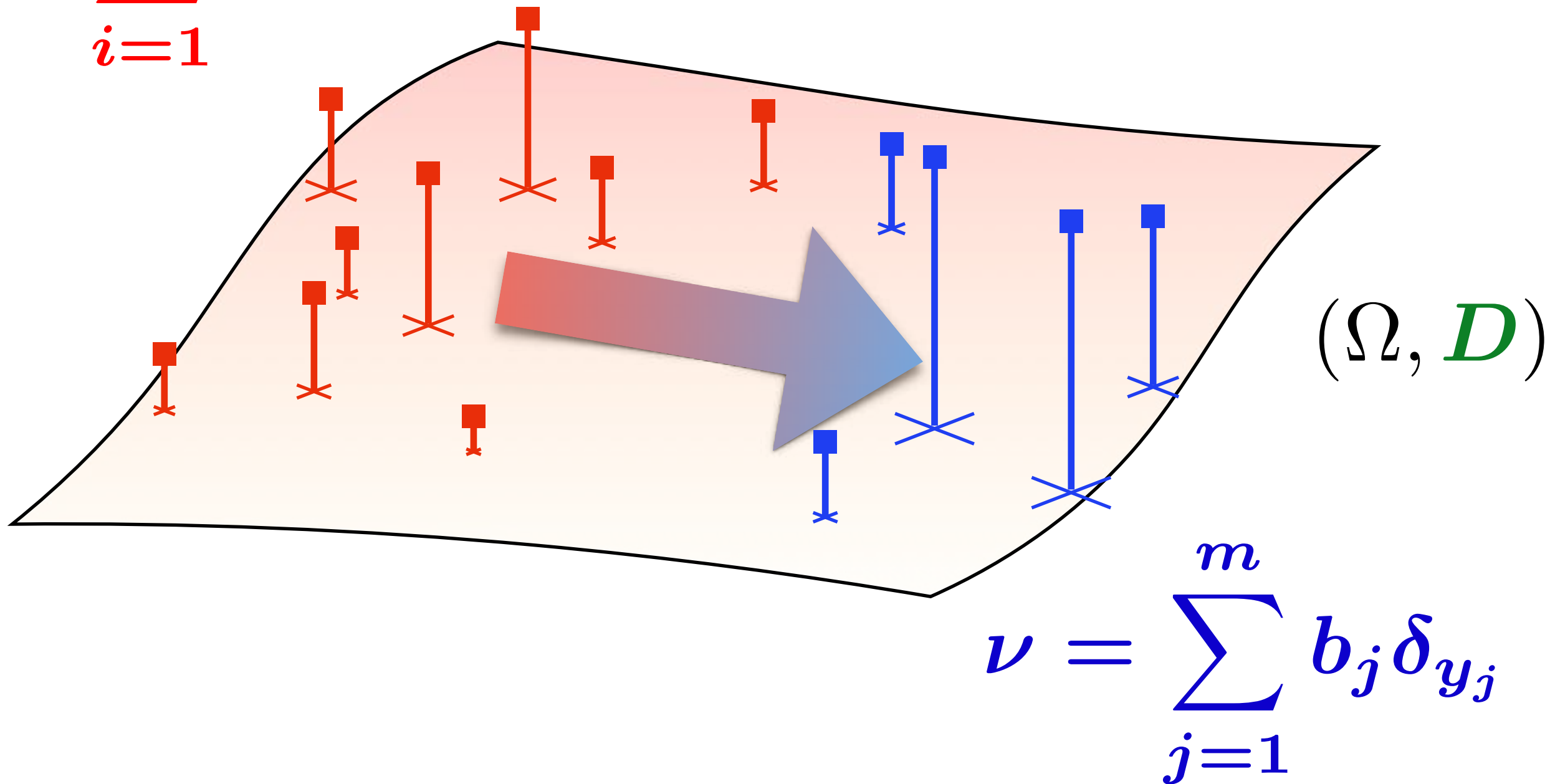
OT on Two Empirical Measures

$$\mu = \sum_{i=1}^n a_i \delta_{x_i}$$



OT on Two Empirical Measures

$$\mu = \sum_{i=1}^n a_i \delta_{x_i}$$



Wasserstein on Empirical Measures

Consider $\mu = \sum_{i=1}^n a_i \delta_{x_i}$ and $\nu = \sum_{j=1}^m b_j \delta_{y_j}$.

$$M_{\mathbf{x}\mathbf{y}} \stackrel{\text{def}}{=} [D(\mathbf{x}_i, \mathbf{y}_j)^p]_{ij}$$

$$U(\mathbf{a}, \mathbf{b}) \stackrel{\text{def}}{=} \{ \mathbf{P} \in \mathbb{R}_+^{n \times m} \mid \mathbf{P} \mathbf{1}_m = \mathbf{a}, \mathbf{P}^T \mathbf{1}_n = \mathbf{b} \}$$

Def. Optimal Transport Problem

$$W_p^p(\mu, \nu) = \min_{\mathbf{P} \in U(\mathbf{a}, \mathbf{b})} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle$$

Dual Kantorovich Problem

$$W_p^p(\mu, \nu) = \min_{\substack{P \in \mathbb{R}_+^{n \times m} \\ P \mathbf{1}_m = \mathbf{a}, P^T \mathbf{1}_n = \mathbf{b}}} \langle P, M_{\mathbf{x}\mathbf{y}} \rangle$$

Dual Kantorovich Problem

$$W_p^p(\mu, \nu) = \min_{\substack{P \in \mathbb{R}_+^{n \times m} \\ P \mathbf{1}_m = \mathbf{a}, P^T \mathbf{1}_n = \mathbf{b}}} \langle P, M_{\mathbf{x}\mathbf{y}} \rangle$$

Def. Dual OT problem

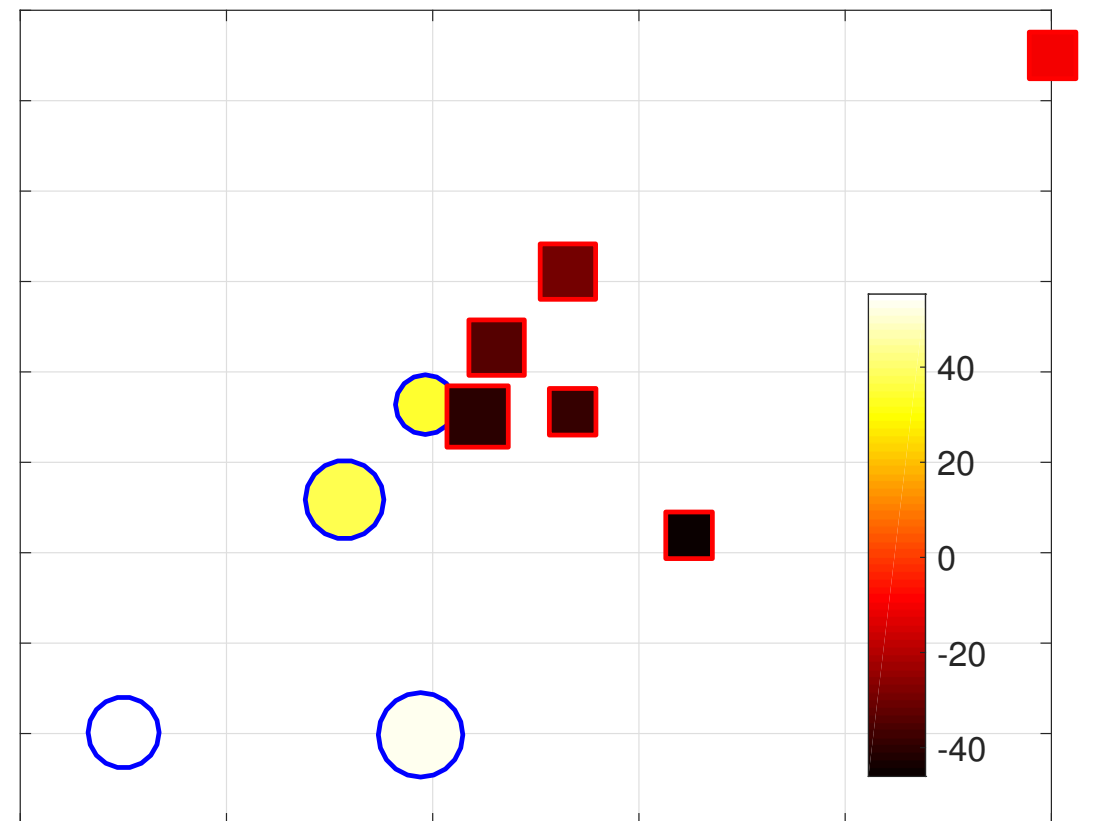
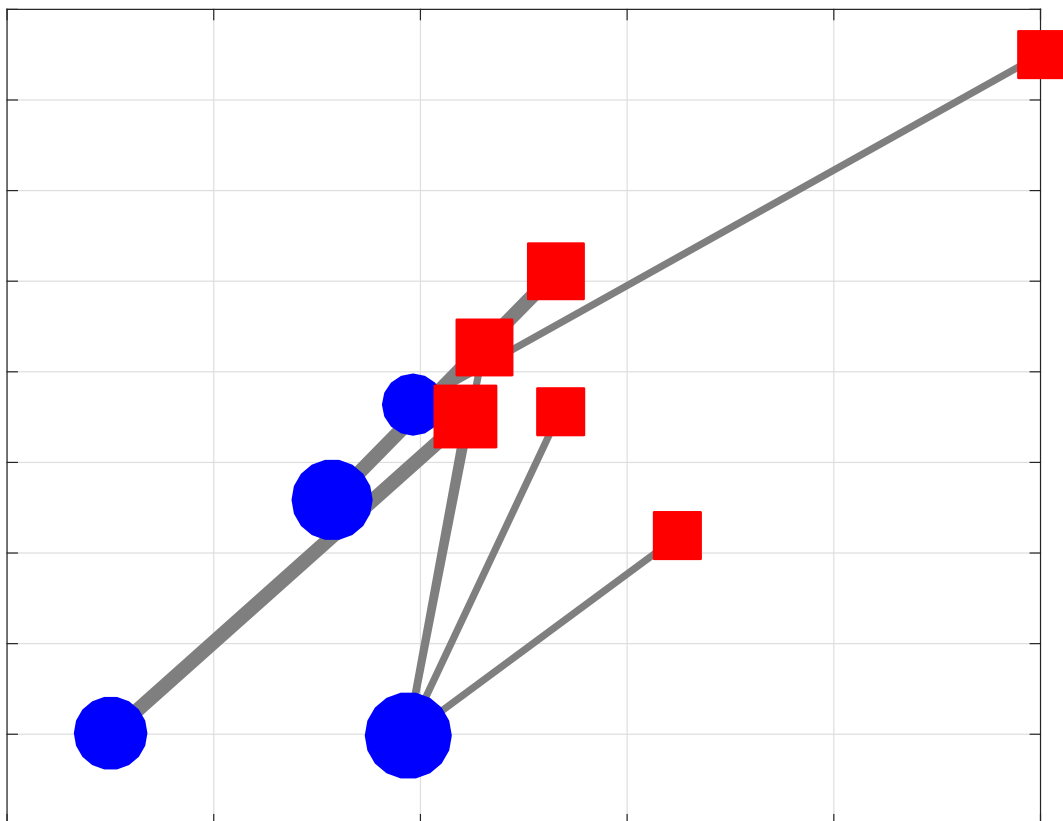
$$W_p^p(\mu, \nu) = \max_{\substack{\alpha \in \mathbb{R}^n, \beta \in \mathbb{R}^m \\ \alpha_i + \beta_j \leq D(\mathbf{x}_i, \mathbf{y}_j)^p}} \alpha^T \mathbf{a} + \beta^T \mathbf{b}$$

Dual Kantorovich Problem

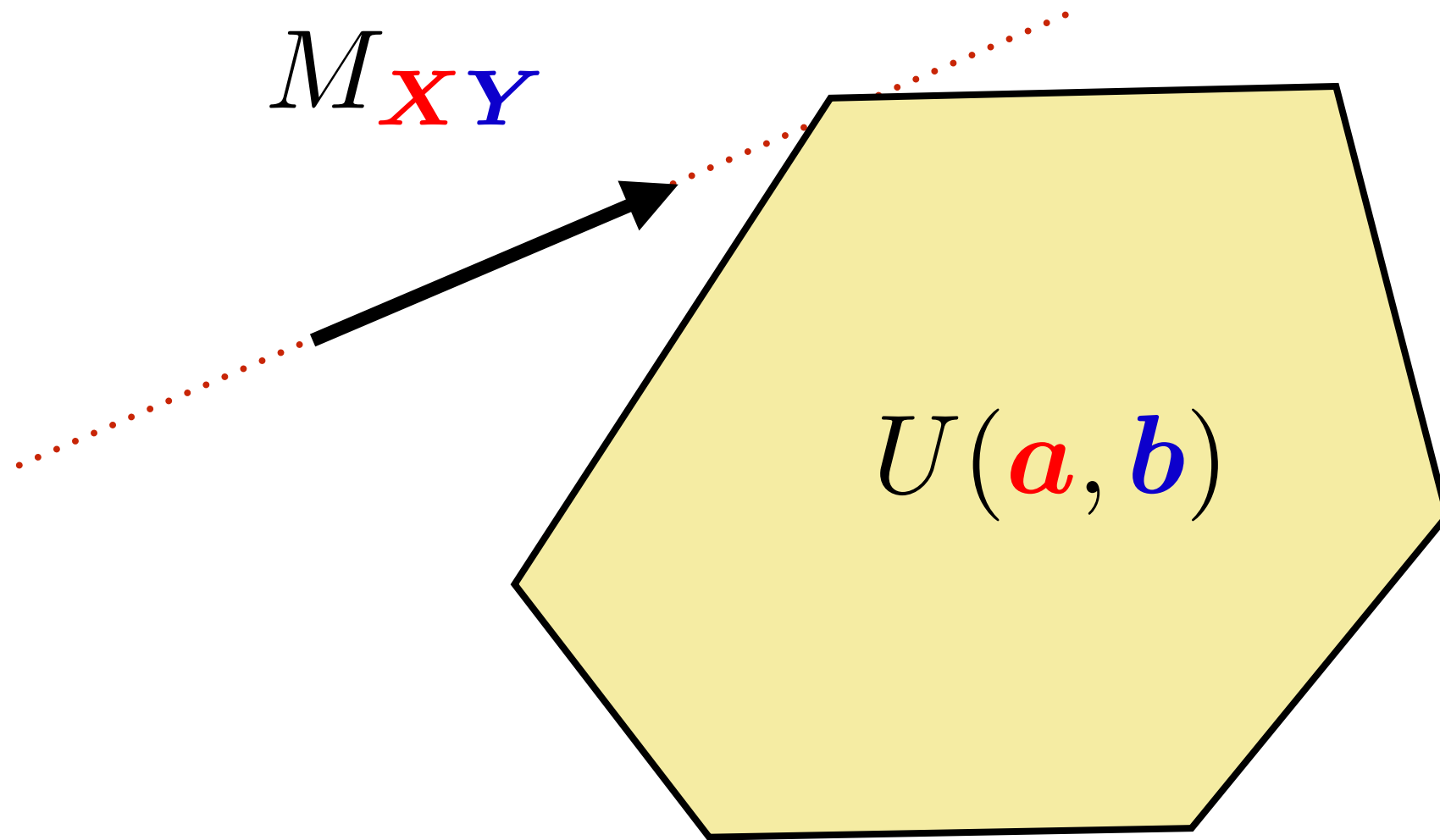
$$W_p^p(\mu, \nu) = \min_{\substack{P \in \mathbb{R}_+^{n \times m} \\ P \mathbf{1}_m = \mathbf{a}, P^T \mathbf{1}_n = \mathbf{b}}} \langle P, M_{\mathbf{x}\mathbf{y}} \rangle$$

Def. Dual OT problem

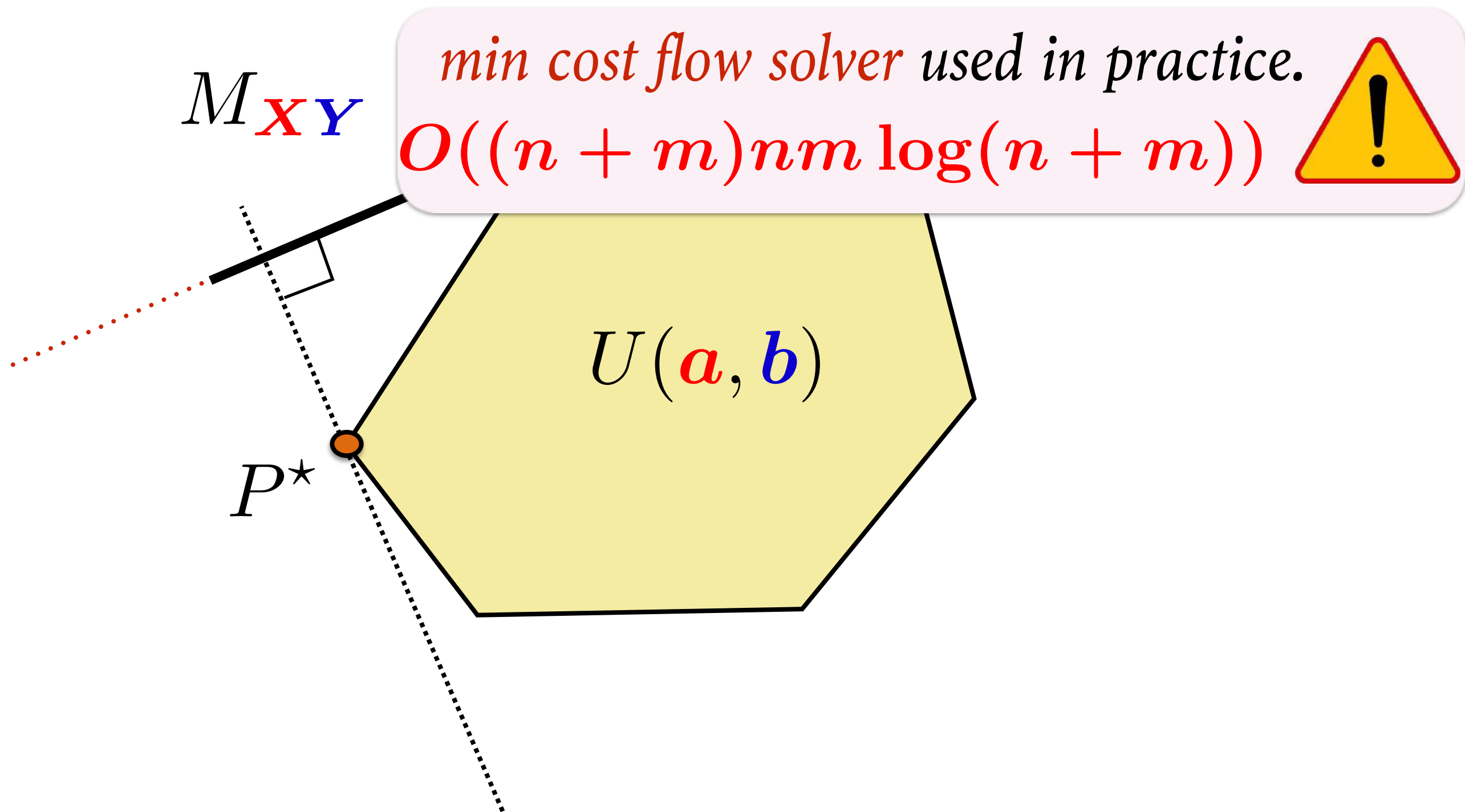
$$W_p^p(\mu, \nu) = \max_{\substack{\alpha \in \mathbb{R}^n, \beta \in \mathbb{R}^m \\ \alpha_i + \beta_j \leq D(\mathbf{x}_i, \mathbf{y}_j)^p}} \alpha^T \mathbf{a} + \beta^T \mathbf{b}$$



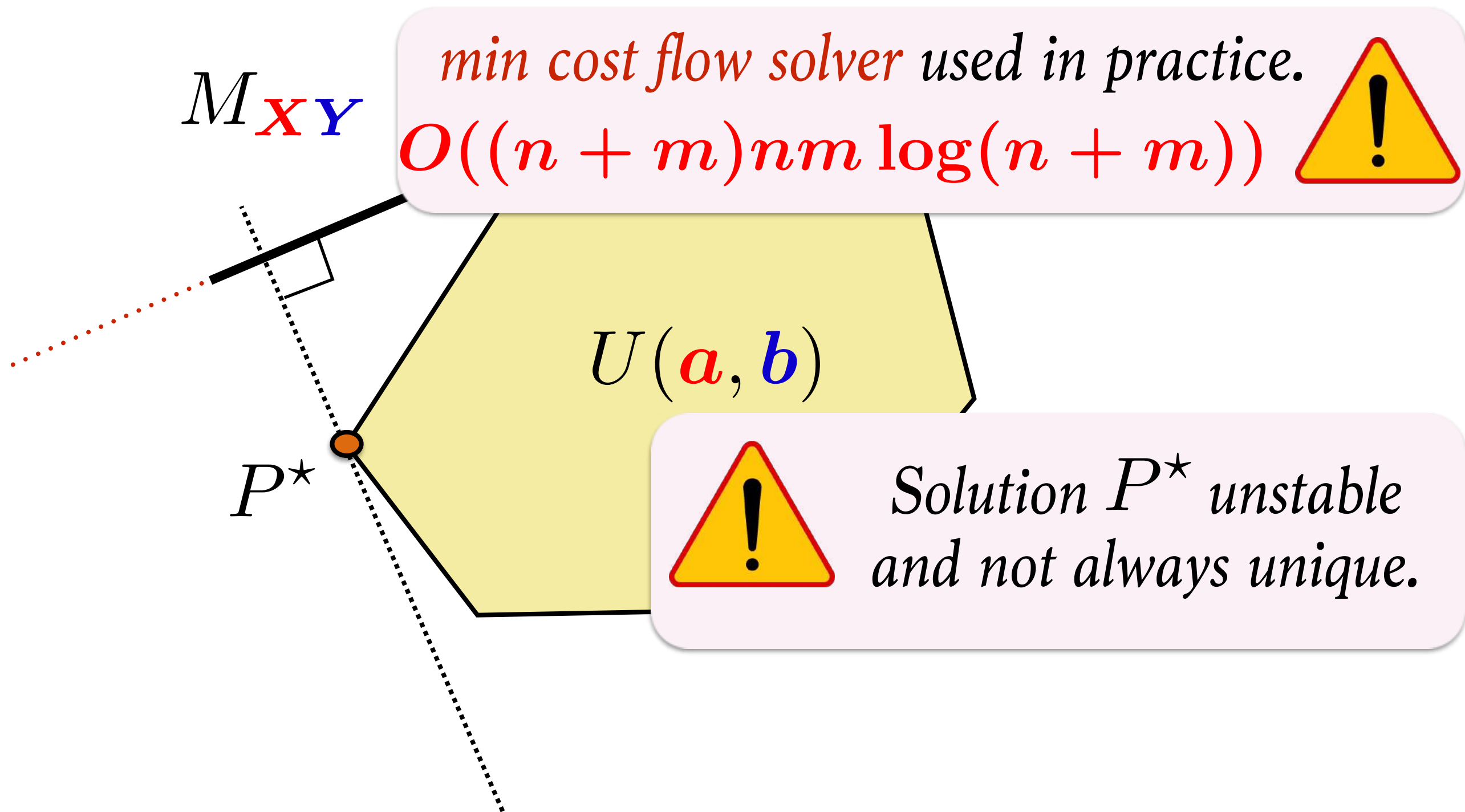
Solving the OT Problem



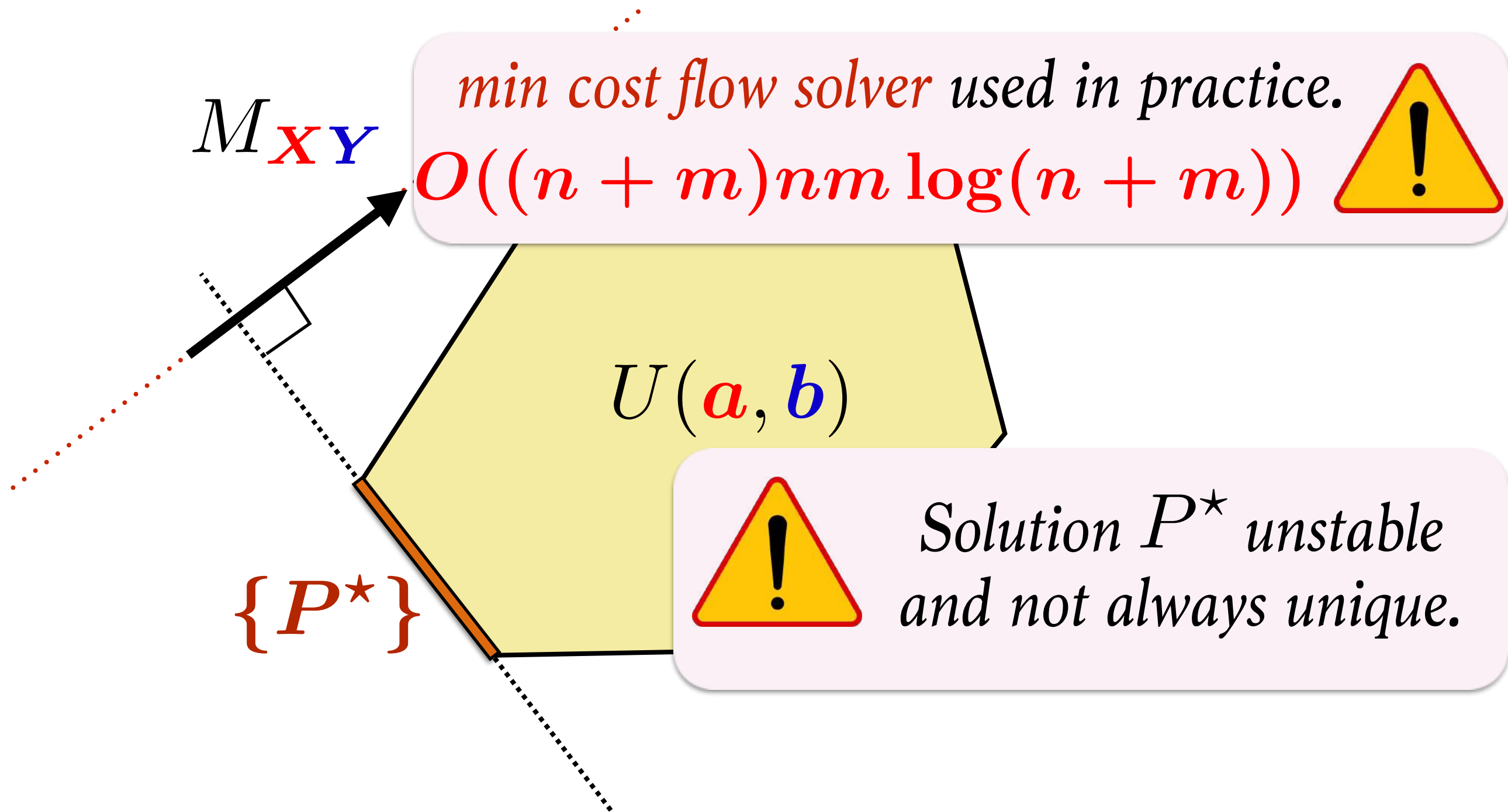
Solving the OT Problem



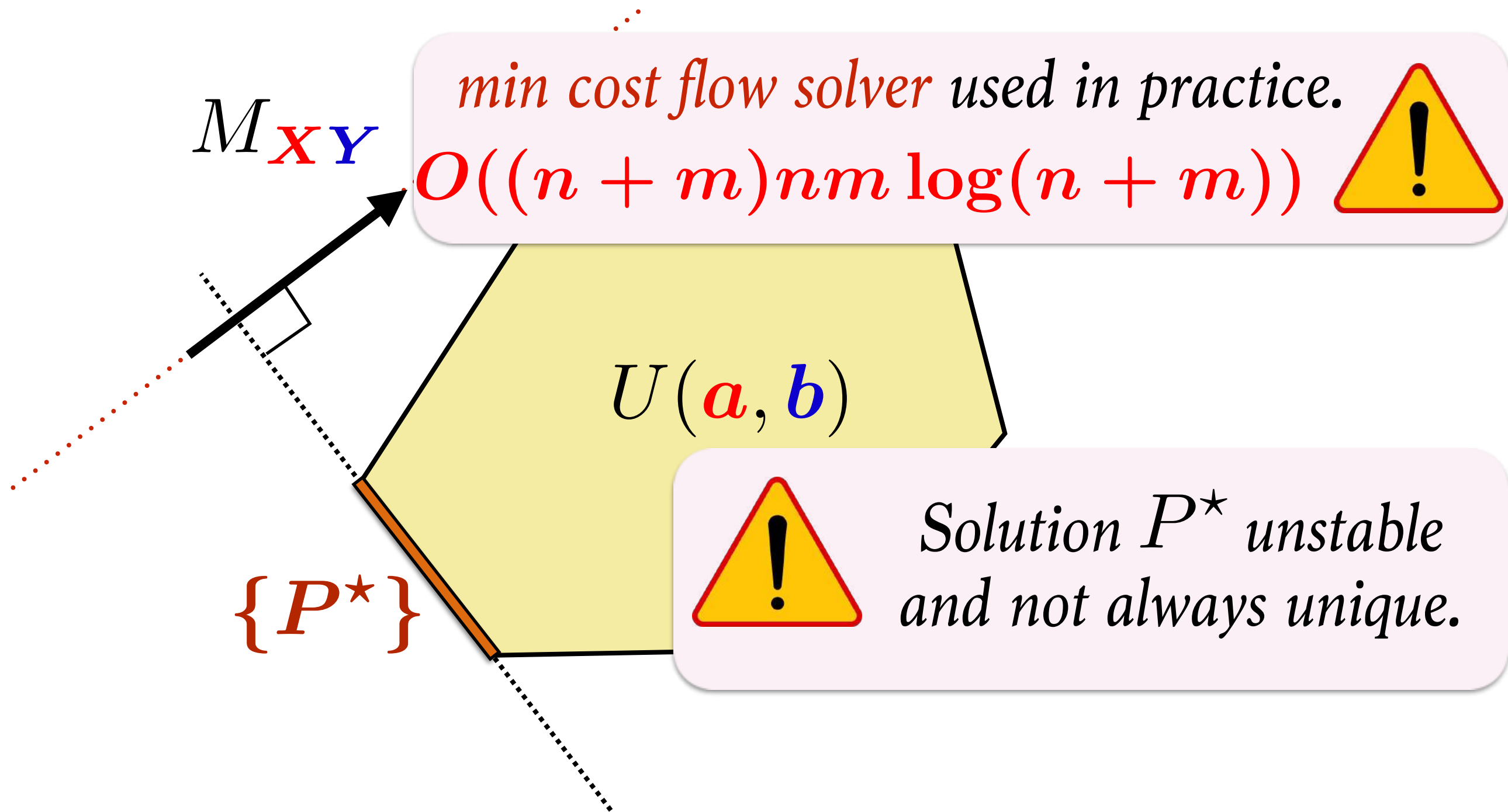
Solving the OT Problem



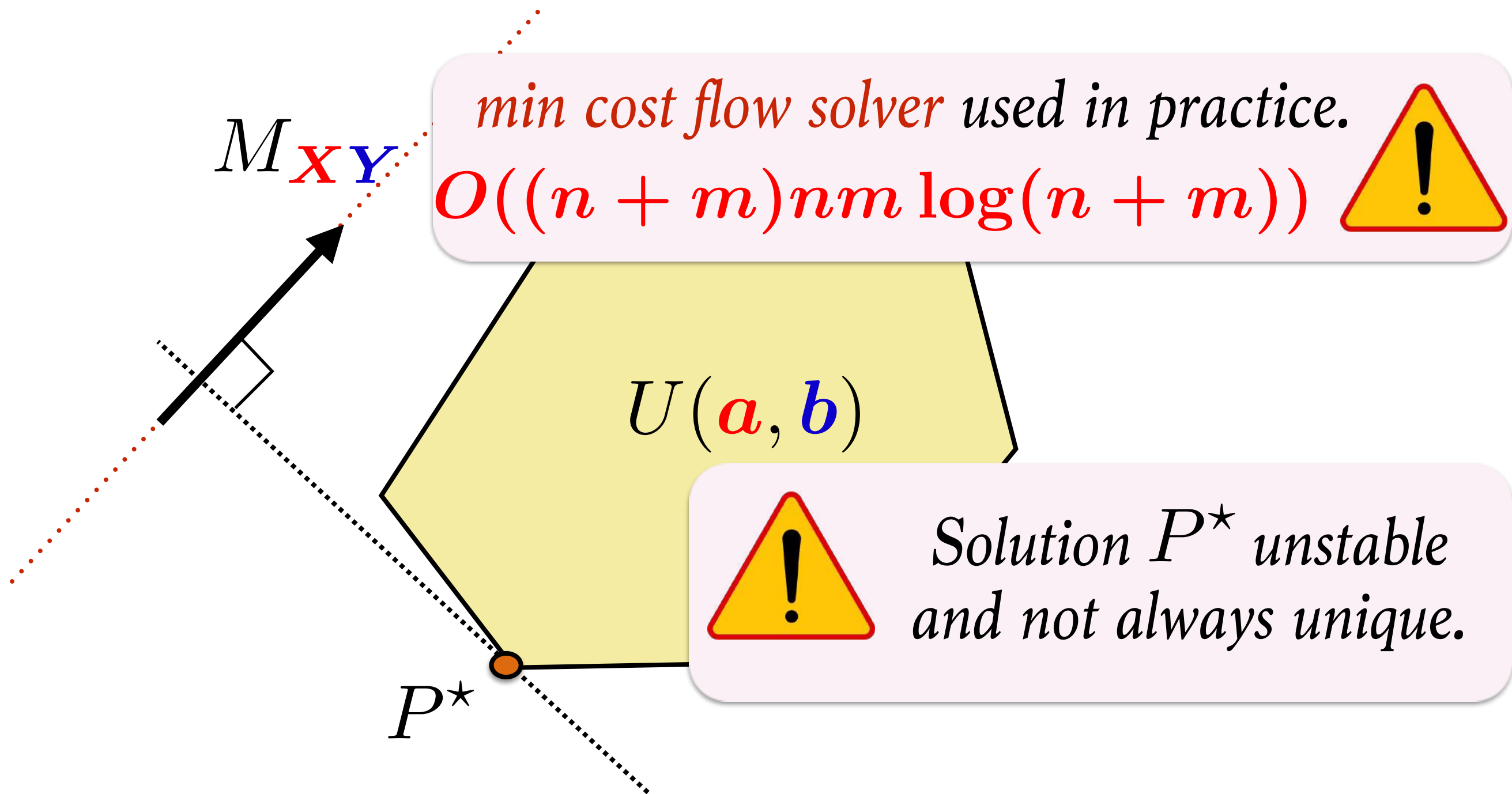
Solving the OT Problem



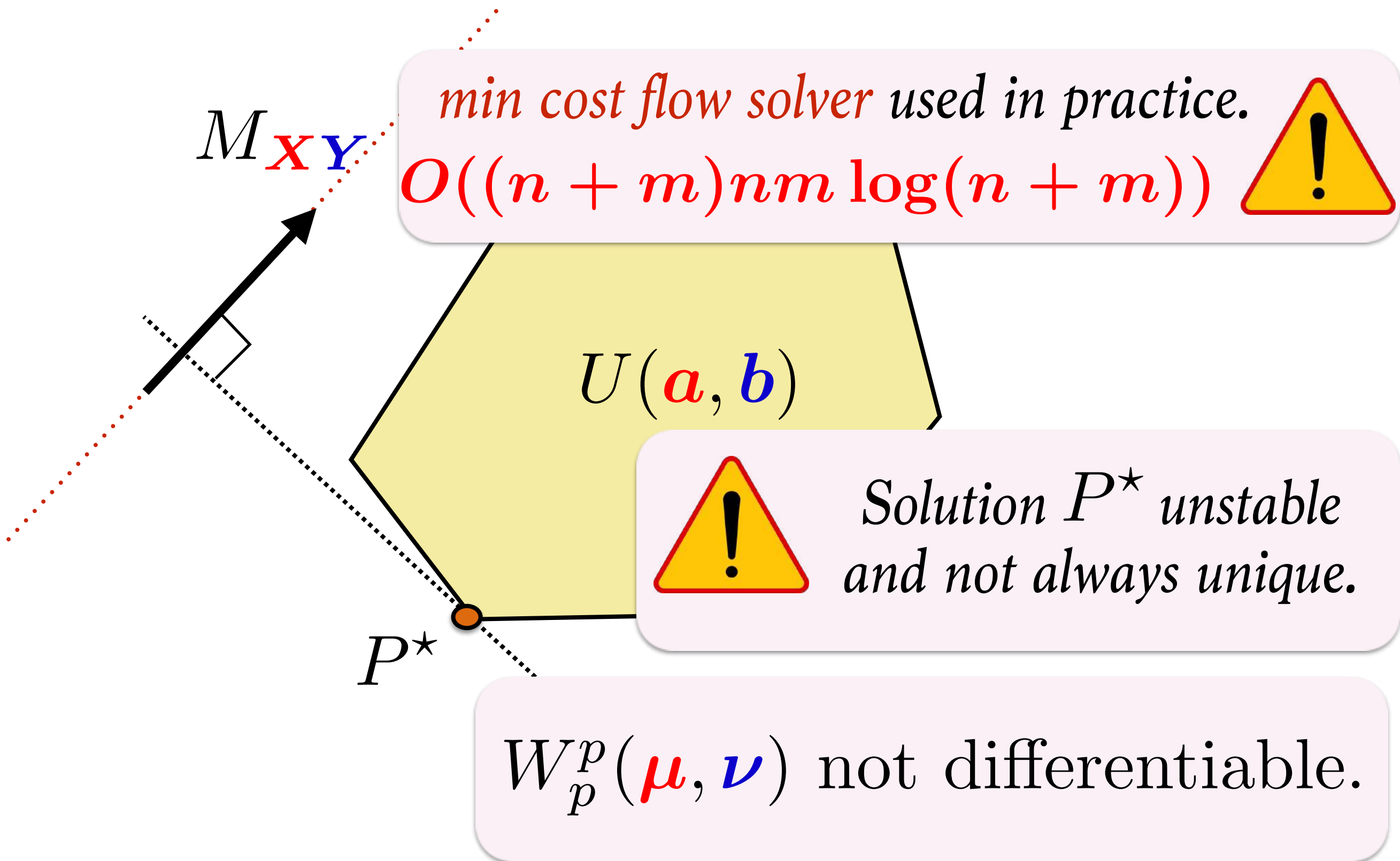
Solving the OT Problem



Solving the OT Problem



Solving the OT Problem



Discrete OT Problem

```
emd.c
Last update: 3/14/98

An implementation of the Earth Movers Distance.
Based of the solution for the Transportation problem as described in
"Introduction to Mathematical Programming" by F. S. Hillier and
G. J. Lieberman, McGraw-Hill, 1990.

Copyright (C) 1998 Yossi Rubner
Computer Science Department, Stanford University
E-Mail: rubner@cs.stanford.edu URL: http://vision.stanford.edu/~rubner

/*
#include <stdio.h>
#include <stdlib.h>
#include <math.h>

#include "emd.h"

#define DEBUG_LEVEL 0
/*
DEBUG_LEVEL:
0 = NO MESSAGES
1 = PRINT THE NUMBER OF ITERATIONS AND THE FINAL RESULT
2 = PRINT THE RESULT AFTER EVERY ITERATION
3 = PRINT ALSO THE FLOW AFTER EVERY ITERATION
4 = PRINT A LOT OF INFORMATION (PROBABLY USEFUL ONLY FOR THE AUTHOR)
*/

#define MAX_SIG_SIZE1 (MAX_SIG_SIZE+1) /* FOR THE POSSIBLE DUMMY FEATURE */

/* NEW TYPES DEFINITION */
/* node1_t IS USED FOR SINGLE-LINKED LISTS */
typedef struct node1_t {
    int i;
    double val;
    struct node1_t *Next;
} node1_t;

/* node1_t IS USED FOR DOUBLE-LINKED LISTS */
typedef struct node2_t {
    int i, j;
    double val;
    struct node2_t *NextC; /* NEXT COLUMN */
    struct node2_t *NextR; /* NEXT ROW */
} node2_t;

/* GLOBAL VARIABLE DECLARATION */
static int _n1, _n2; /* SIGNATURES SIZES */
static float _C[MAX_SIG_SIZE1][MAX_SIG_SIZE1]; /* THE COST MATRIX */
static node2_t _X[MAX_SIG_SIZE1*2]; /* THE BASIC VARIABLES VECTOR */
```

ice.



Discrete OT Problem

```
emd.c
Last update: 3/14/98

An implementation of the Earth Movers Distance.
Based of the solution for the Transportation problem as described in
"Introduction to Mathematical Programming" by F. S. Hillier and
G. J. Lieberman, McGraw-Hill, 1990.

Copyright (C) 1998 Yossi Rubner
Computer Science Department, Stanford University
E-Mail: rubner@cs.stanford.edu URL: http://vision.stanford.edu/~rubner

/*
#include <stdio.h>
#include <stdlib.h>
#include <math.h>

#include "emd.h"

#define DEBUG_LEVEL 0
/*
DEBUG_LEVEL:
0 = NO MESSAGES
1 = PRINT THE NUMBER OF ITERATIONS AND THE FINAL RESULT
2 = PRINT THE RESULT AFTER EVERY ITERATION
3 = PRINT ALSO THE FLOW AFTER EVERY ITERATION
4 = PRINT A LOT OF INFORMATION (PROBABLY USEFUL ONLY FOR THE AUTHOR)
*/

#define MAX_SIG_SIZE1 (MAX_SIG_SIZE+1) /* FOR THE POSSIBLE DUMMY FEATURE */

/* NEW TYPES DEFINITION */

/* node1_t IS USED FOR SINGLE-LINKED LISTS */
typedef struct node1_t {
    int i;
    double val;
    struct node1_t *Next;
} node1_t;

/* node1_t IS USED FOR DOUBLE-LINKED LISTS */
typedef struct node2_t {
    int i, j;
    double val;
    struct node2_t *NextC; /* NEXT COLUMN */
    struct node2_t *NextR; /* NEXT ROW */
} node2_t;

/* GLOBAL VARIABLE DECLARATION */
static int _n1, _n2; /* SIGNATURES SIZES */
static float _C[MAX_SIG_SIZE1][MAX_SIG_SIZE1]; /* THE COST MATRIX */
static node2_t _X[MAX_SIG_SIZE1*2]; /* THE BASIC VARIABLES VECTOR */
```

ice.



3. Computing OT for data sciences

- The need for regularization
- Entropic regularization
- Differentiation

What matters for practitioners?

i.i.d samples $\mathbf{x}_1, \dots, \mathbf{x}_n \sim \mu$, $\mathbf{y}_1, \dots, \mathbf{y}_m \sim \nu$,

$$\hat{\mu}_n \stackrel{\text{def}}{=} \frac{1}{n} \sum_i \delta_{\mathbf{x}_i}, \hat{\nu}_m \stackrel{\text{def}}{=} \frac{1}{m} \sum_j \delta_{\mathbf{y}_j}$$

Computational properties

Compute/approximate $W_p(\hat{\mu}_n, \hat{\nu}_m)$?

Statistical properties

$$\mathbb{E} [|W_p(\mu, \nu) - W_p(\hat{\mu}_n, \hat{\nu}_m)|] \leq f(n, m)?$$

What matters for practitioners?

i.i.d samples $\mathbf{x}_1, \dots, \mathbf{x}_n \sim \mu$, $\mathbf{y}_1, \dots, \mathbf{y}_m \sim \nu$,

$$\hat{\mu}_n \stackrel{\text{def}}{=} \frac{1}{n} \sum_i \delta_{\mathbf{x}_i}, \hat{\nu}_m \stackrel{\text{def}}{=} \frac{1}{m} \sum_j \delta_{\mathbf{y}_j}$$

Computational properties



$$O((n + m)nm \log(n + m))$$

Statistical properties

$$\mathbb{E} [|\mathbf{W}_p(\mu, \nu) - \mathbf{W}_p(\hat{\mu}_n, \hat{\nu}_m)|] \leq f(n, m)?$$

Sample Complexity



If $\Omega = \mathbb{R}^d, d > 3$

$$\mathbb{E} [|\mathbf{W}_p(\mu, \nu) - \mathbf{W}_p(\hat{\mu}_n, \hat{\nu}_n)|] = O(n^{-1/d})$$

- **[Dudley'69][Dereich+'11][Fournier+'13] & others..**
- **[Weed/Bach'17]:** sharper results when measures' support has “low effective d ” in metric spaces
- **[Weed/Berthet'19]** for smooth densities
- **Lower bounds:** optimal quantization error.

From theory to practice ?

$$\hat{\mu}_n \stackrel{\text{def}}{=} \frac{1}{n} \sum_i \delta_{\mathbf{x}_i}, \hat{\nu}_m \stackrel{\text{def}}{=} \frac{1}{m} \sum_j \delta_{\mathbf{y}_j}$$

Computational properties



$$O((n + m)nm \log(n + m))$$

Statistical properties



$$\mathbb{E} [||W_p(\mu, \nu) - W_p(\hat{\mu}_n, \hat{\nu}_n)||] = O(n^{-1/d})$$

From theory to practice ?

$$\hat{\mu}_n \stackrel{\text{def}}{=} \frac{1}{n} \sum_i \delta_{\mathbf{x}_i}, \hat{\nu}_m \stackrel{\text{def}}{=} \frac{1}{m} \sum_j \delta_{\mathbf{y}_j}$$

Computational properties



$$O((n + m)nm \log(n + m))$$

Statistical properties



$$\mathbb{E} [||W_p(\mu, \nu) - W_p(\hat{\mu}_n, \hat{\nu}_n)||] = O(n^{-1/d})$$

For data sciences, we must regularize the problem to improve on either/both aspects.

Many ways to regularize (dual) OT

$$\sup_{\varphi(x) + \psi(y) \leq c(x,y)} \int \varphi d\mu + \int \psi d\nu.$$

Many ways to regularize (dual) OT

$$\sup_{\varphi(x) + \psi(y) \leq c(x,y)} \int \varphi d\mu + \int \psi d\nu.$$

- RKHS for potentials [\[GCBP'16\]](#), dualize/smooth indicator constraint, [\[VMVRB'21\]](#)

Many ways to regularize (dual) OT

$$\sup_{\varphi(x) + \psi(y) \leq c(x,y)} \int \varphi d\mu + \int \psi d\nu.$$

- RKHS for potentials [\[GCBP'16\]](#), dualize/smooth indicator constraint, [\[VMVRB'21\]](#)

$$W_1(\mu, \nu) = \sup_{\varphi \text{ 1-Lipschitz}} \int \varphi (d\mu - d\nu).$$

- Parameterize functions using ReLU Deep net with bounded weights [\[Arjovsky+'17\]](#) or use Wavelet decompositions [\[Shirdonkhar+'08\]](#) for low d .

Many ways to regularize (primal) OT

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

Many ways to regularize (primal) OT

$$\inf_{P \in \Pi(\mu, \nu)} \iint \mathbf{c}(x, y) P(dx, dy).$$

- Change cost function: threshold metric [**Pele+'09**], use geodesic distance on graphs [**Beckman'52**] [**Lin+'07**], [**Solomon+'14**], simplifies the LP.

Many ways to regularize (primal) OT

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

- Change cost function: threshold metric [**Pele+'09**], use geodesic distance on graphs [**Beckman'52**] [**Lin+'07**], [**Solomon+'14**], simplifies the LP.
- Quantize measures first [**Canas+'12**]; use Gaussians [**Gelbrich'92**][**Chen+17**]; projections [**Rabin+'11**] & k -dimensional subspaces [**Paty+'19**][**Weed+'19**].

Many ways to regularize (primal) OT

$$\inf_{P \in \Pi(\mu, \nu)} \iint c(x, y) P(dx, dy).$$

- Change cost function: threshold metric [**Pele+'09**], use geodesic distance on graphs [**Beckman'52**] [**Lin+'07**], [**Solomon+'14**], simplifies the LP.
- Quantize measures first [**Canas+'12**]; use Gaussians [**Gelbrich'92**][**Chen+17**]; projections [**Rabin+'11**] & k -dimensional subspaces [**Paty+'19**][**Weed+'19**].
- Add regularization on coupling [**C'13**][**GP'16**] [**GCBP'16**] [**GCBP'19**] [**DPR'16**] [**BSR'18**]

A Different Route: Projection

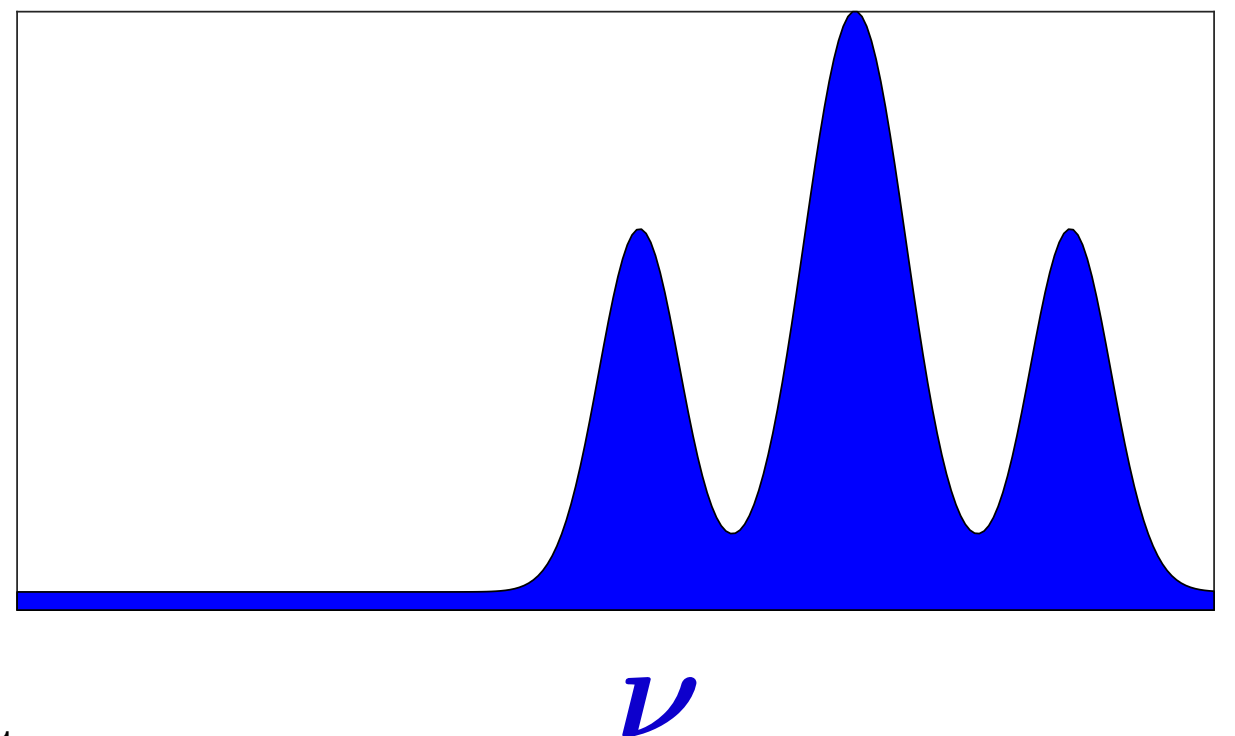
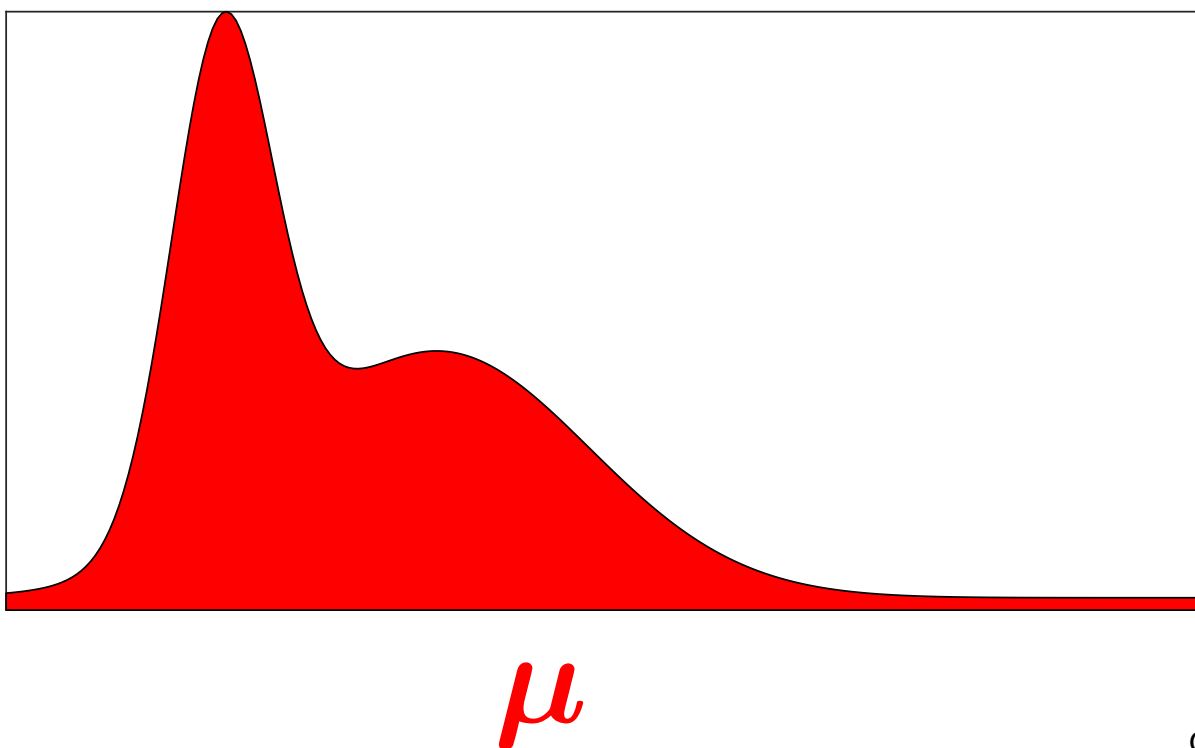
Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$

A Different Route: Projection

Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

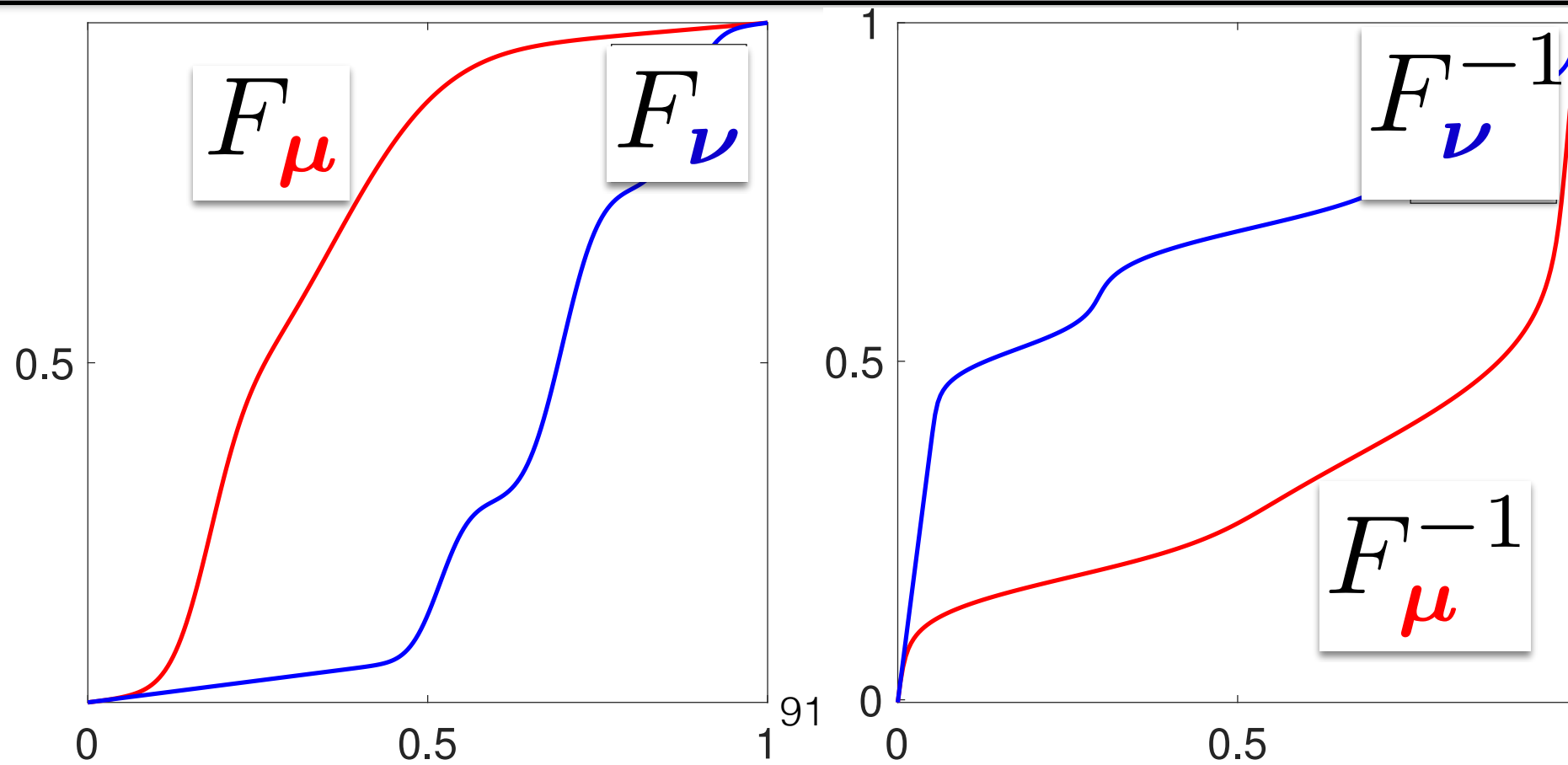
$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$



A Different Route: Projection

Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

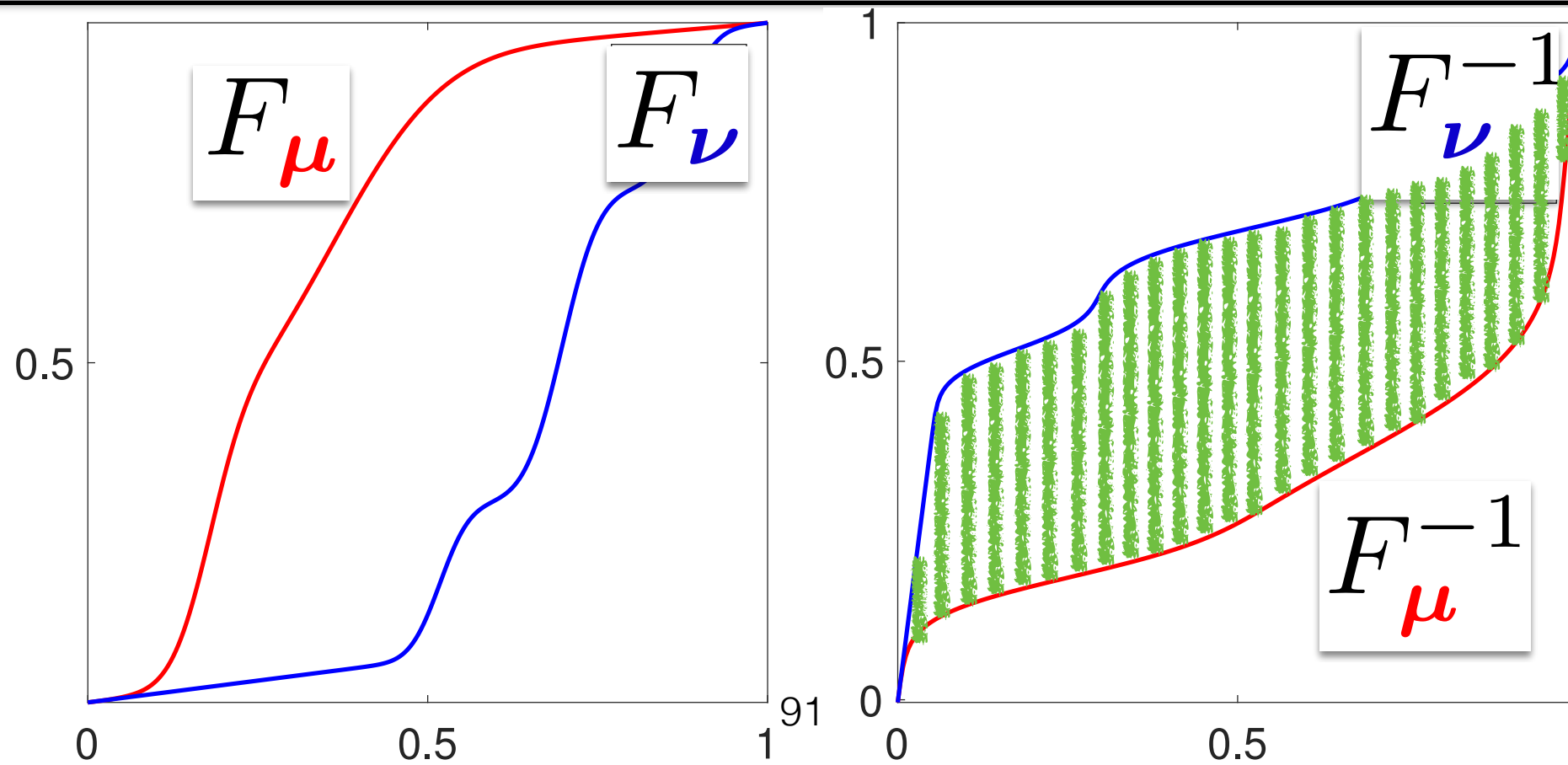
$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$



A Different Route: Projection

Remark. If $\Omega = \mathbb{R}$, $\mathbf{c}(x, y) = \mathbf{c}(|x - y|)$, $\mathbf{c}$ convex, $F_{\mu}^{-1}, F_{\nu}^{-1}$ quantile functions,

$$W(\mu, \nu) = \int_0^1 \mathbf{c}(|F_{\mu}^{-1}(x) - F_{\nu}^{-1}(x)|) dx$$



A simple baseline

- Dodges the high-dimensionality curse, by simplifying considerably the measures.
- Effective in practice, fast and easy, but far from OT. Induced matchings are extremely blurry.

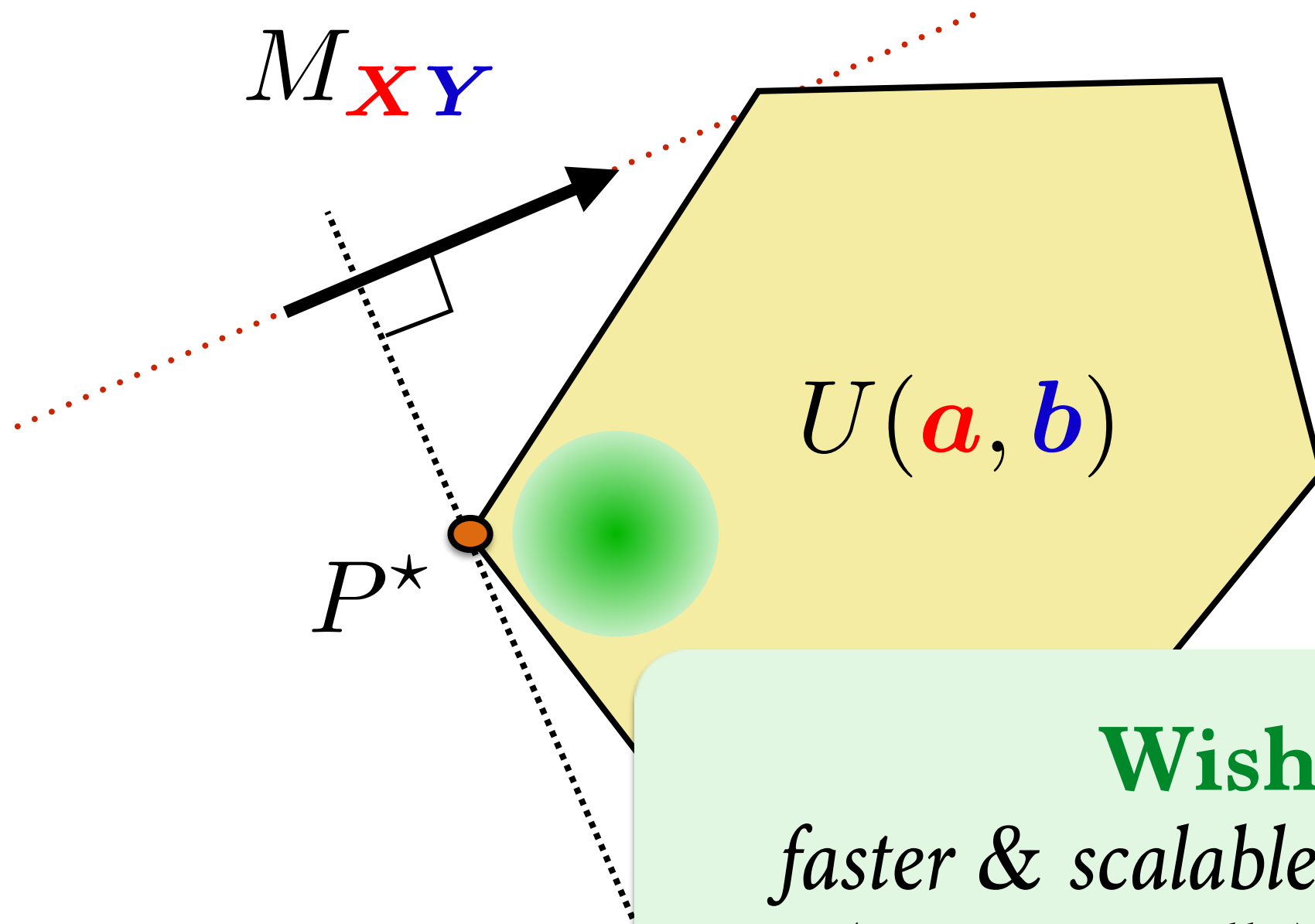
A simple baseline

Sliced Wasserstein Distance [Rabin+'11]

$$SW(\mu, \nu) = \mathbb{E}_{\theta \sim \mathcal{S}^{d-1}} \left[\int_0^1 c(|F_{\theta_{\#}^T \mu}^{-1}(x) - F_{\theta_{\#}^T \nu}^{-1}(x)|) dx \right]$$

- Dodges the high-dimensionality curse, by simplifying considerably the measures.
- Effective in practice, fast and easy, but far from OT. Induced matchings are extremely blurry.

Regularization on the Primal



Wishlist:

*faster & scalable, more stable,
(automatically) differentiable*

Entropic Regularization [Wilson'62]

Def. Regularized Wasserstein, $\gamma \geq 0$

$$W_\gamma(\mu, \nu) \stackrel{\text{def}}{=} \min_{P \in U(a, b)} \langle P, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(P)$$

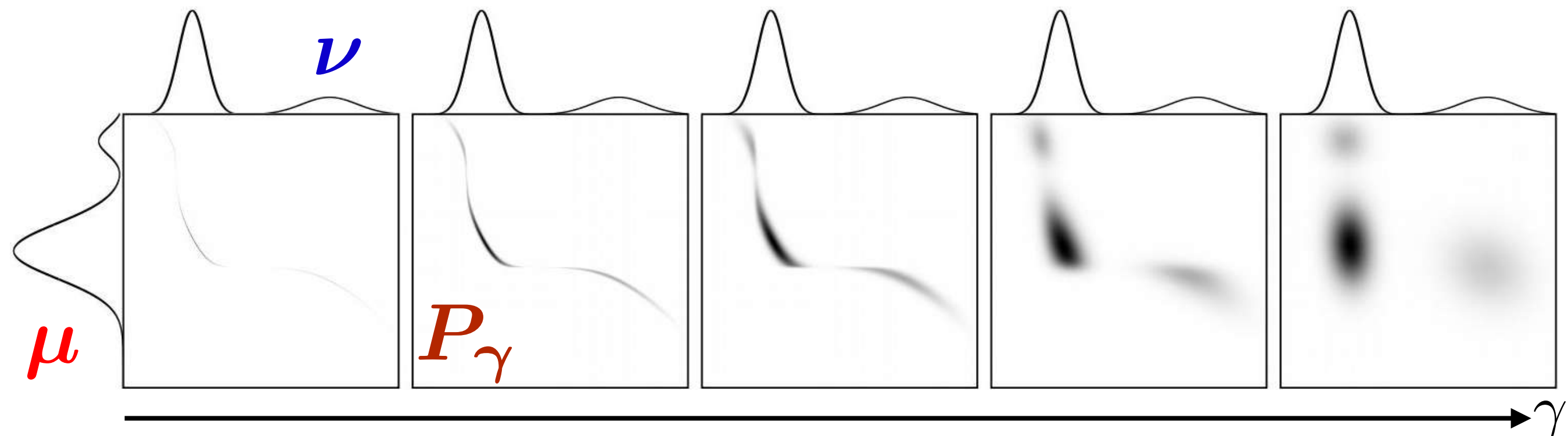
$$E(P) \stackrel{\text{def}}{=} - \sum_{i,j=1}^{nm} P_{ij} (\log P_{ij} - 1)$$

Note: Unique optimal solution because of strong concavity of entropy

Entropic Regularization [Wilson'62]

Def. Regularized Wasserstein, $\gamma \geq 0$

$$W_\gamma(\mu, \nu) \stackrel{\text{def}}{=} \min_{P \in U(\mu, \nu)} \langle P, M_{XY} \rangle - \gamma E(P)$$

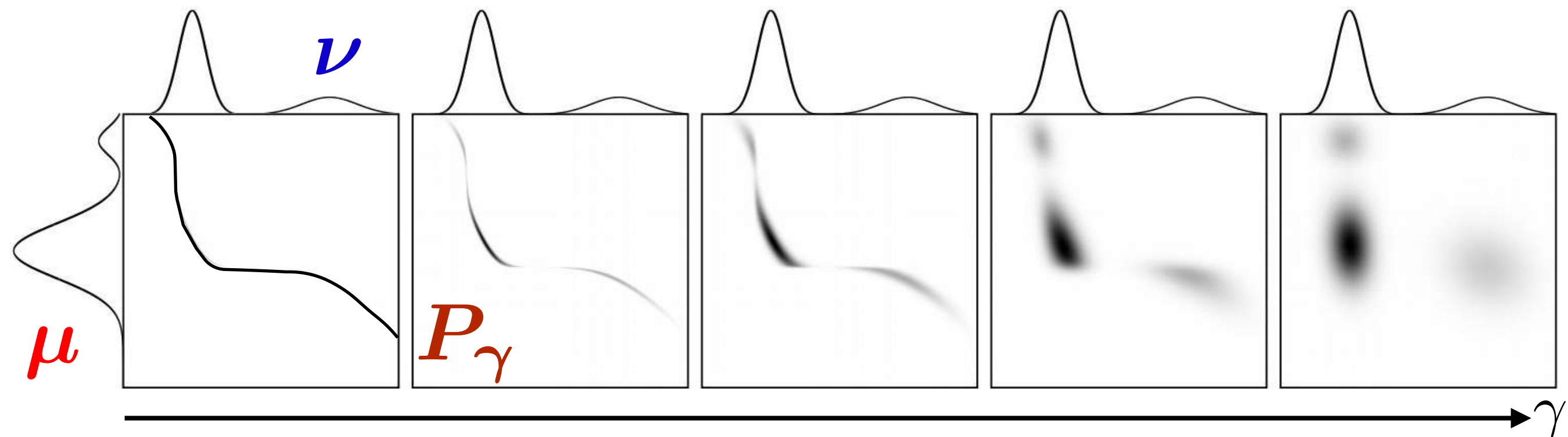


Note: Unique optimal solution because of strong concavity of entropy

Entropic Regularization [Wilson'62]

Def. Regularized Wasserstein, $\gamma \geq 0$

$$W_\gamma(\mu, \nu) \stackrel{\text{def}}{=} \min_{P \in U(\mu, \nu)} \langle P, M_{XY} \rangle - \gamma E(P)$$

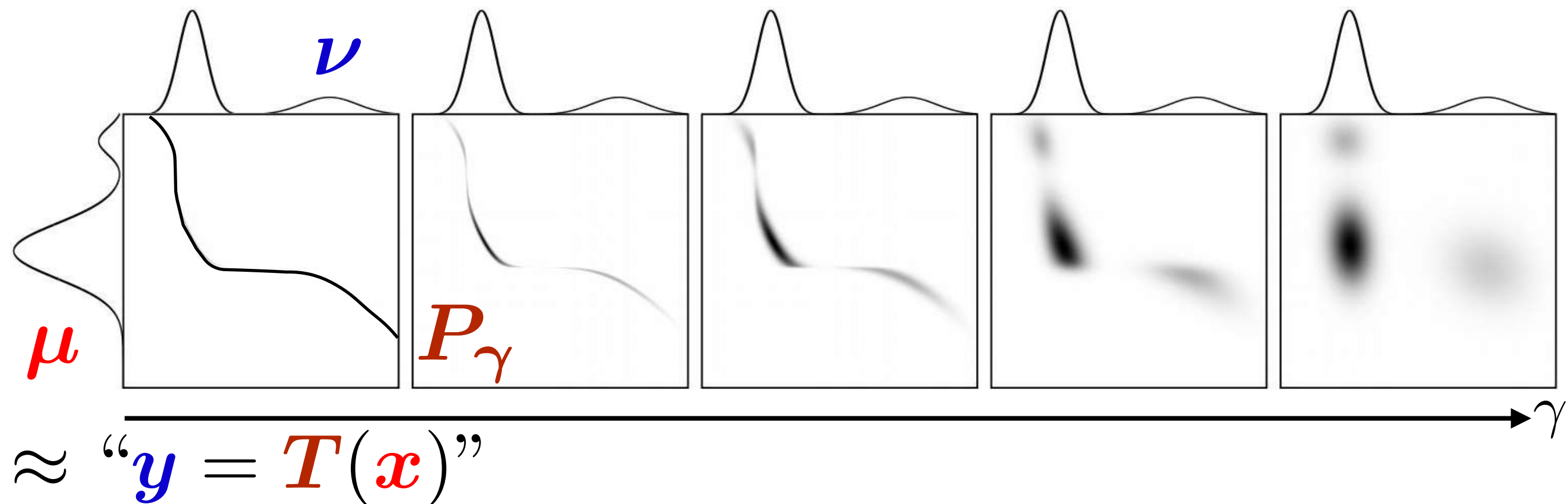


Note: Unique optimal solution because of strong concavity of entropy

Entropic Regularization [Wilson'62]

Def. Regularized Wasserstein, $\gamma \geq 0$

$$W_\gamma(\mu, \nu) \stackrel{\text{def}}{=} \min_{P \in U(\mu, \nu)} \langle P, M_{XY} \rangle - \gamma E(P)$$



Note: Unique optimal solution because of strong concavity of entropy

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$L(P, \alpha, \beta) = \sum_{ij} P_{ij} M_{ij} + \gamma P_{ij} (\log P_{ij} - 1) + \alpha^T (P \mathbf{1} - \mathbf{a}) + \beta^T (P^T \mathbf{1} - \mathbf{b})$$

$$\partial L / \partial P_{ij} = M_{ij} + \gamma \log P_{ij} + \alpha_i + \beta_j$$

$$(\partial L / \partial P_{ij} = 0) \Rightarrow P_{ij} = e^{\frac{\alpha_i}{\gamma}} e^{-\frac{M_{ij}}{\gamma}} e^{\frac{\beta_j}{\gamma}} = \mathbf{u}_i \mathbf{K}_{ij} \mathbf{v}_j$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$P_\gamma \in U(\mathbf{a}, \mathbf{b}) \Leftrightarrow \begin{cases} \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}) \mathbf{1}_m & = \mathbf{a} \\ \operatorname{diag}(\mathbf{v}) \mathbf{K}^T \operatorname{diag}(\mathbf{u}) \mathbf{1}_n & = \mathbf{b} \end{cases}$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$P_\gamma \in U(\mathbf{a}, \mathbf{b}) \Leftrightarrow \begin{cases} \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}) \mathbf{1}_m & = \mathbf{a} \\ \operatorname{diag}(\mathbf{v}) \underbrace{\mathbf{K}^T \operatorname{diag}(\mathbf{u}) \mathbf{1}_n}_{\mathbf{u}} & = \mathbf{b} \end{cases}$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$P_\gamma \in U(\mathbf{a}, \mathbf{b}) \Leftrightarrow \begin{cases} \overbrace{\operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}) \mathbf{1}_m}^{\mathbf{v}} = \mathbf{a} \\ \underbrace{\operatorname{diag}(\mathbf{v}) \mathbf{K}^T \operatorname{diag}(\mathbf{u}) \mathbf{1}_n}_{\mathbf{u}} = \mathbf{b} \end{cases}$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$P_\gamma \in U(\mathbf{a}, \mathbf{b}) \Leftrightarrow \begin{cases} \operatorname{diag}(\mathbf{u}) \mathbf{K} \mathbf{v} & = \mathbf{a} \\ \operatorname{diag}(\mathbf{v}) \mathbf{K}^T \mathbf{u} & = \mathbf{b} \end{cases}$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$P_\gamma \in U(\mathbf{a}, \mathbf{b}) \Leftrightarrow \begin{cases} \mathbf{u} \odot \mathbf{K} \mathbf{v} & = \mathbf{a} \\ \mathbf{v} \odot \mathbf{K}^T \mathbf{u} & = \mathbf{b} \end{cases}$$

Fast & Scalable Algorithm

Prop. If $P_\gamma \stackrel{\text{def}}{=} \underset{P \in U(\mathbf{a}, \mathbf{b})}{\operatorname{argmin}} \langle \mathbf{P}, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(\mathbf{P})$

then $\exists! \mathbf{u} \in \mathbb{R}_+^n, \mathbf{v} \in \mathbb{R}_+^m$, such that

$$P_\gamma = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad \mathbf{K} \stackrel{\text{def}}{=} e^{-M_{\mathbf{x}\mathbf{y}} / \gamma}$$

$$P_\gamma \in U(\mathbf{a}, \mathbf{b}) \Leftrightarrow \begin{cases} \mathbf{u} = \mathbf{a} / \mathbf{K} \mathbf{v} \\ \mathbf{v} = \mathbf{b} / \mathbf{K}^T \mathbf{u} \end{cases}$$

Fast & Scalable Algorithm

Sinkhorn's Algorithm : Repeat

$$1. \quad \mathbf{u} = \mathbf{a} / \mathbf{K} \mathbf{v}$$

$$2. \quad \mathbf{v} = \mathbf{b} / \mathbf{K}^T \mathbf{u}$$

Fast & Scalable Algorithm

Sinkhorn's Algorithm : Repeat

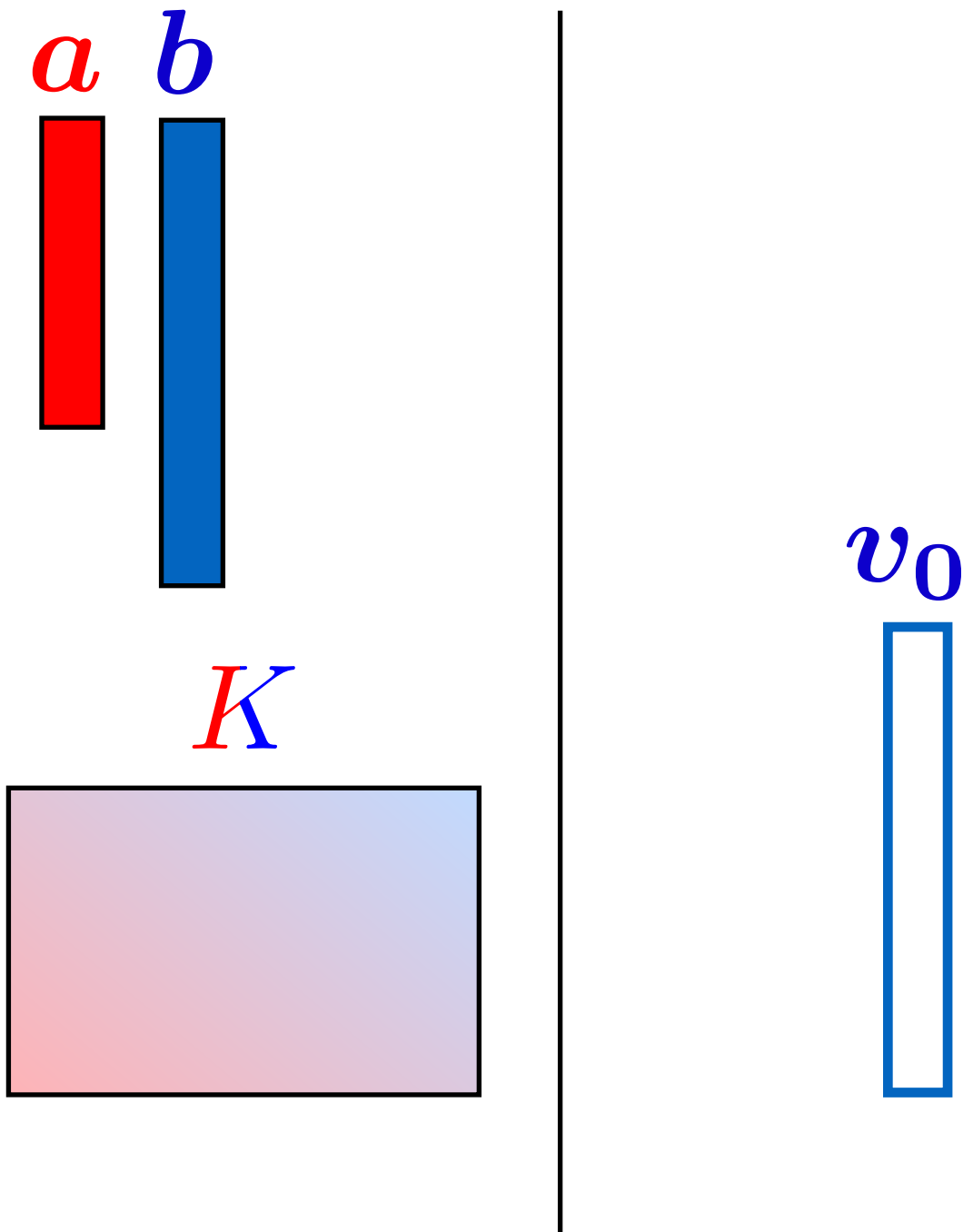
$$1. \quad \mathbf{u} = \mathbf{a} / \mathbf{K} \mathbf{v}$$

$$2. \quad \mathbf{v} = \mathbf{b} / \mathbf{K}^T \mathbf{u}$$

- [Sinkhorn'64] proved first convergence result
[Lorenz+'89] characterised linear convergence
- Recent wave of great results by [Altschuler+'17]
[Dvurechensky+18][Lin+19]
- $O(nm)$ complexity, GPGPU parallel [C'13] .
- $O(n \log n)$ on gridded spaces using convolutions.
[Solomon'+15]

Fast & Scalable Algorithm

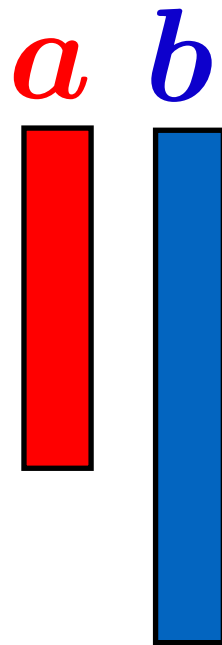
- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$
 $\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$

$$\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$$



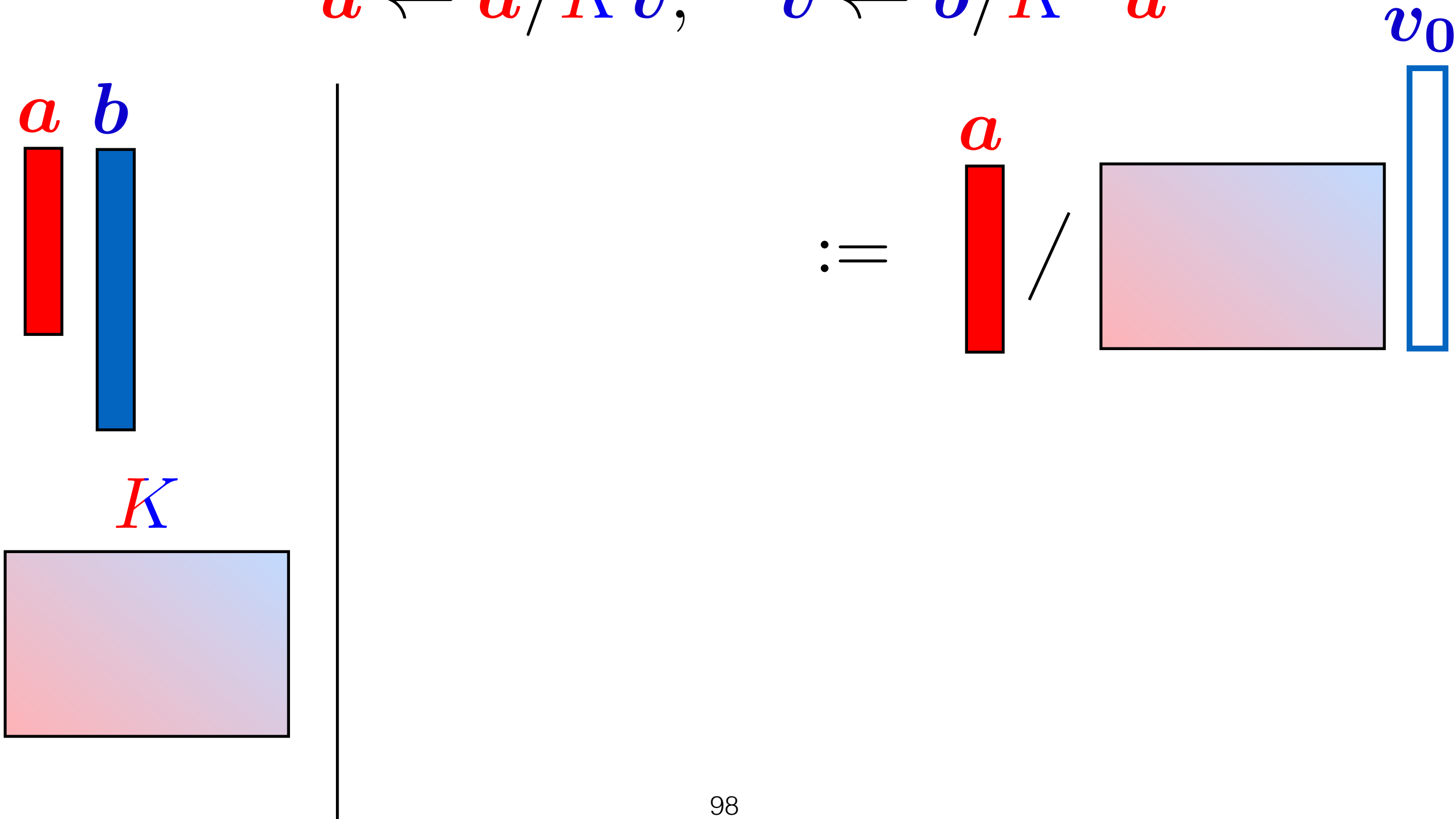
$\mathbf{K}$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$

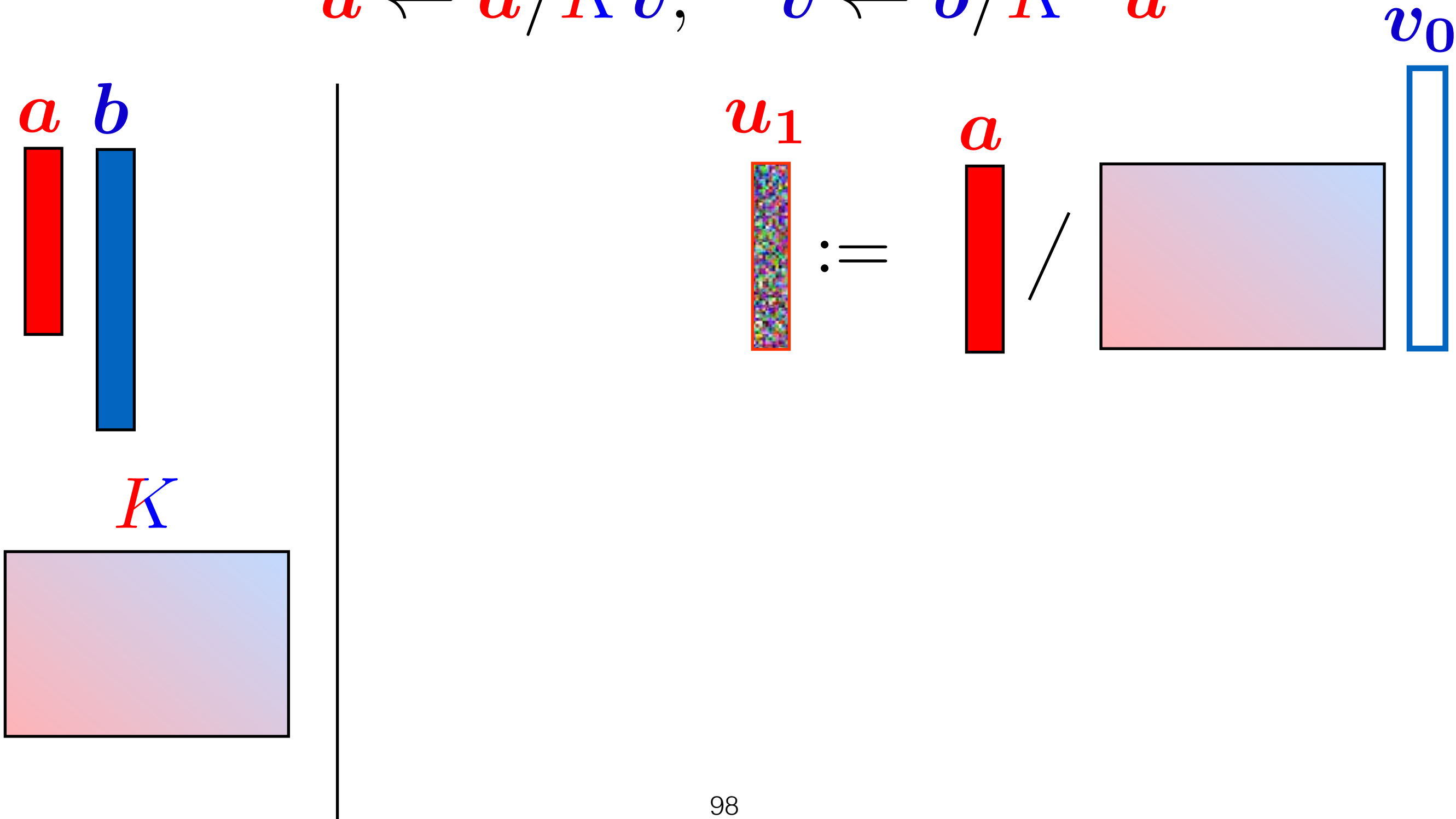
$$\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$

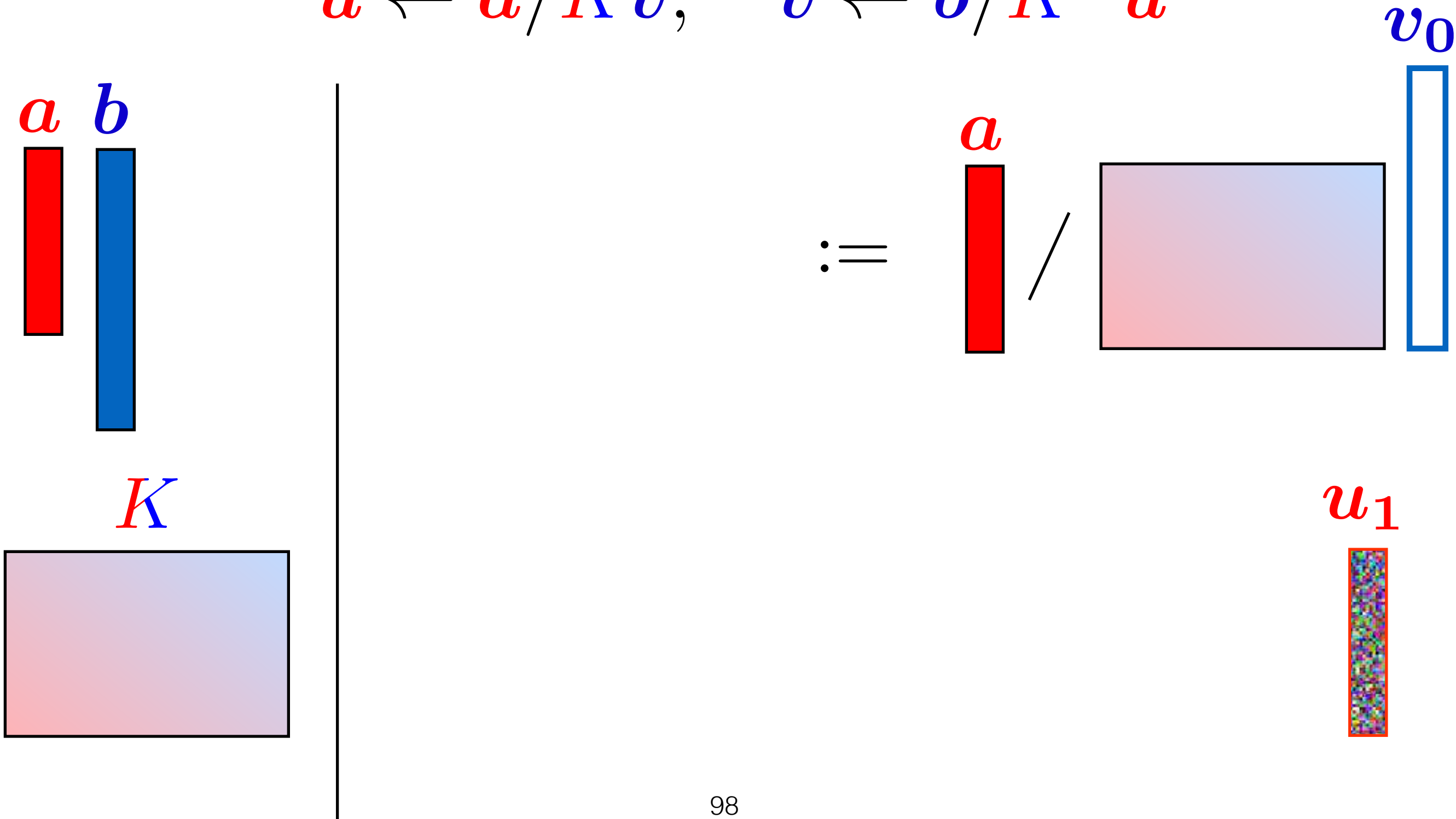
$$\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$

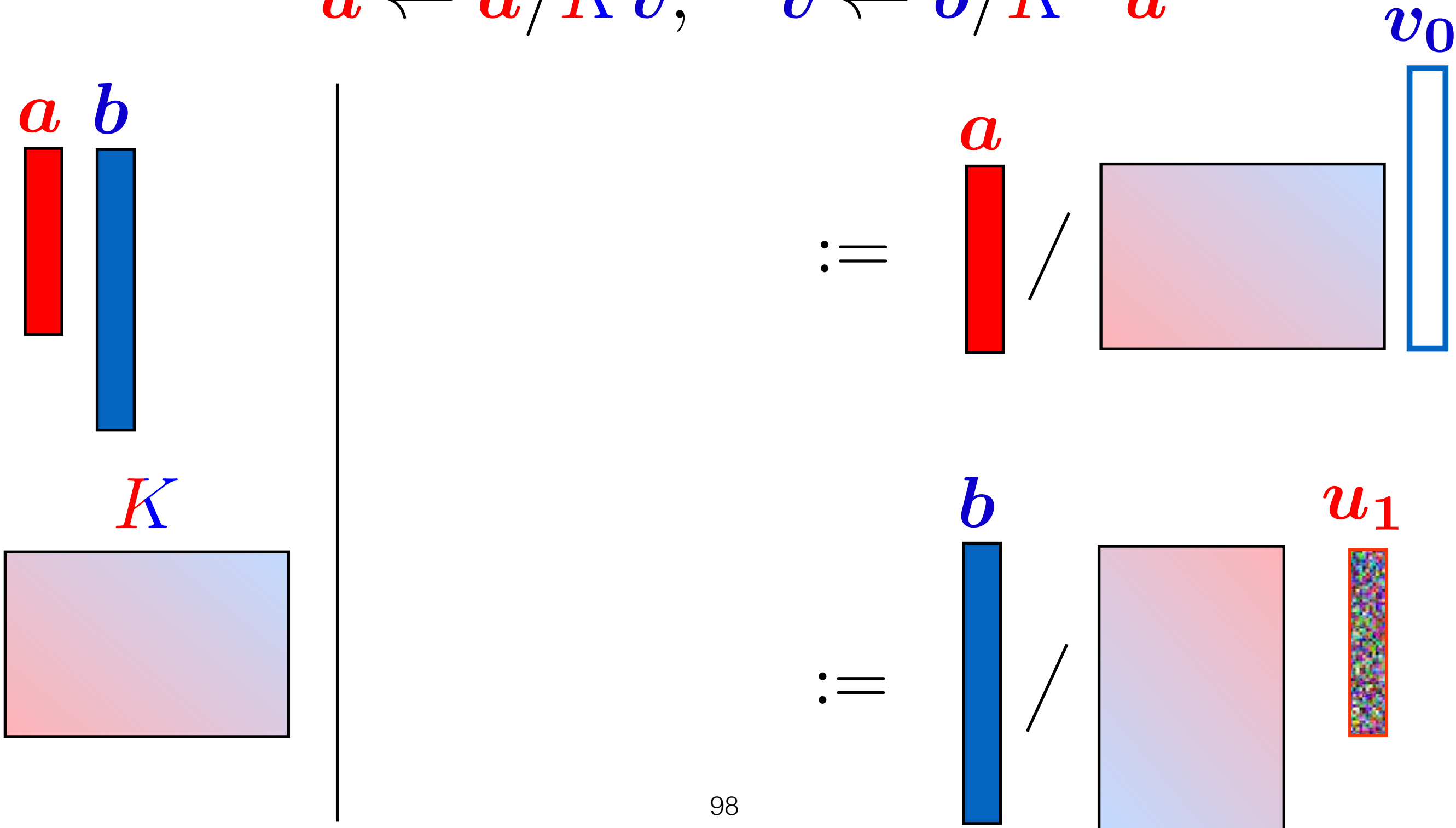
$$\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

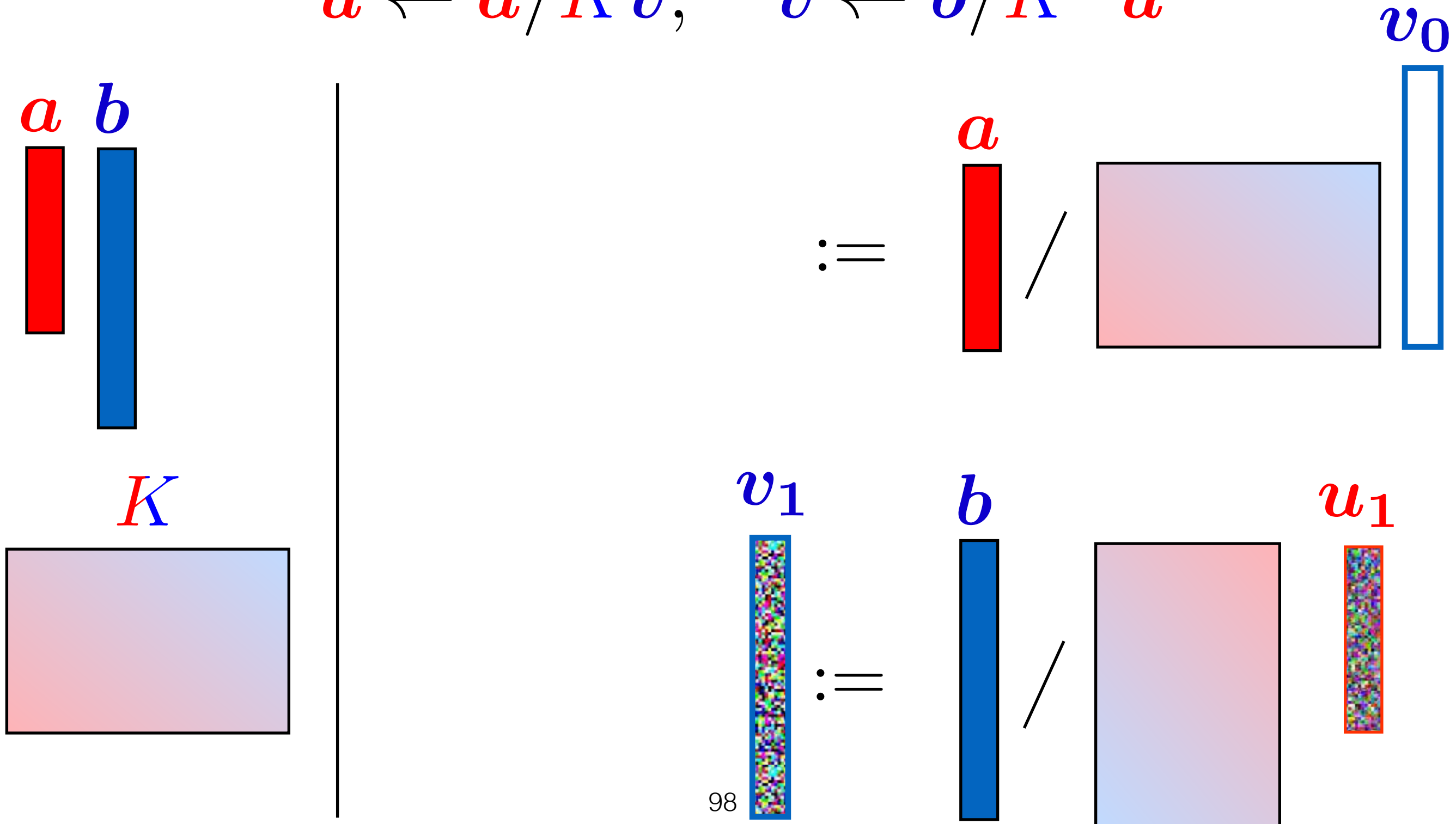
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

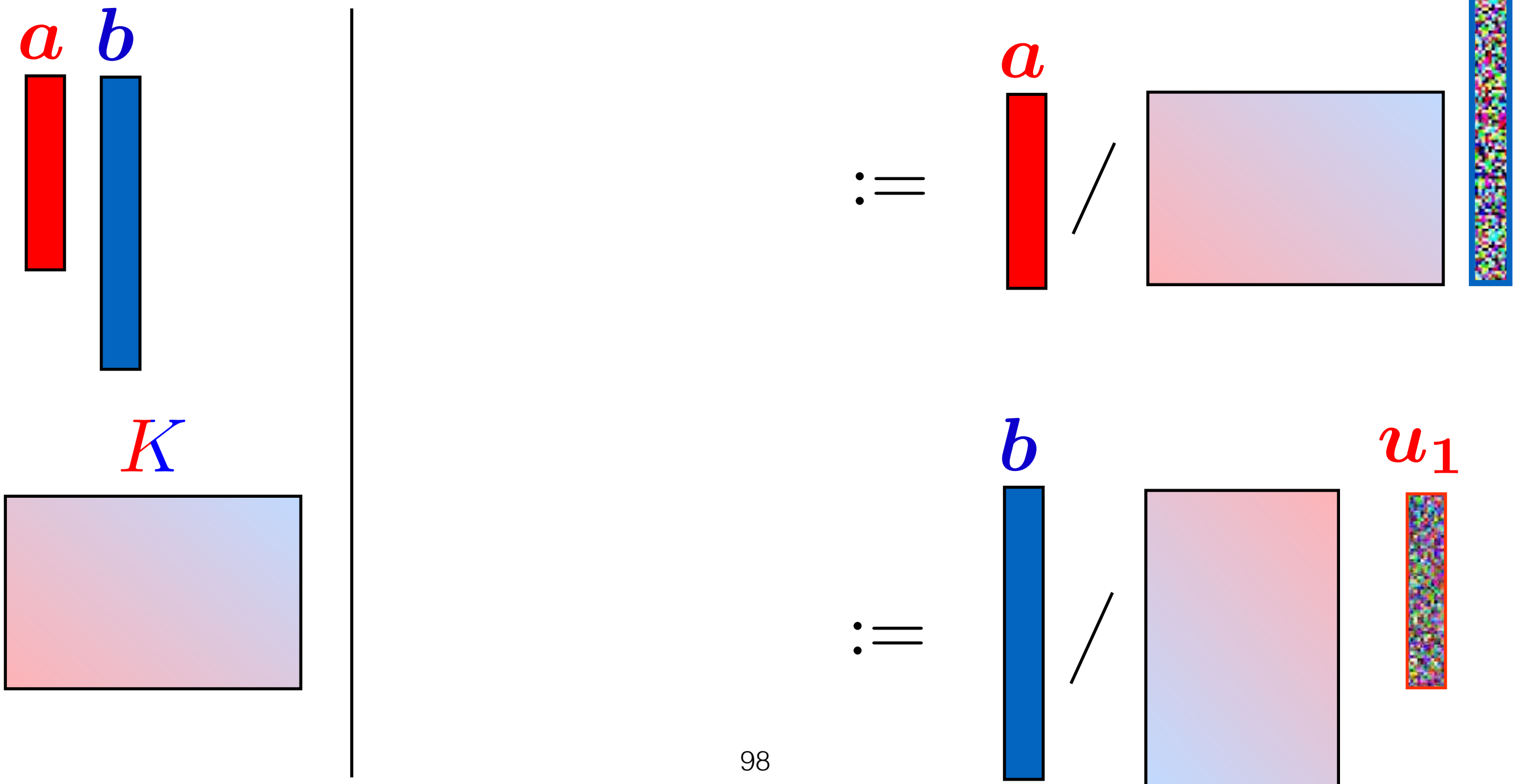
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$

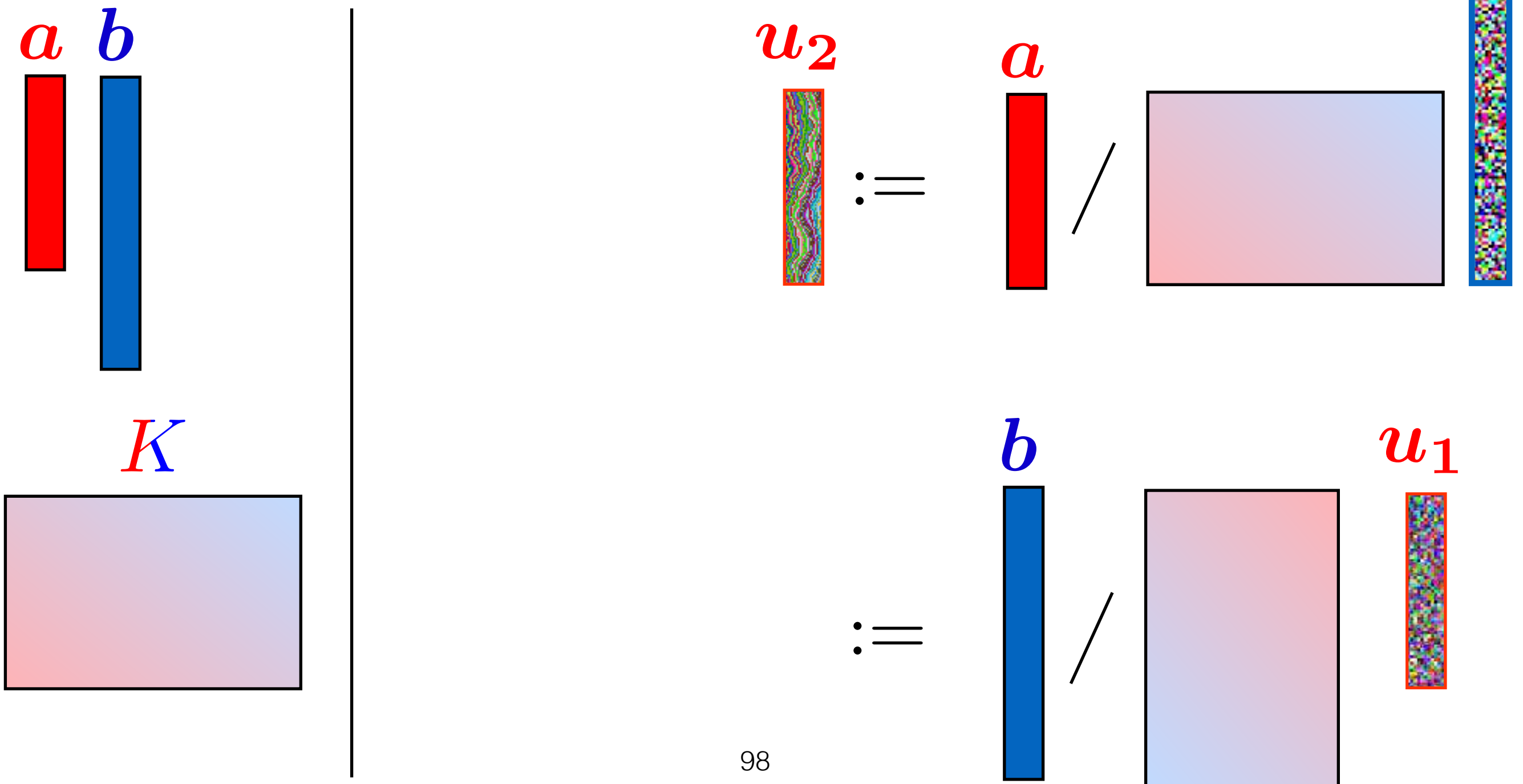
$$\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

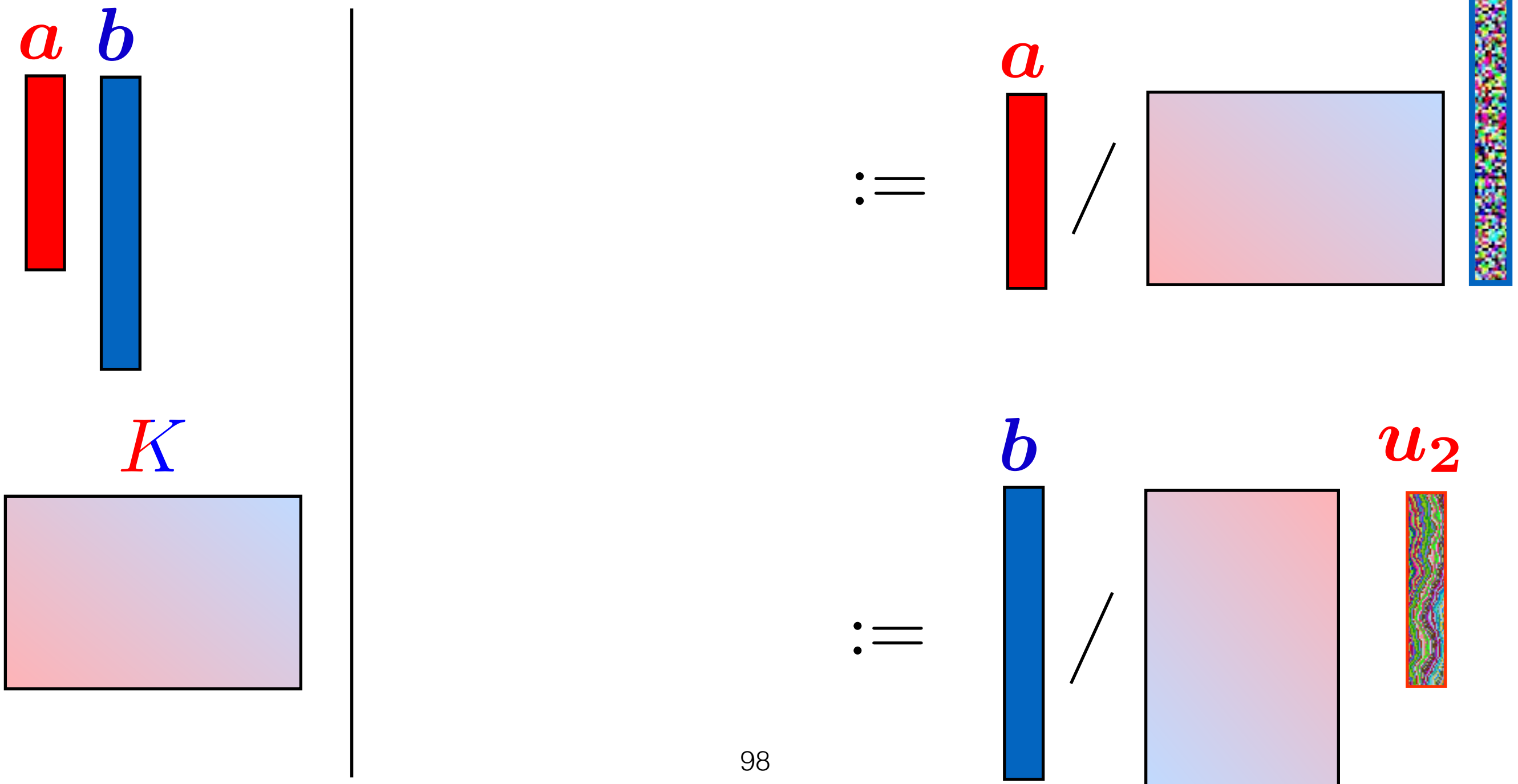
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

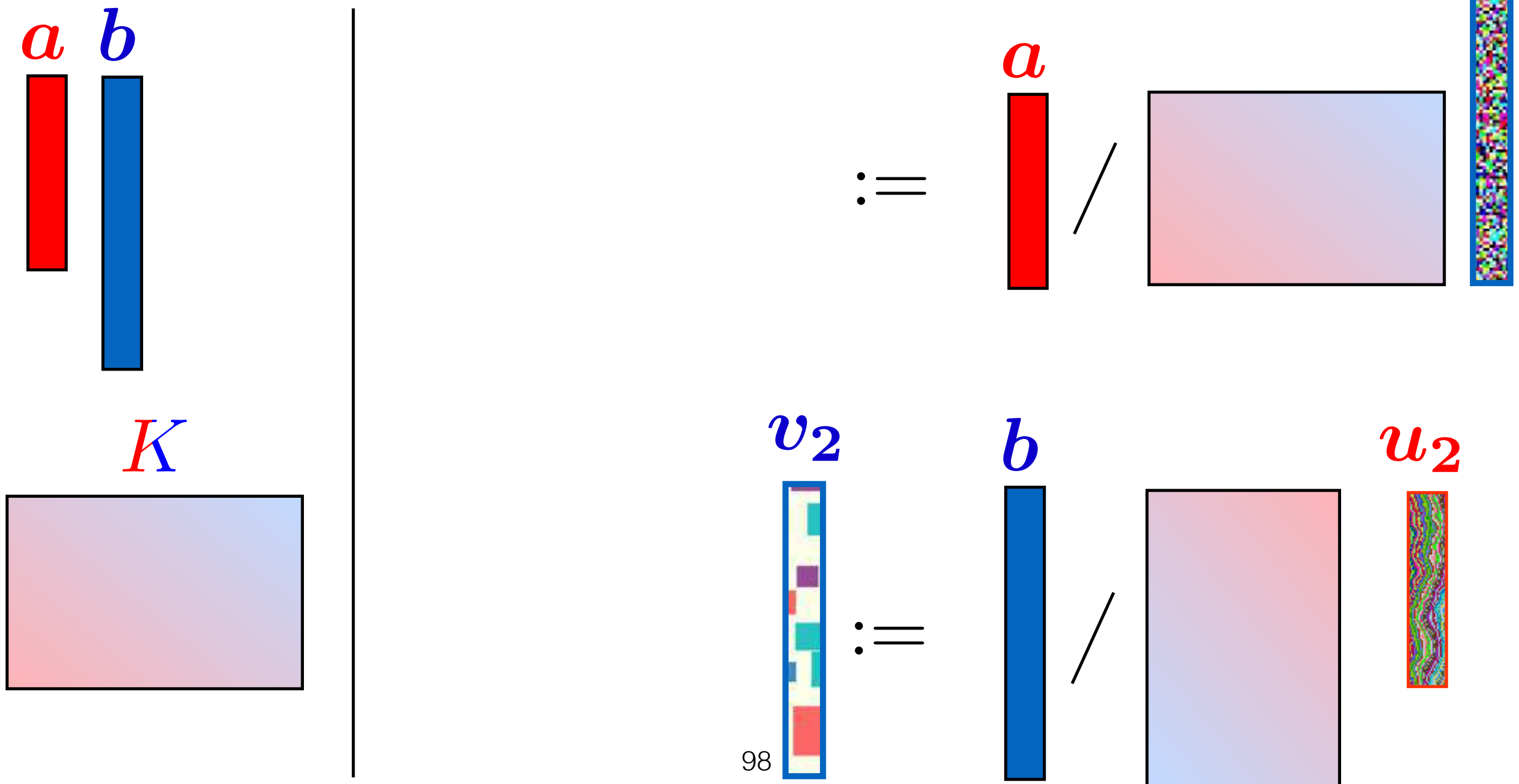
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

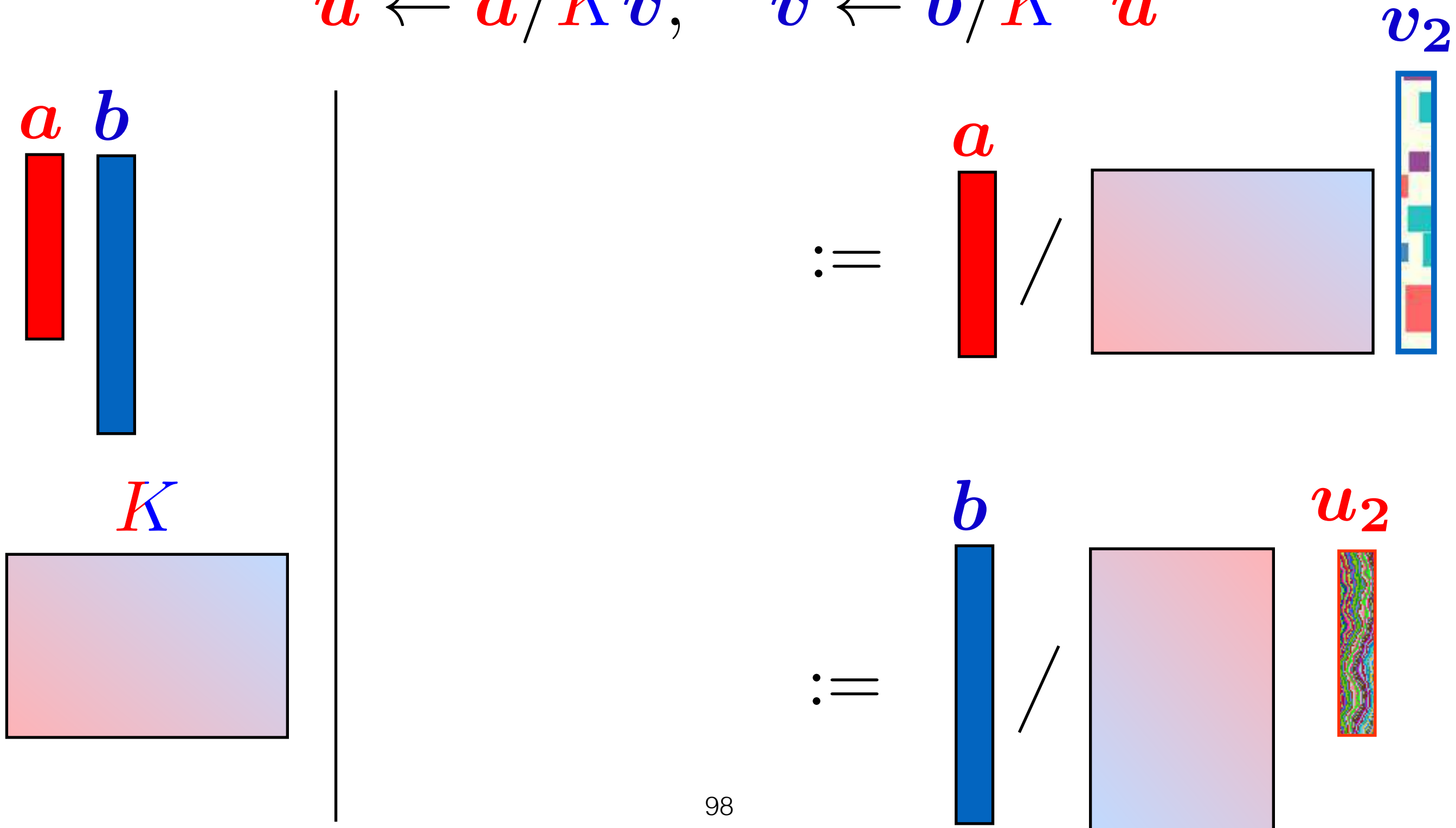
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

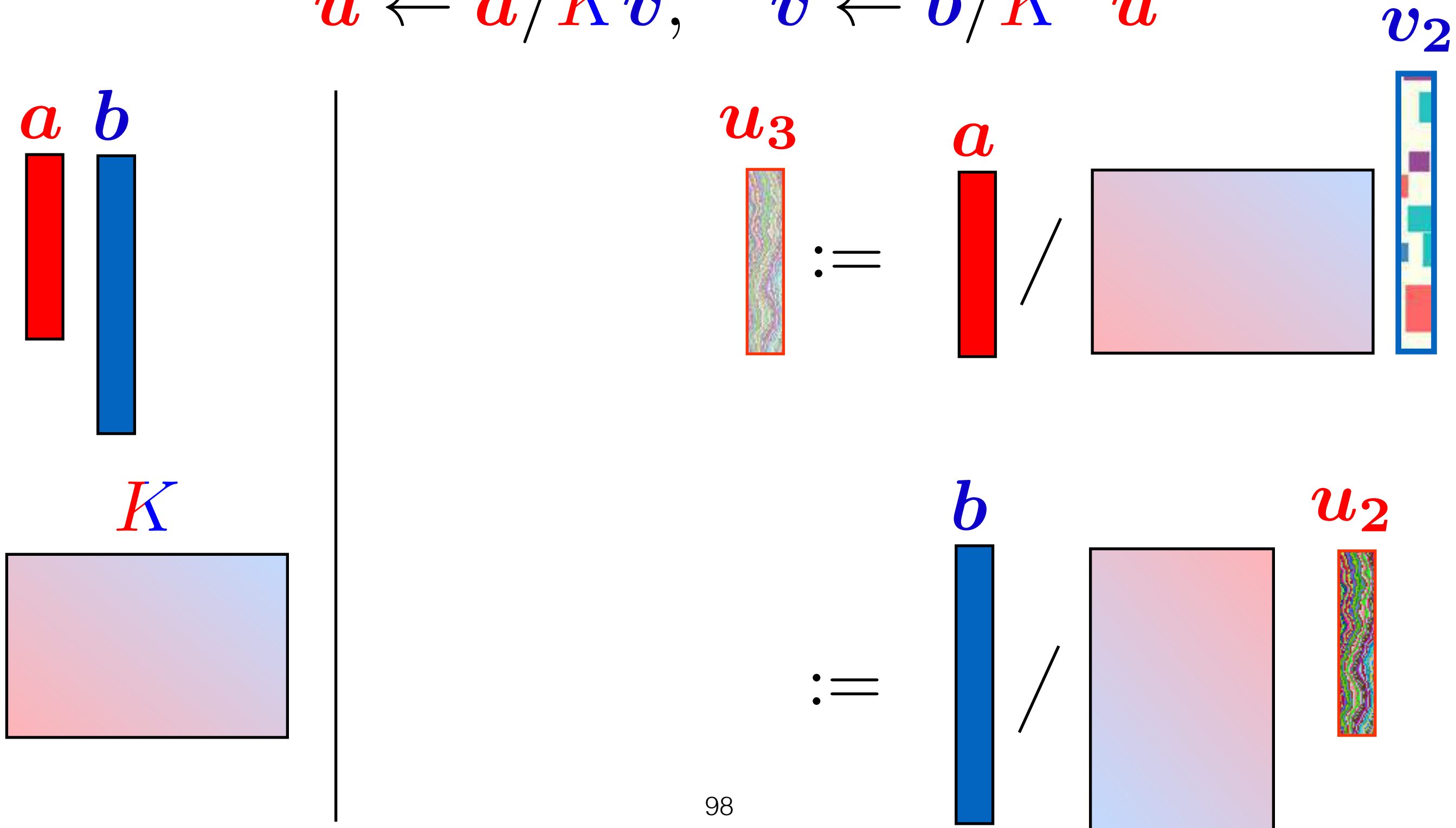
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for $(\mathbf{u}, \mathbf{v})$

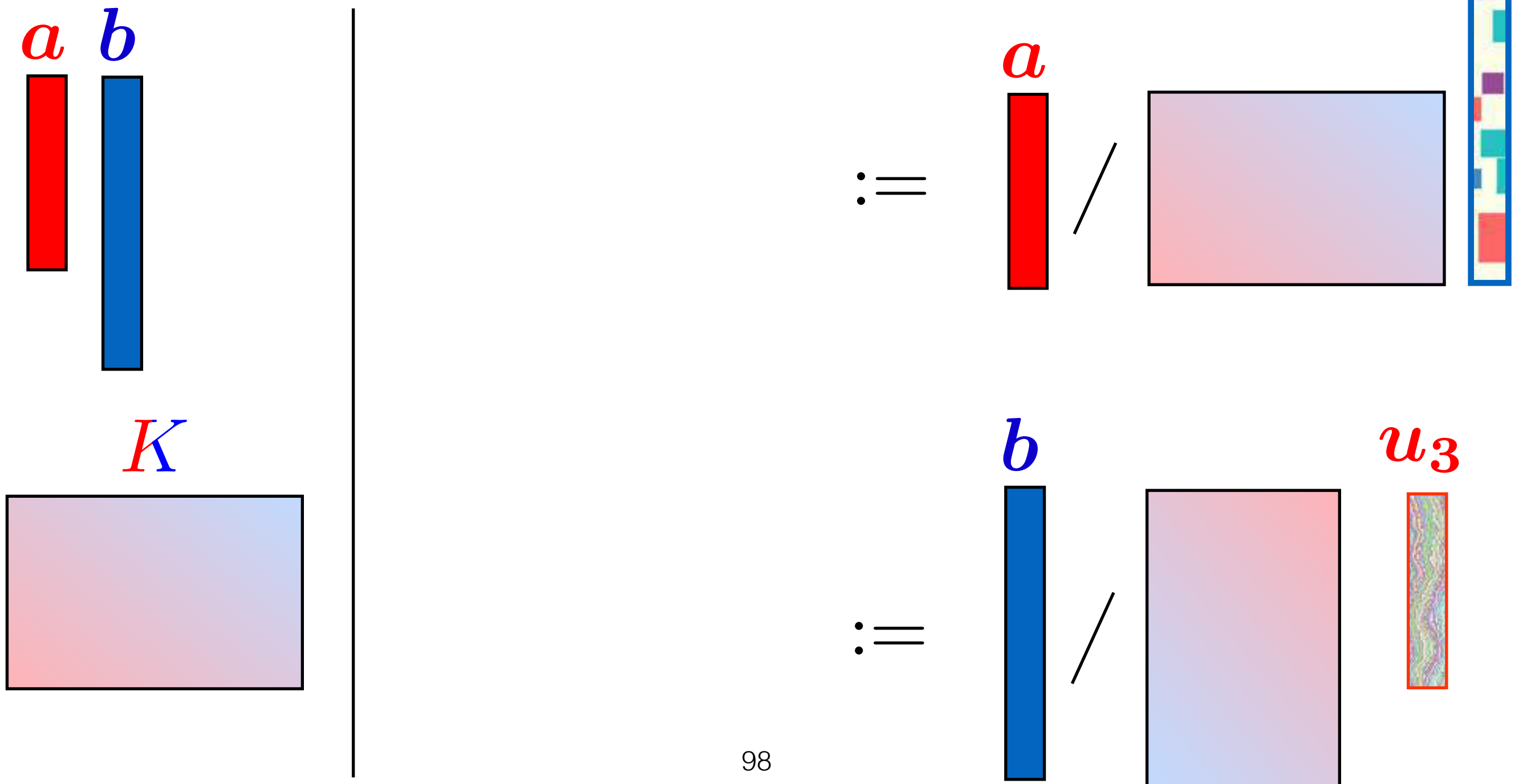
$$\mathbf{u} \leftarrow \mathbf{a} / \mathbf{K} \mathbf{v}, \quad \mathbf{v} \leftarrow \mathbf{b} / \mathbf{K}^T \mathbf{u}$$



Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations for (u, v)

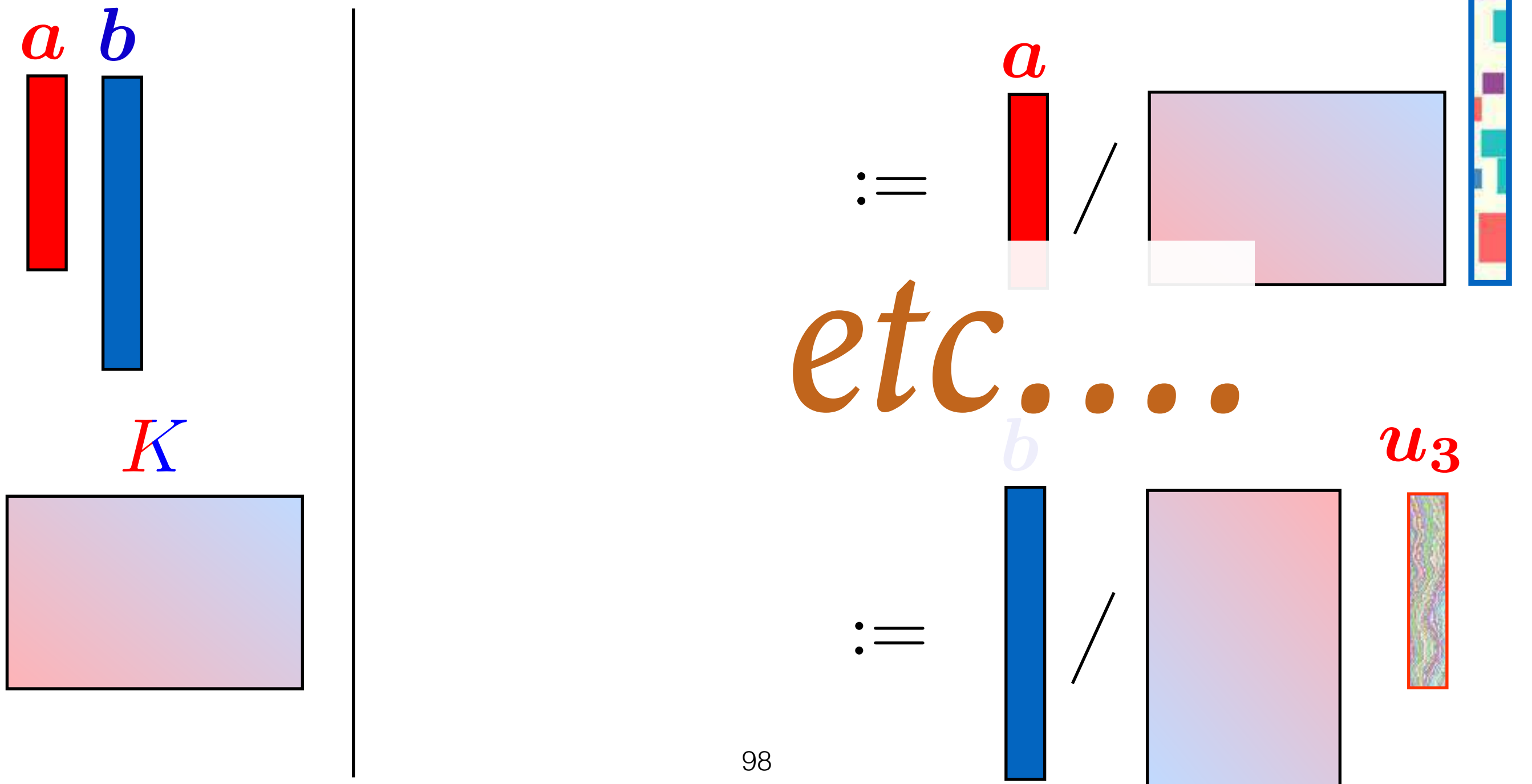
$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



Fast & Scalable Algorithm

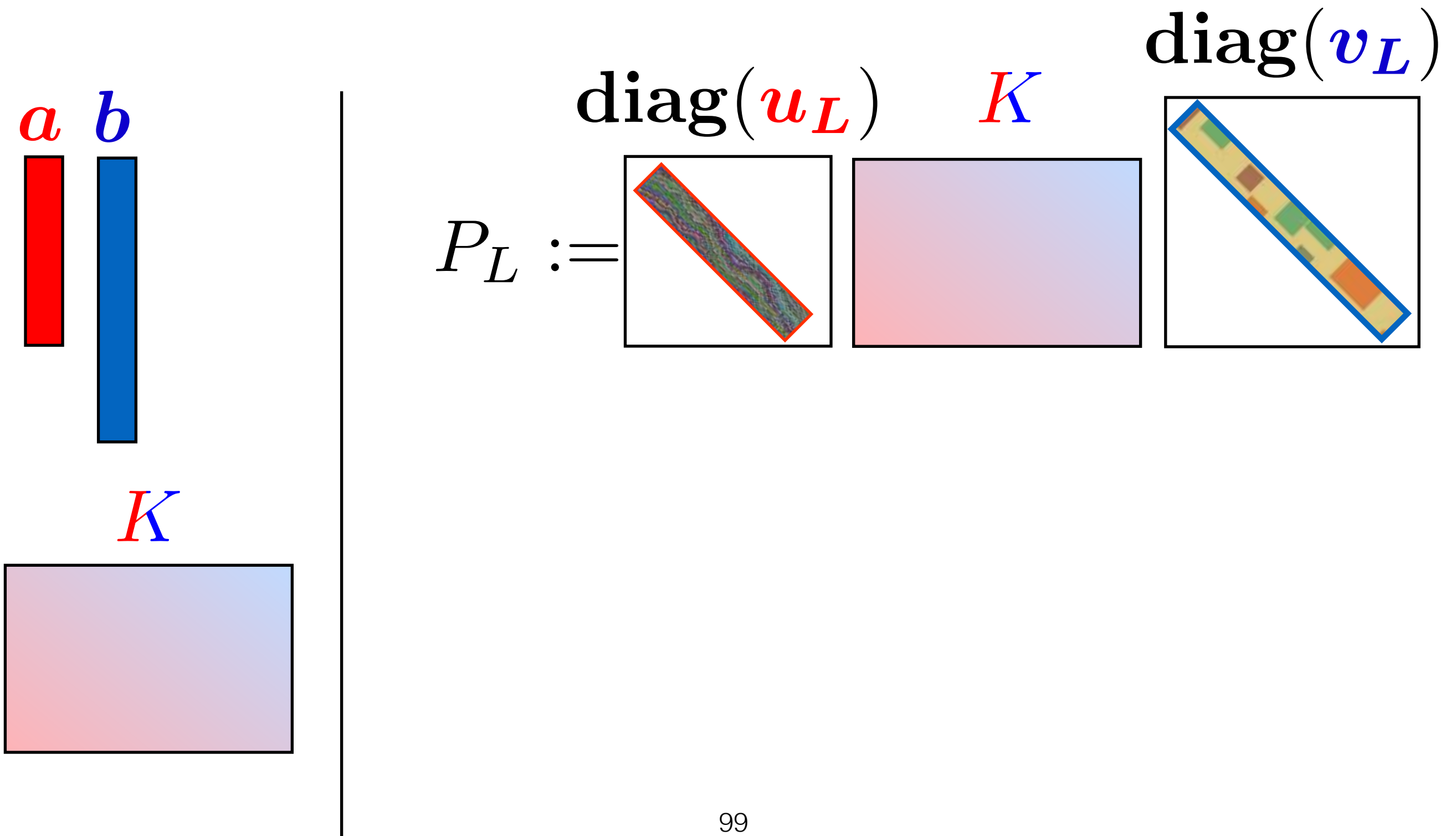
- [Sinkhorn'64] fixed-point iterations for (u, v)

$$u \leftarrow a / K v, \quad v \leftarrow b / K^T u$$



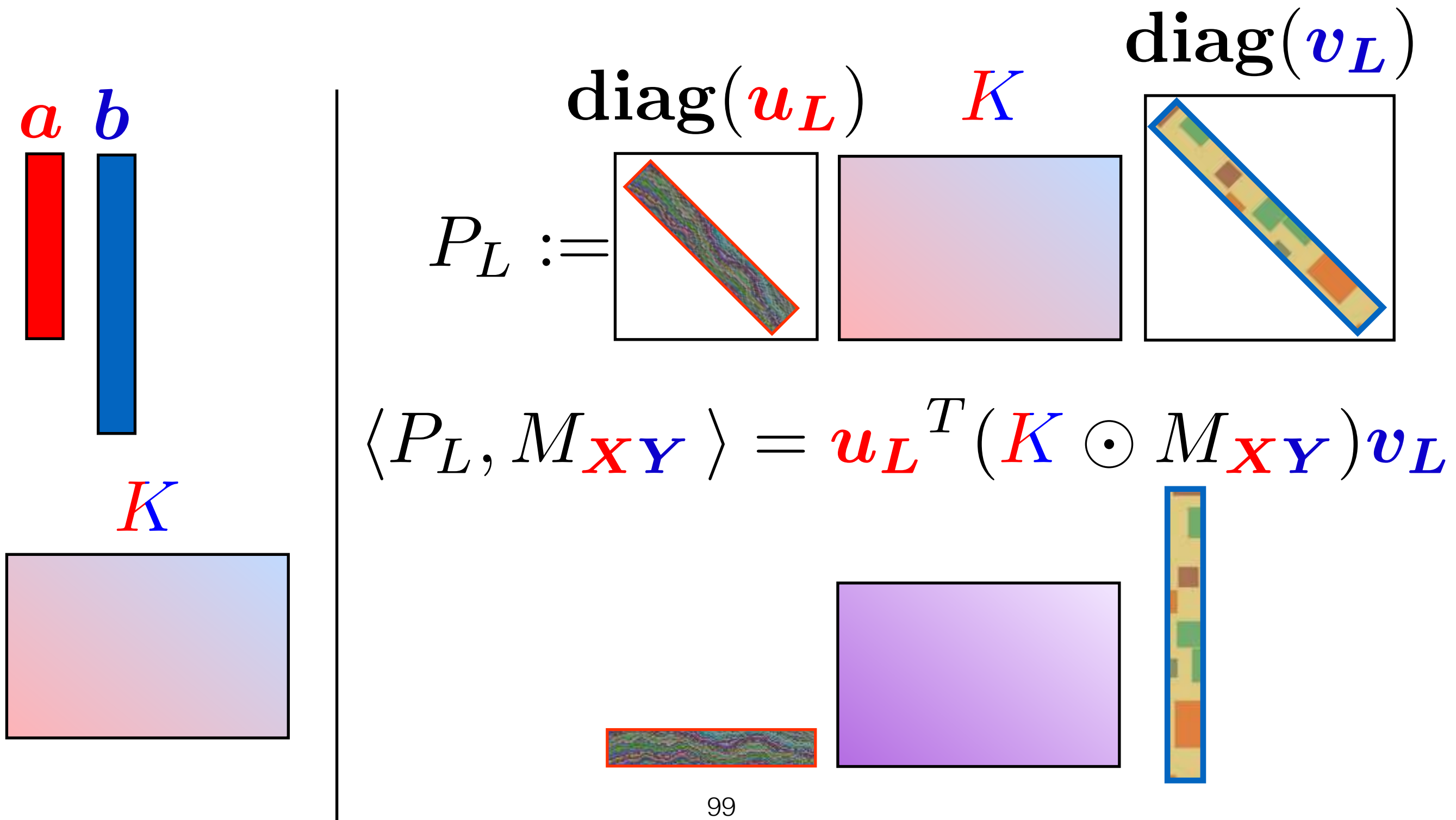
Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations.



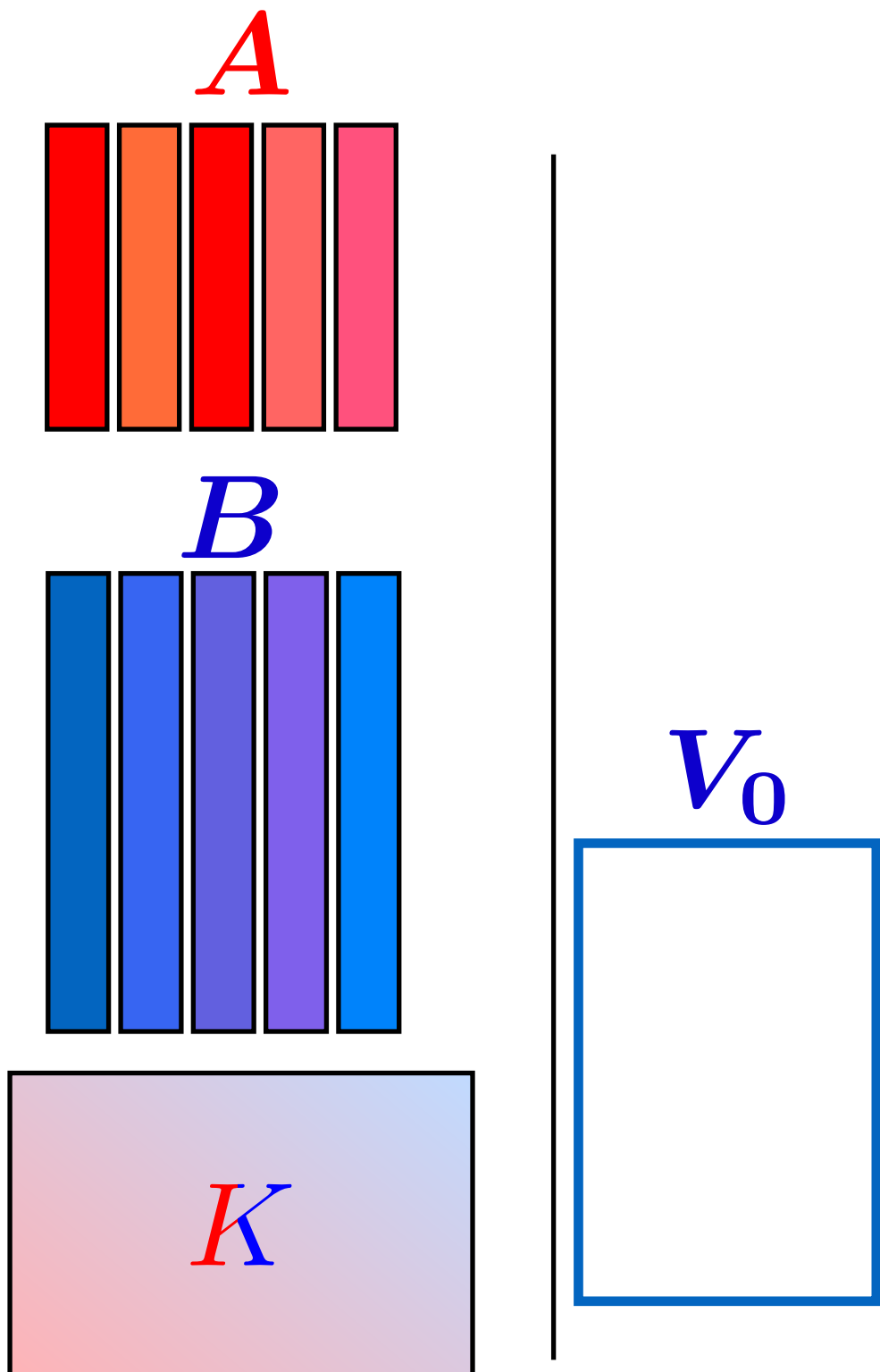
Fast & Scalable Algorithm

- [Sinkhorn'64] fixed-point iterations.



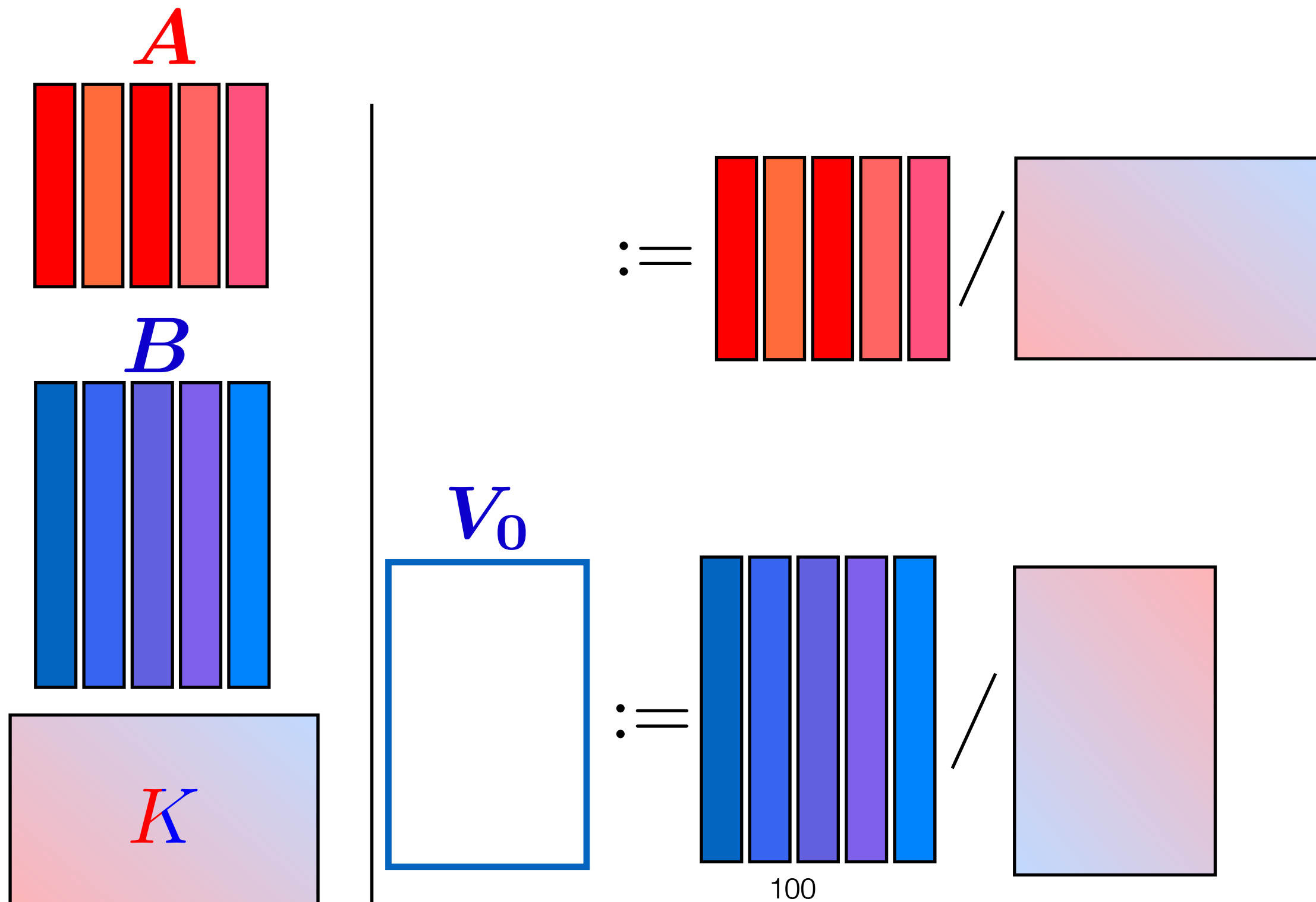
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



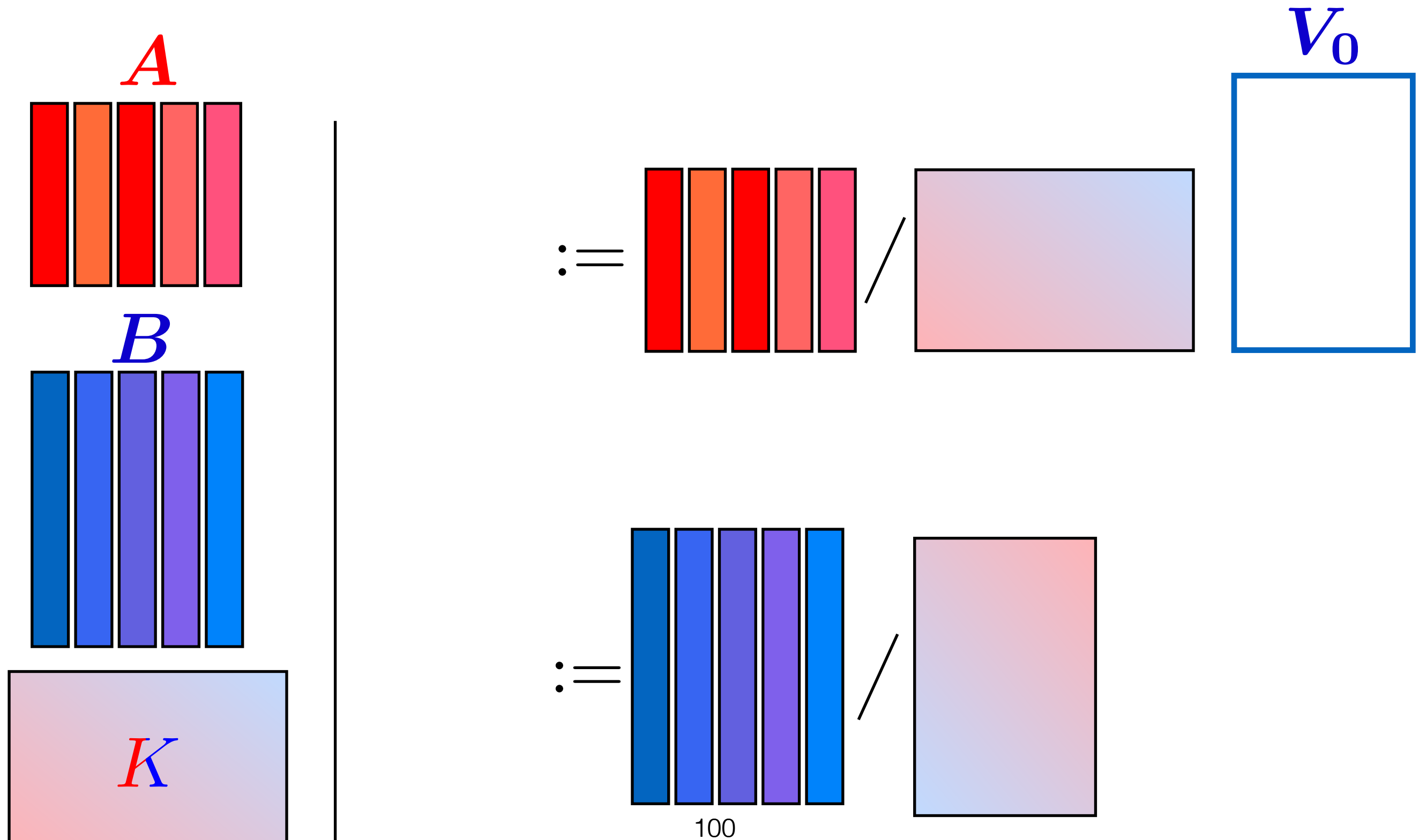
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



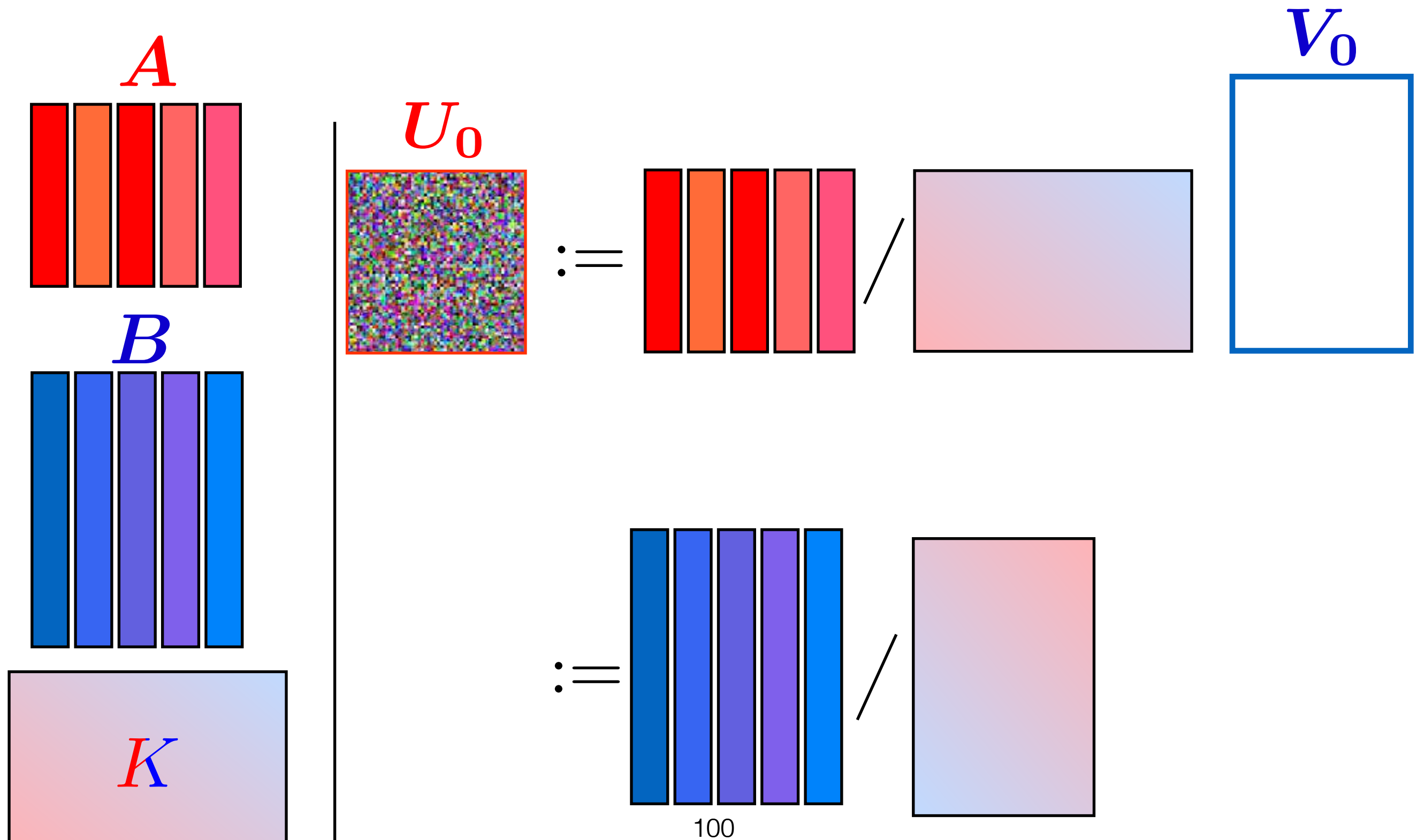
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



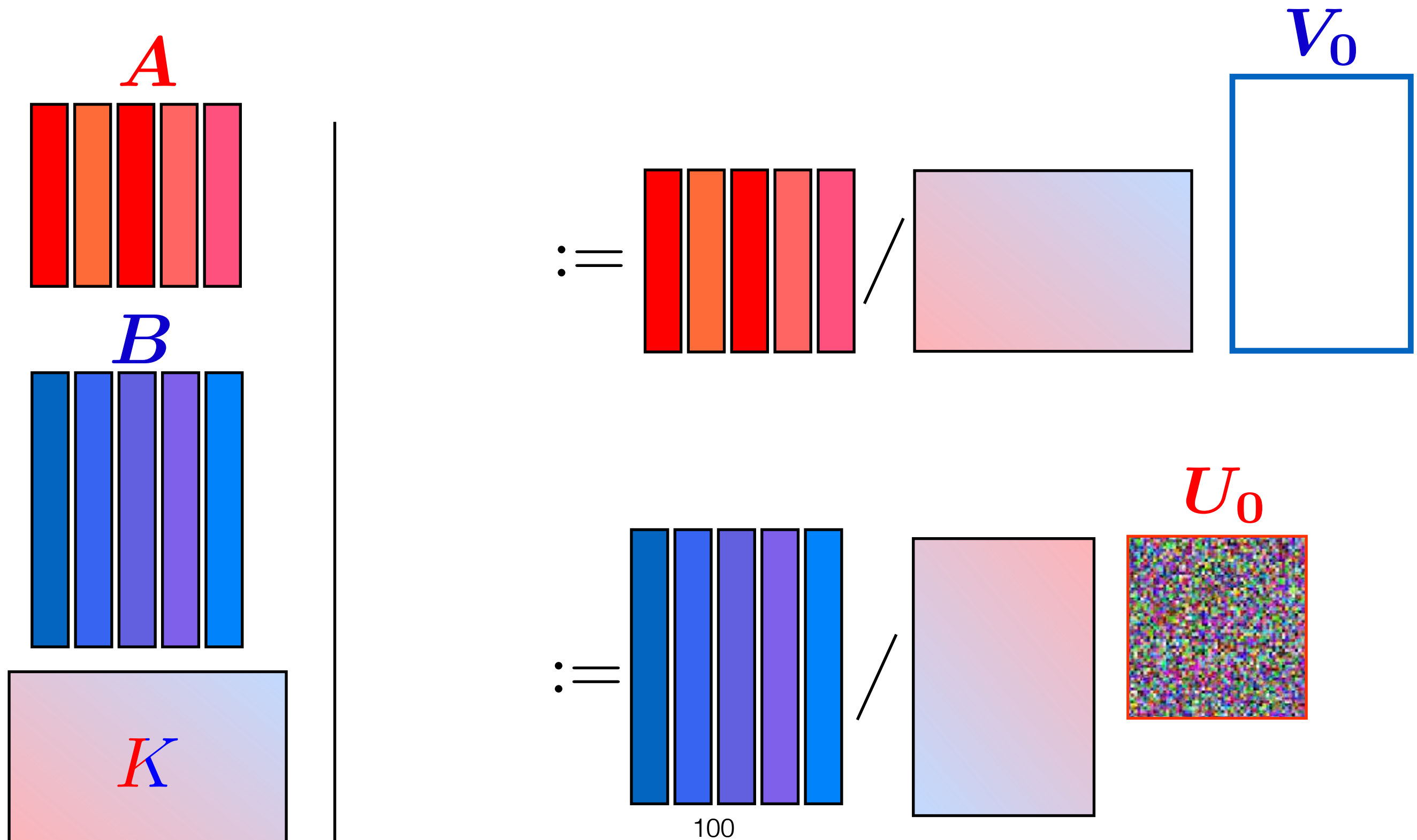
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



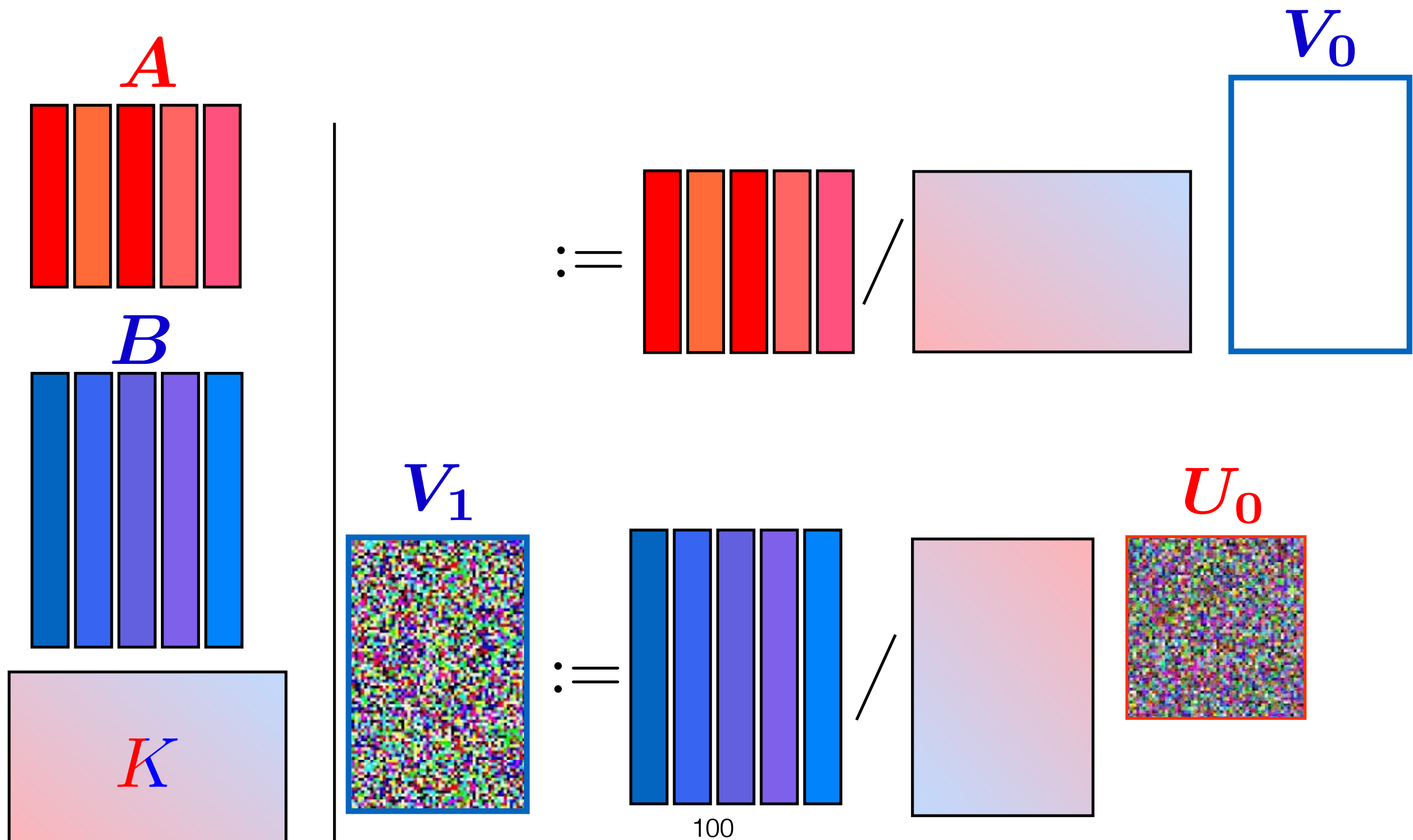
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



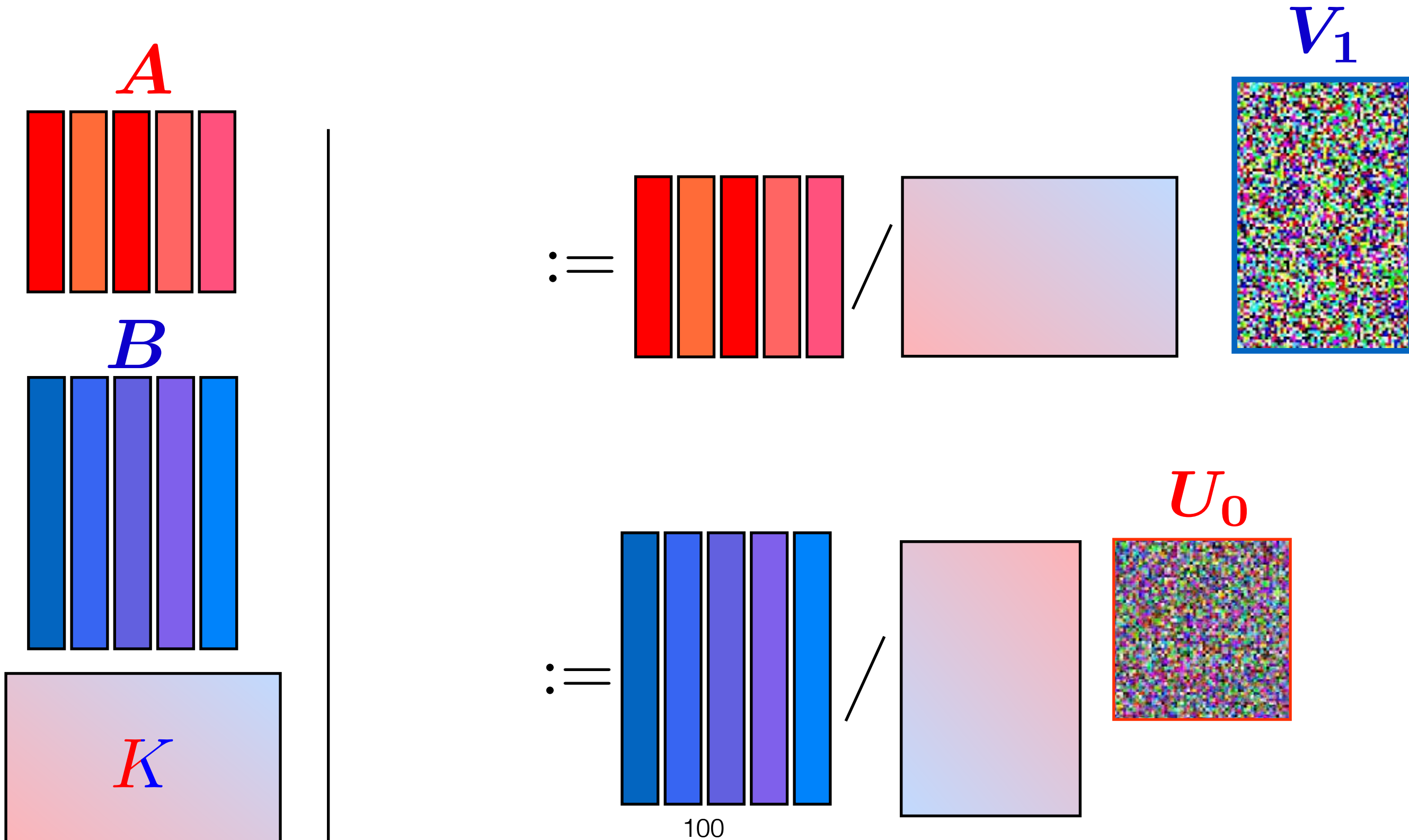
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



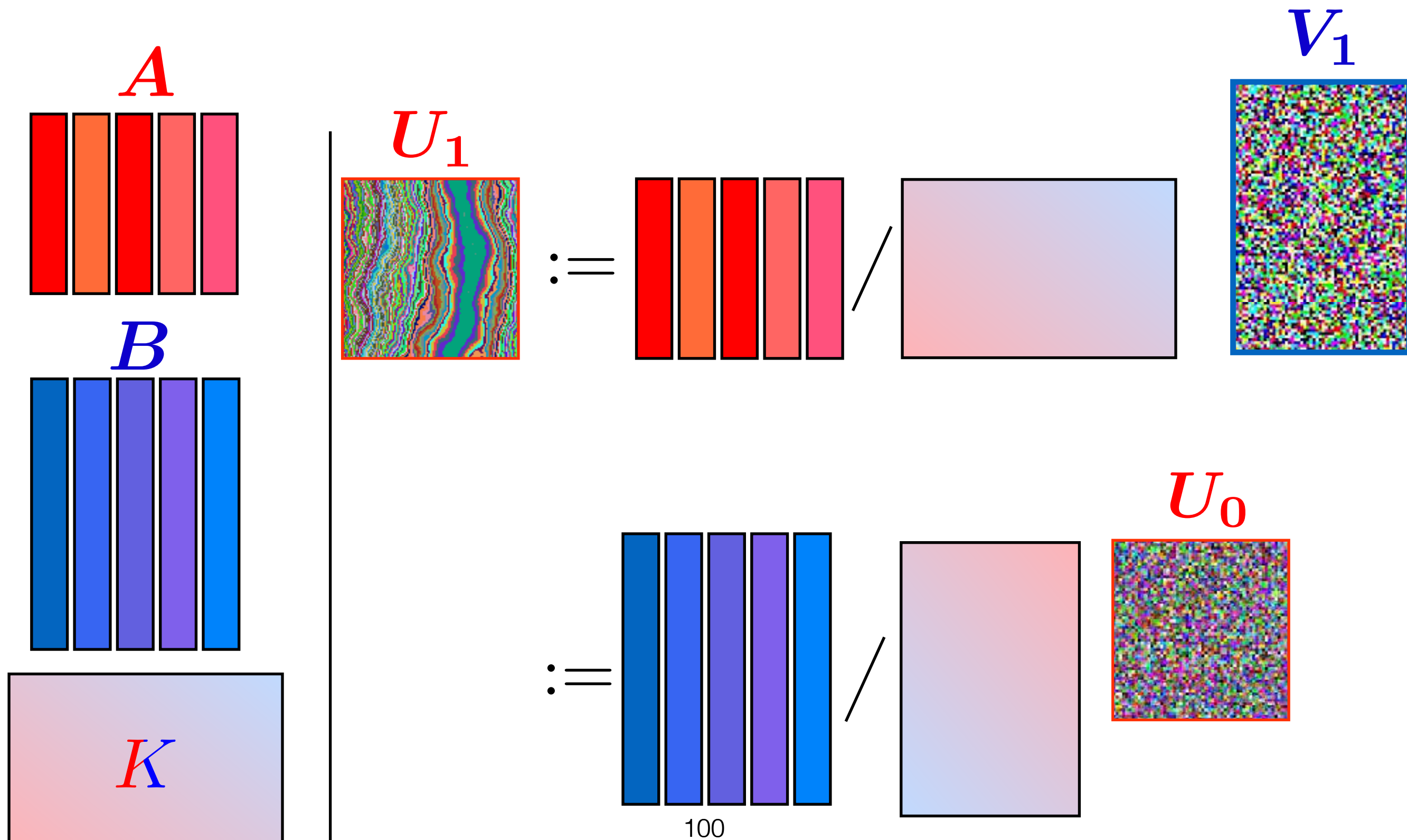
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



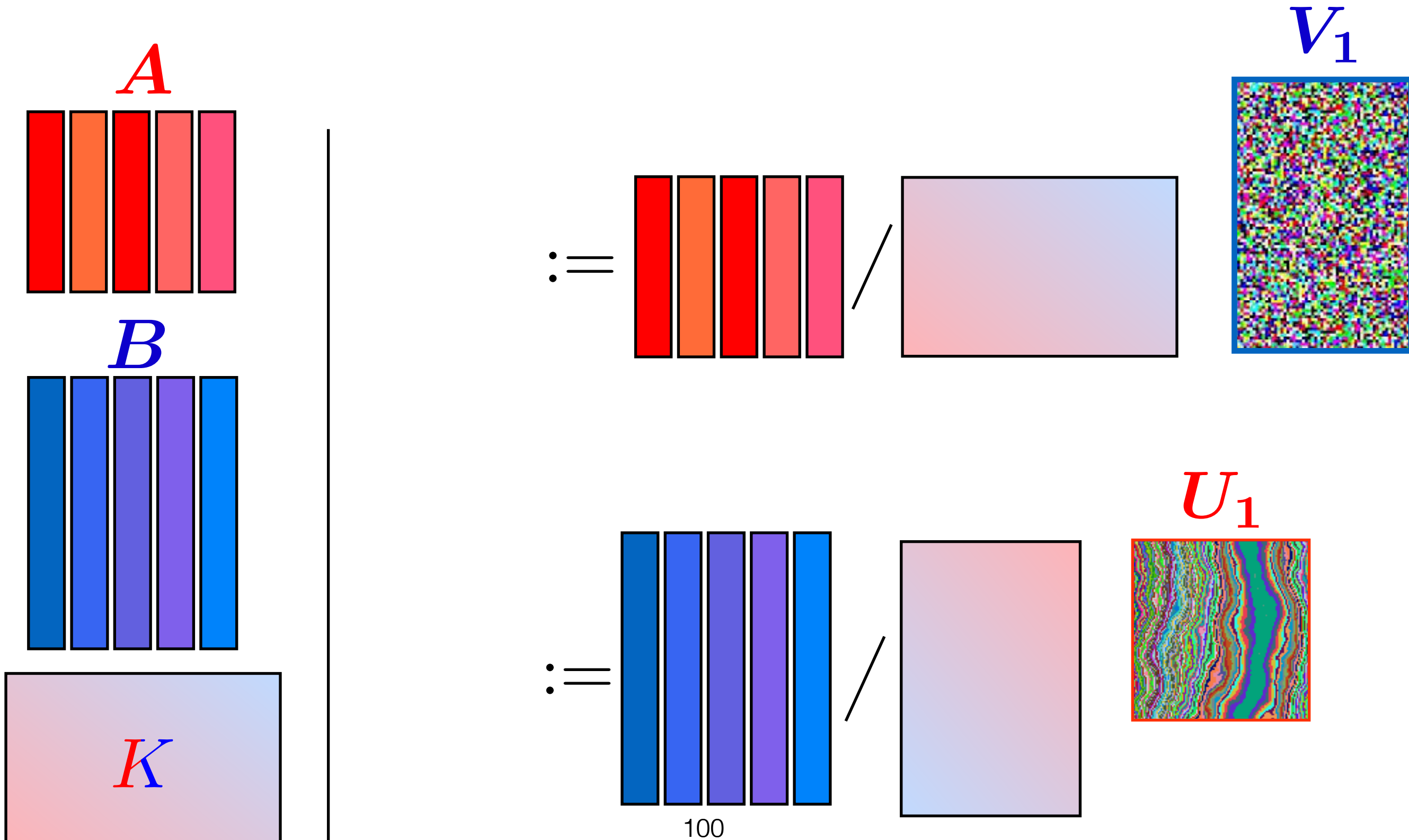
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



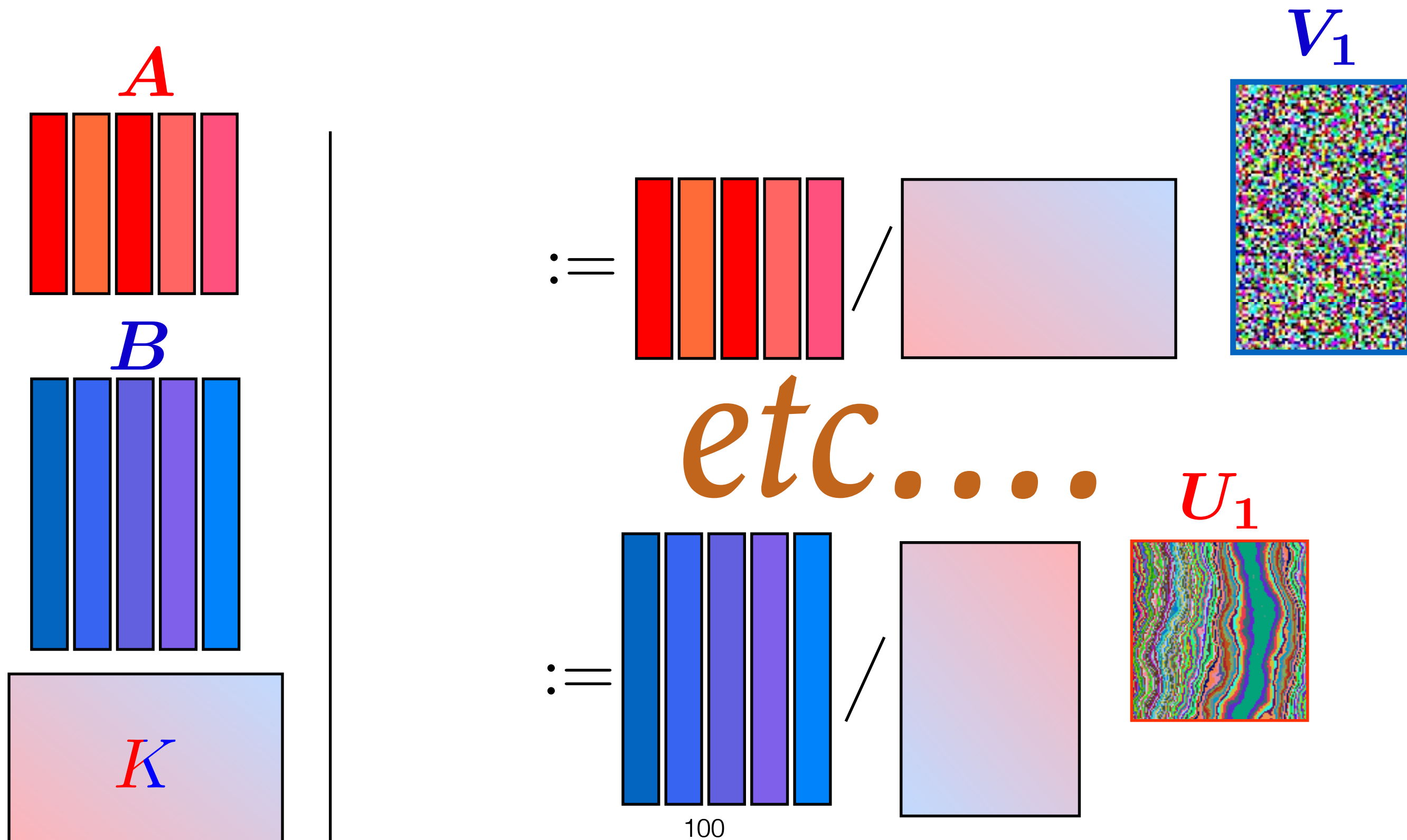
Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations

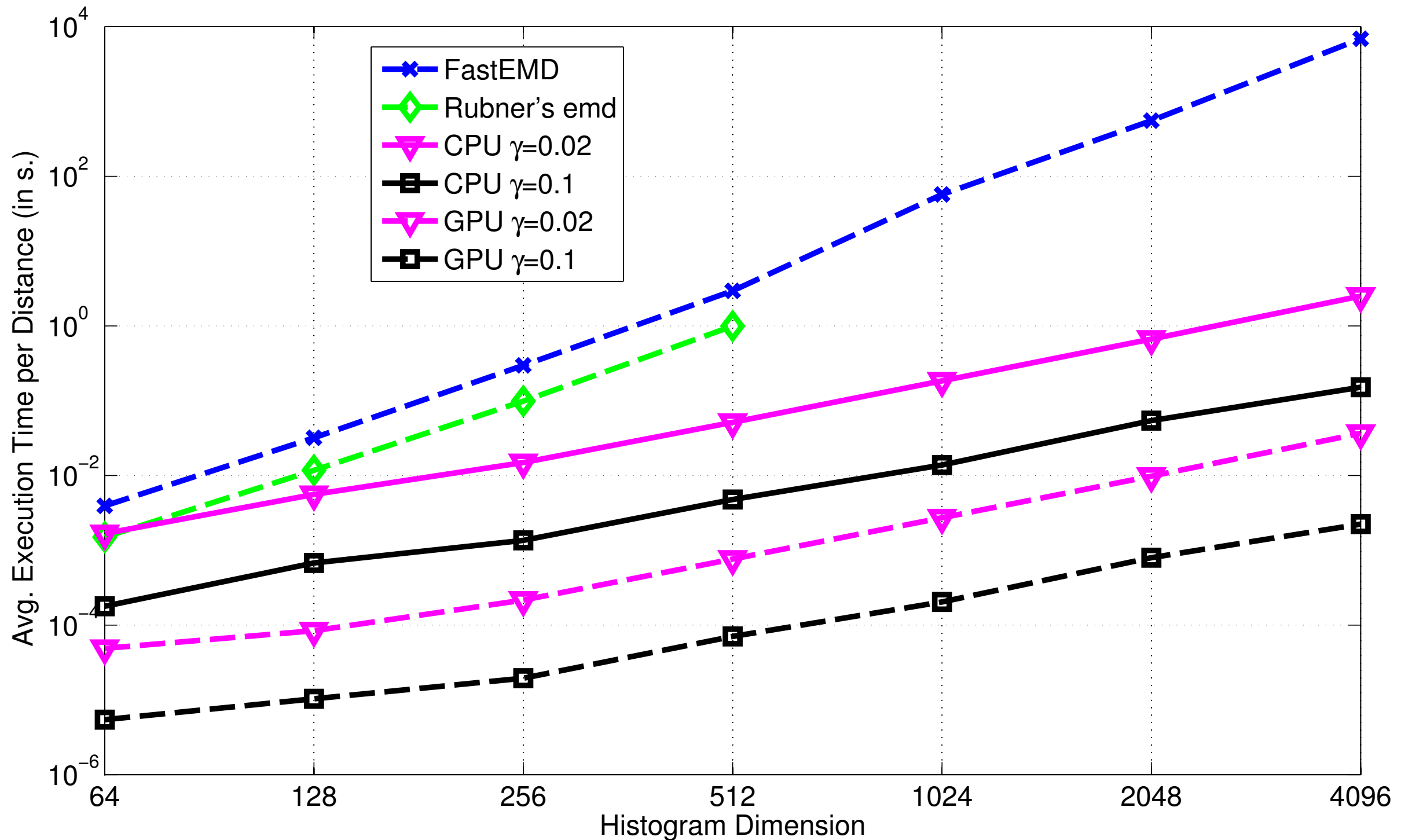


Also embarrassingly parallel

- [Sinkhorn'64] with *matrix* fixed-point iterations



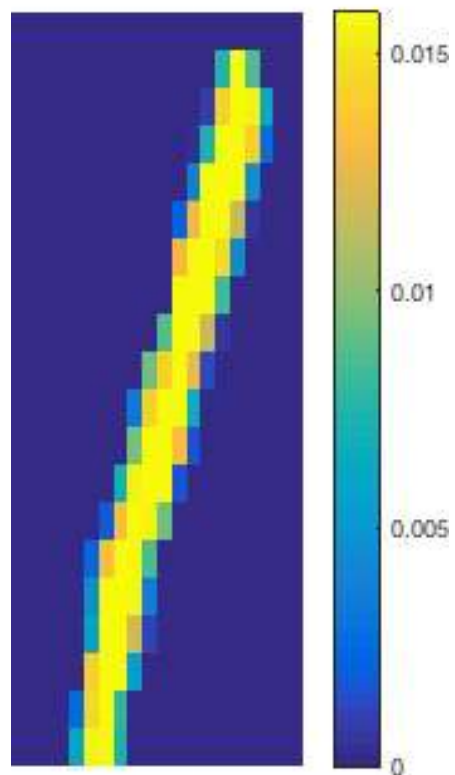
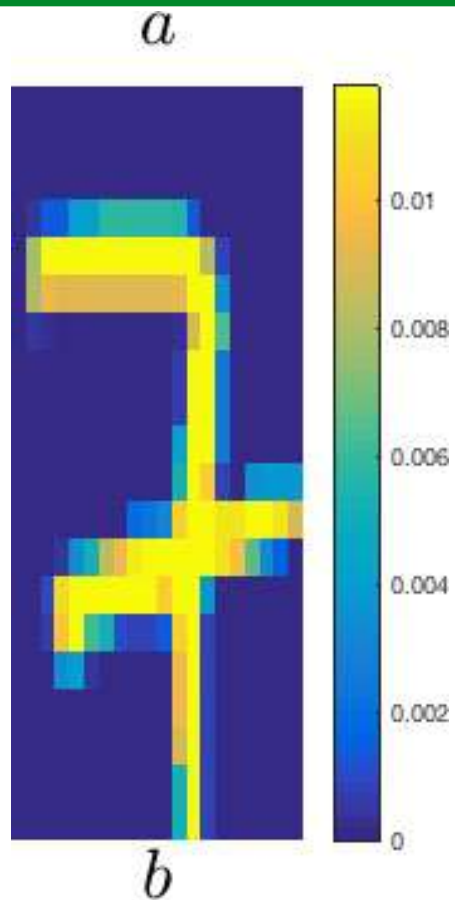
Very Fast EMD Approx. Solver



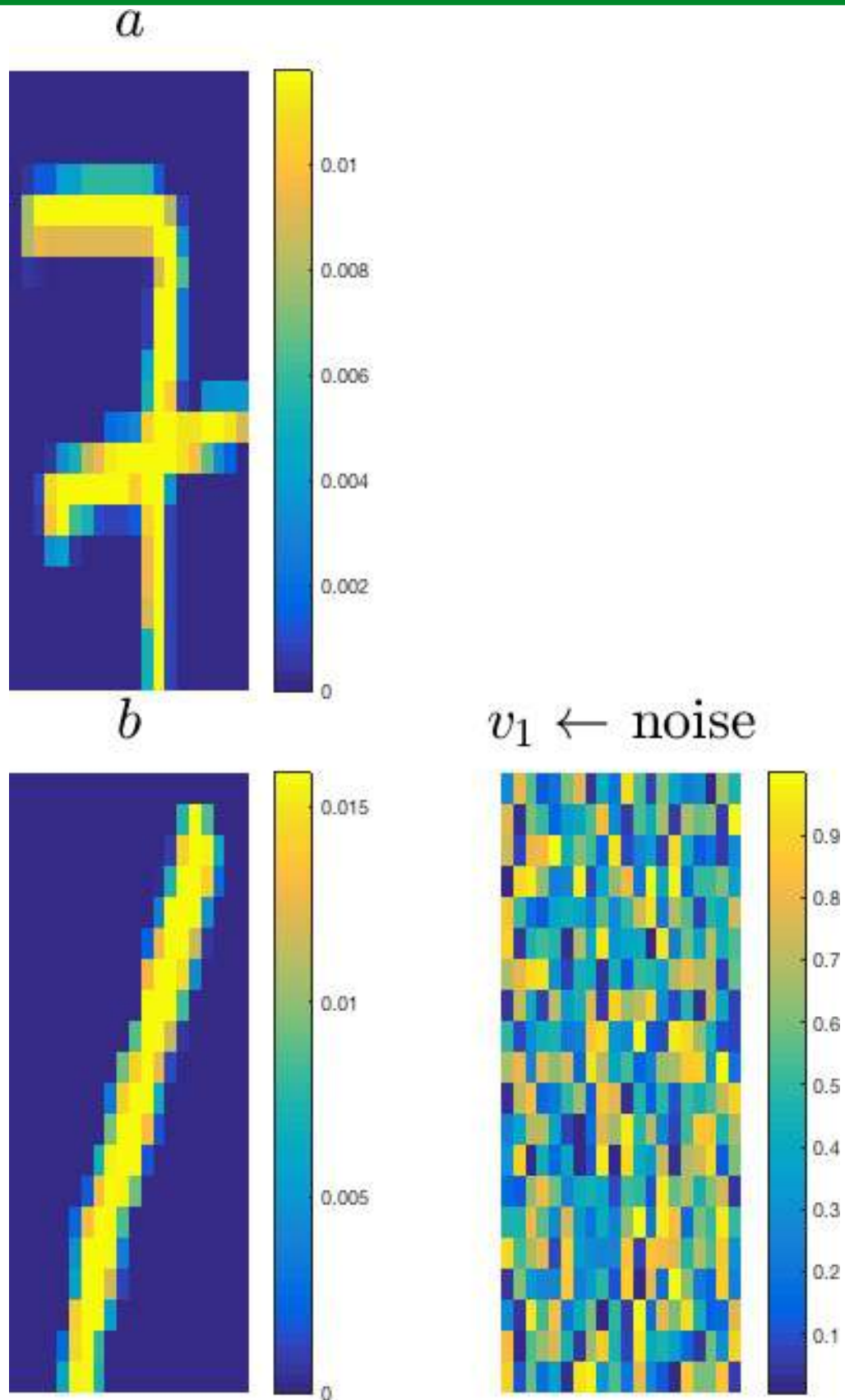
Setup. $(\Omega, \mathbf{D})$ is a random graph with shortest path metric, histograms sampled uniformly on simplex, Sinkhorn tolerance 10^{-2} .

Very Fast EMD Approx. Solver

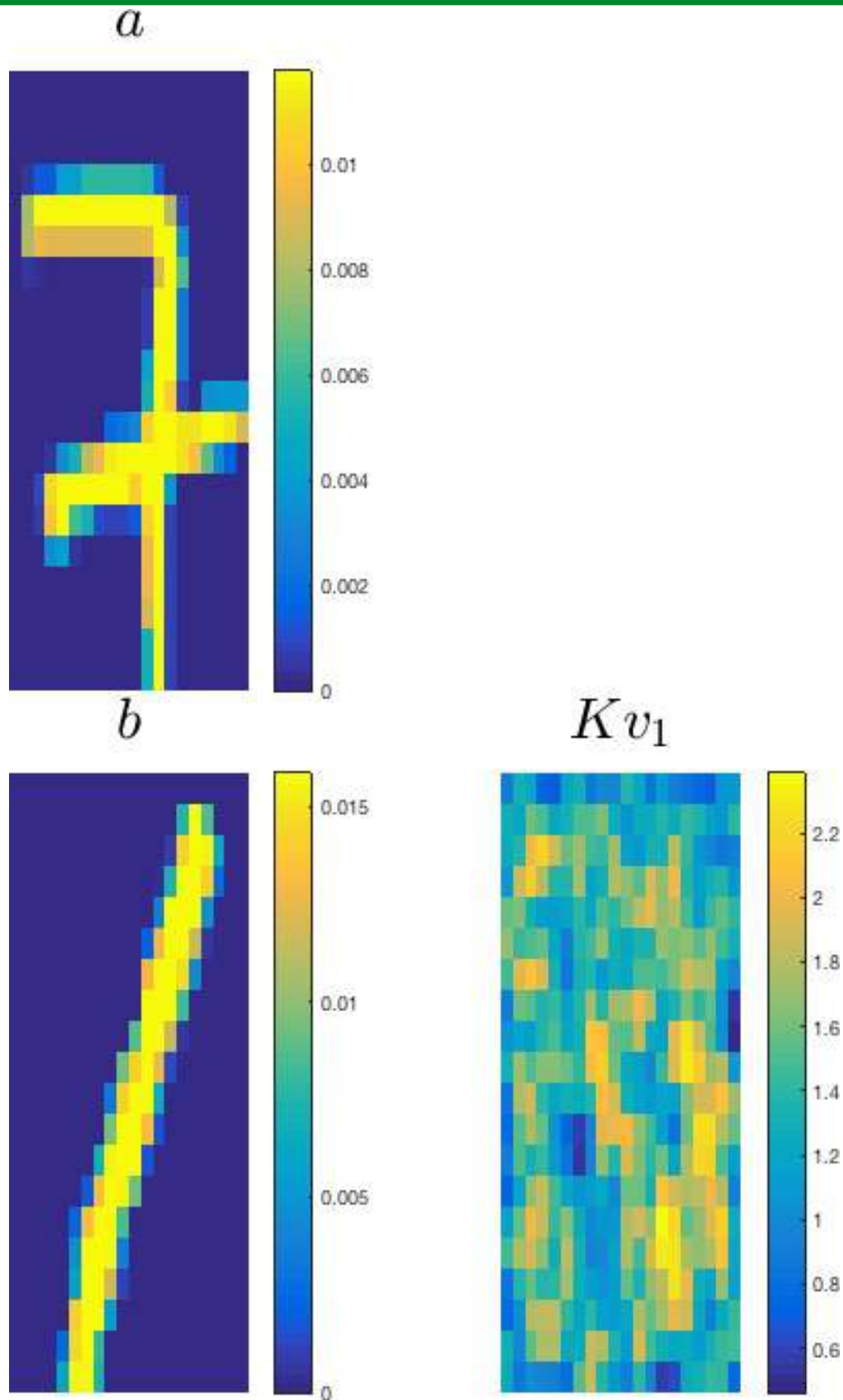
Very Fast EMD Approx. Solver



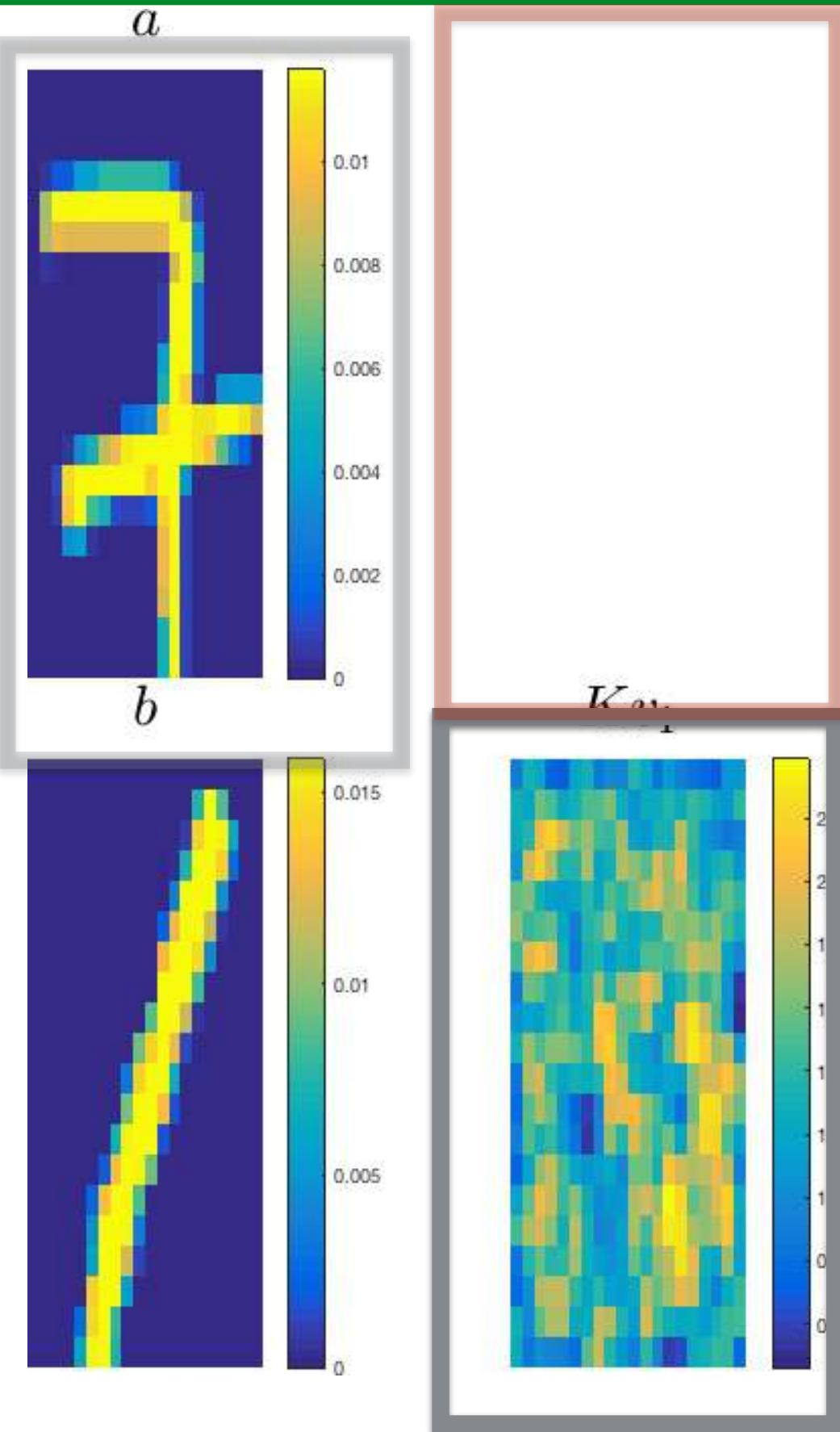
Very Fast EMD Approx. Solver



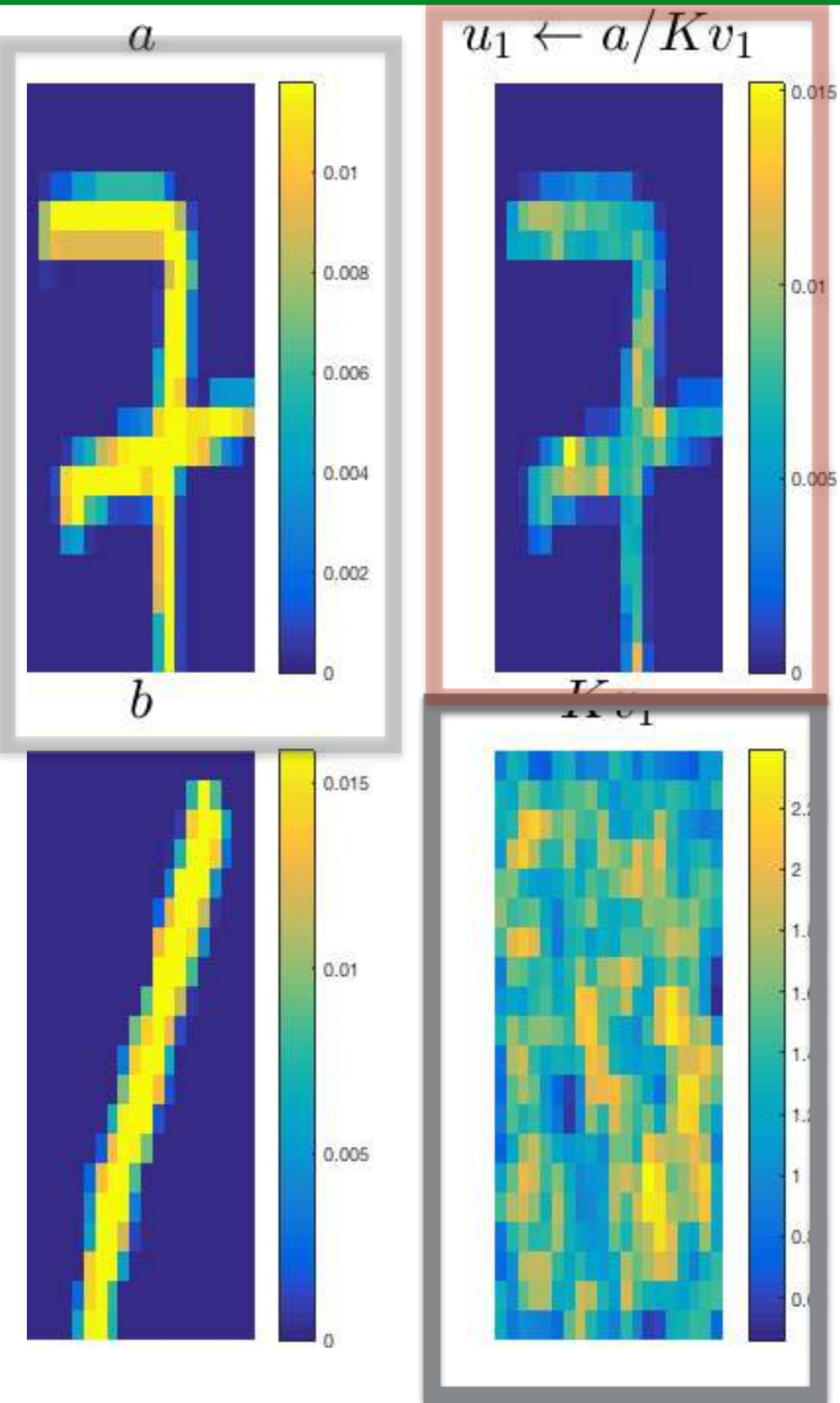
Very Fast EMD Approx. Solver



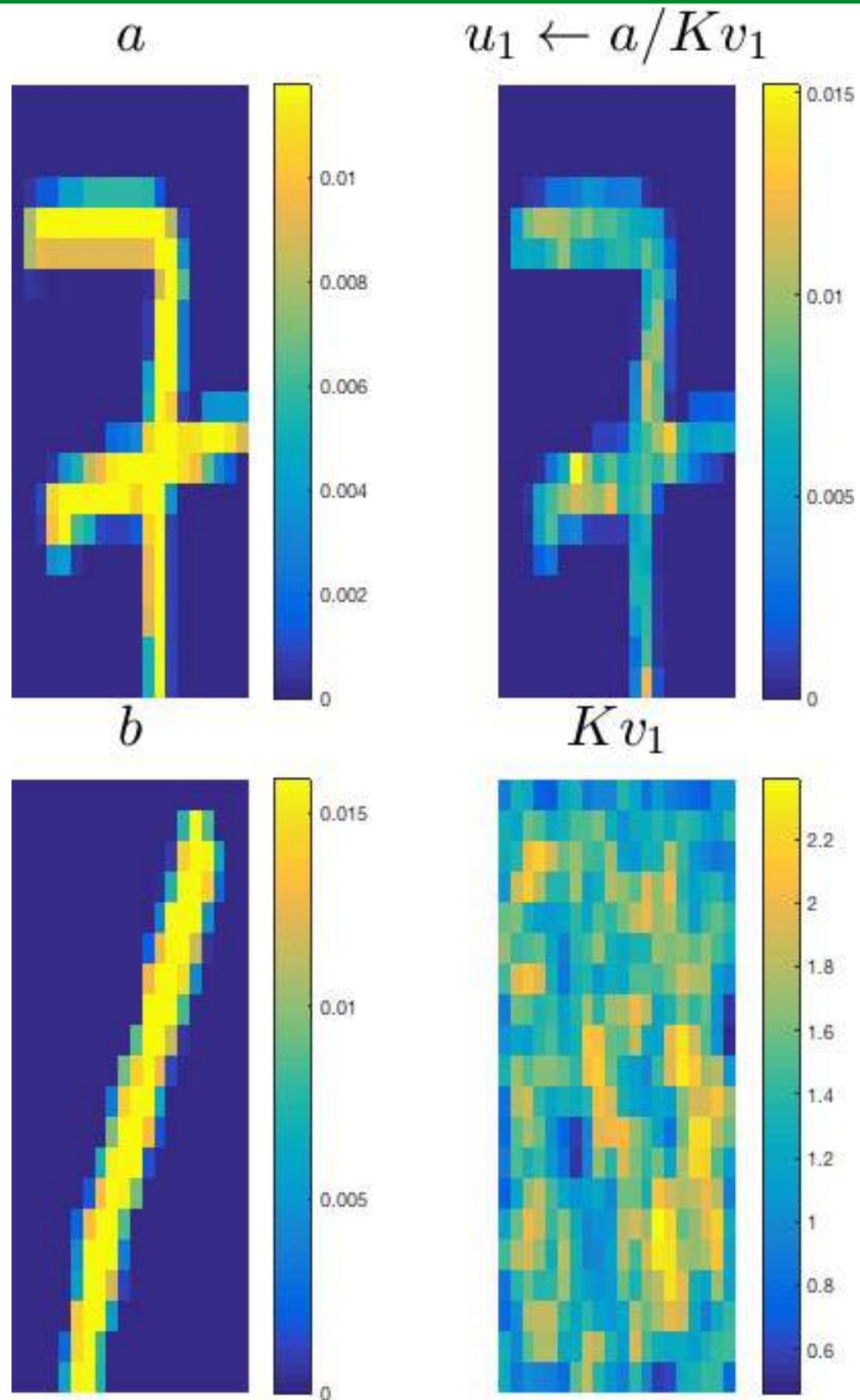
Very Fast EMD Approx. Solver



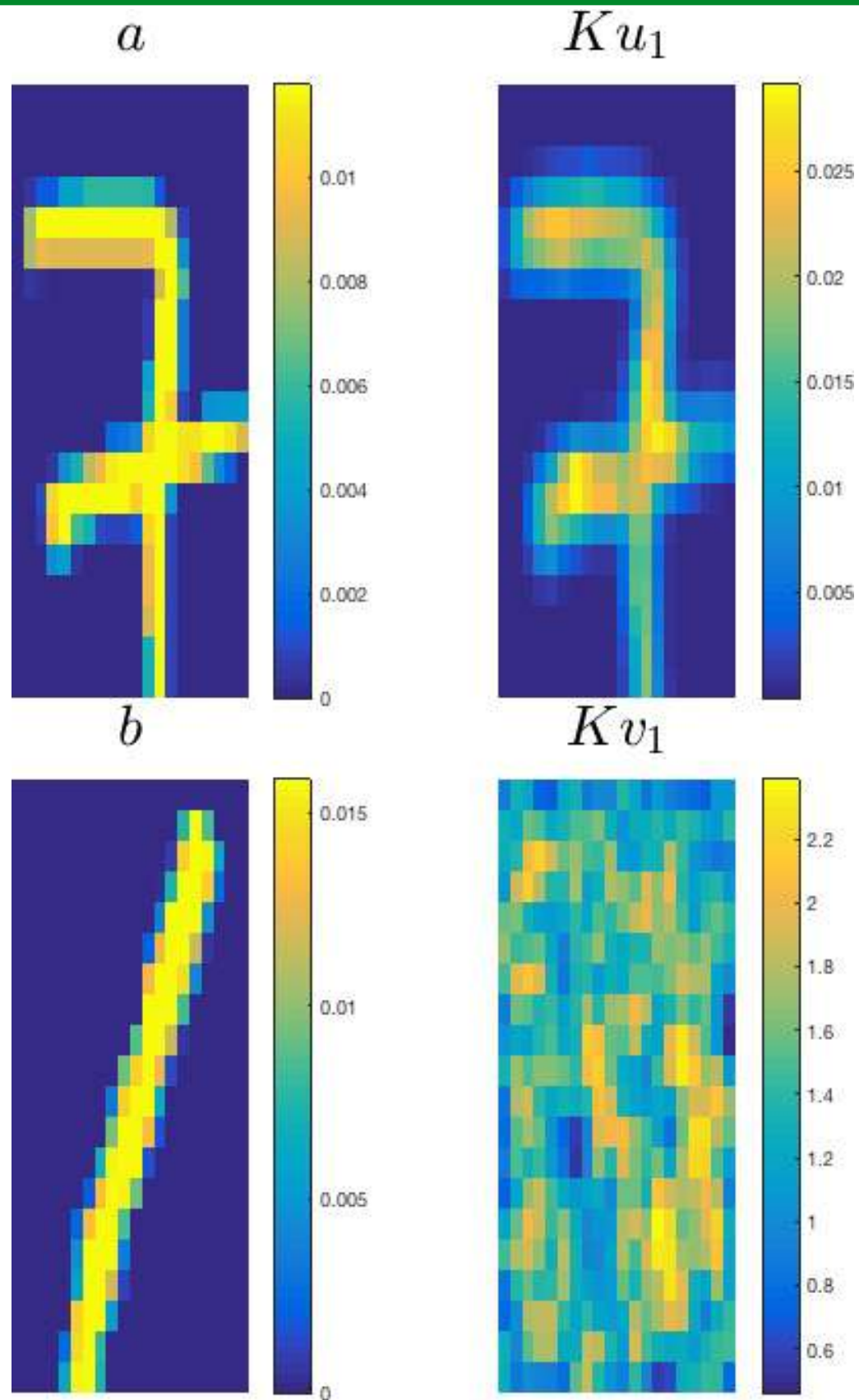
Very Fast EMD Approx. Solver



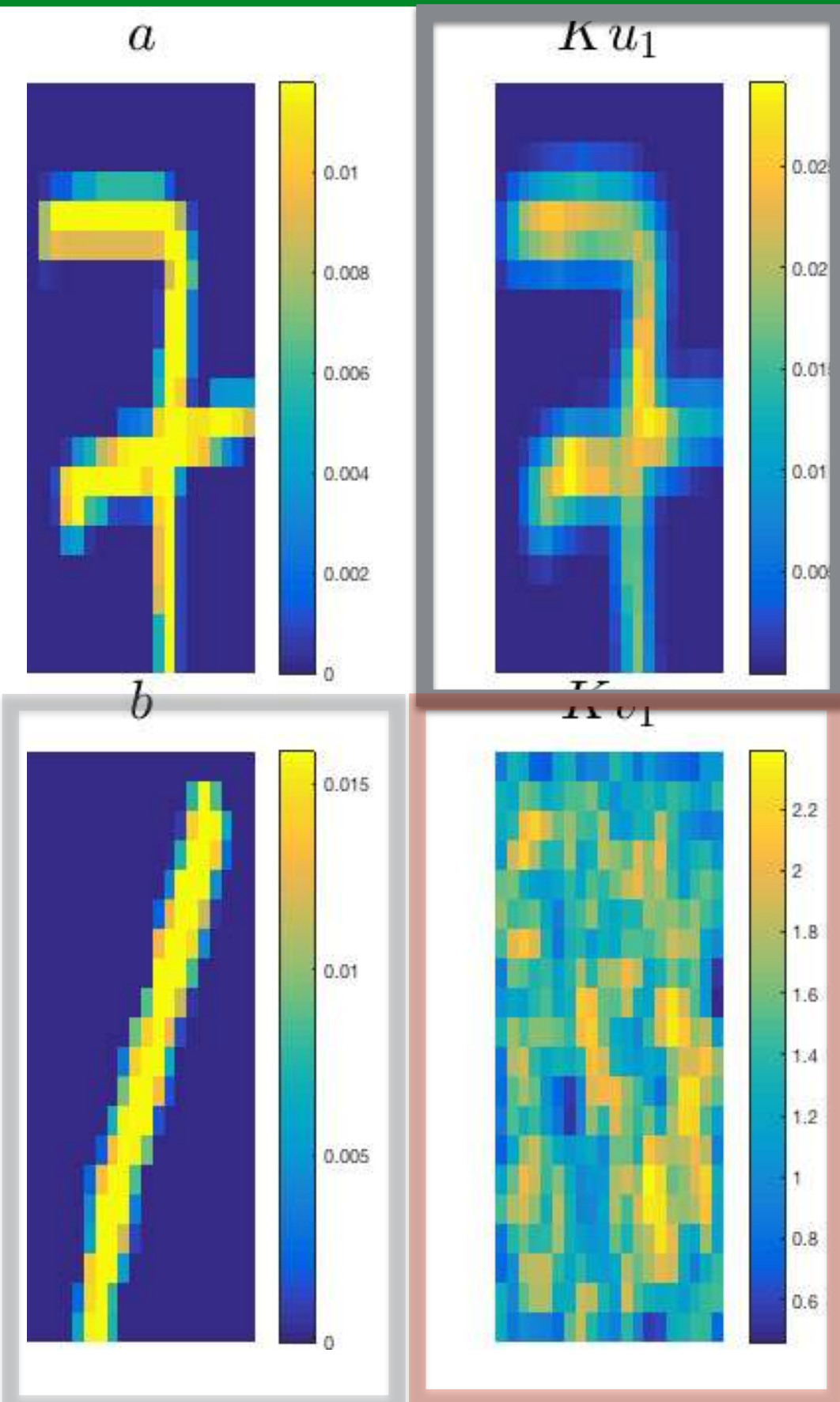
Very Fast EMD Approx. Solver



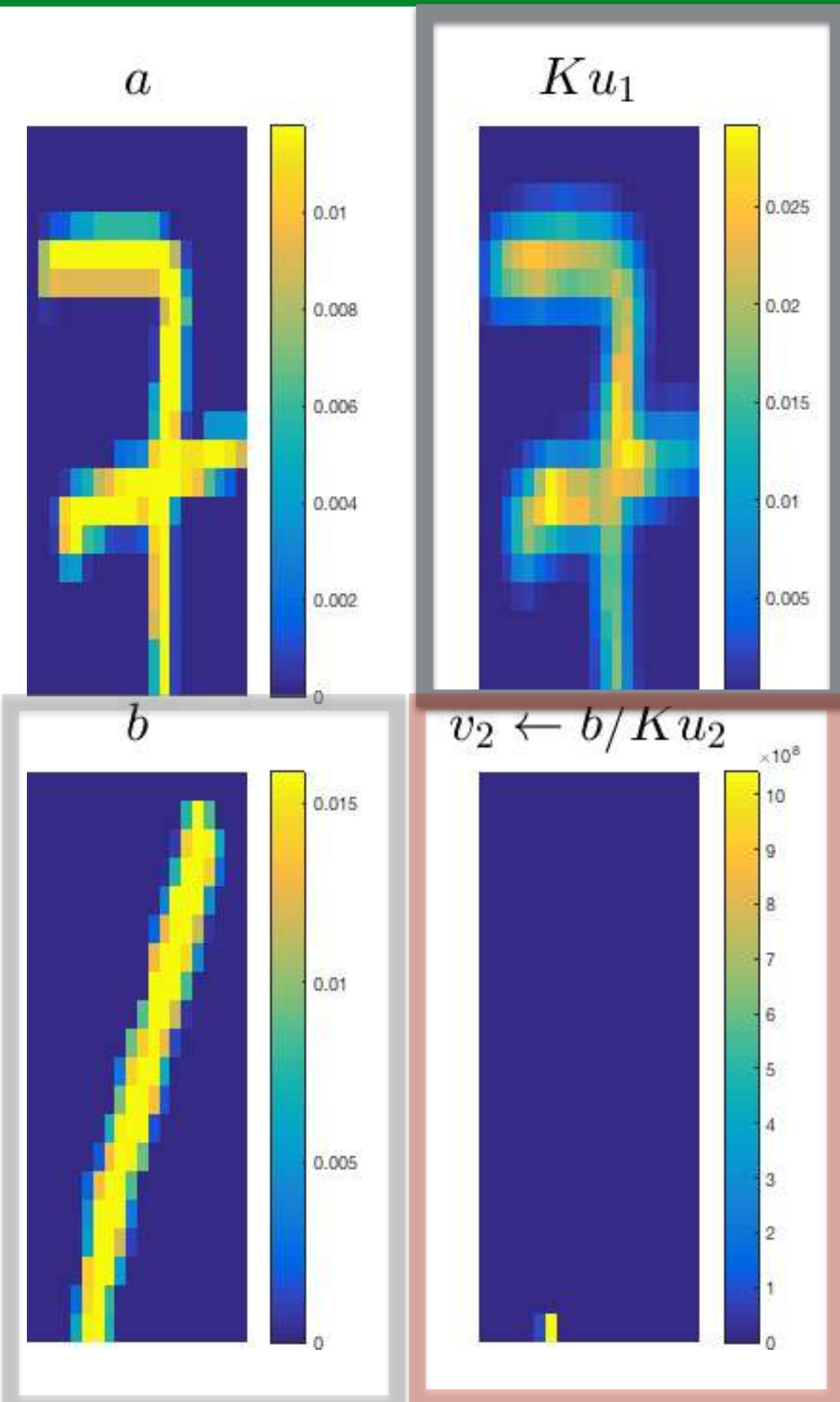
Very Fast EMD Approx. Solver



Very Fast EMD Approx. Solver

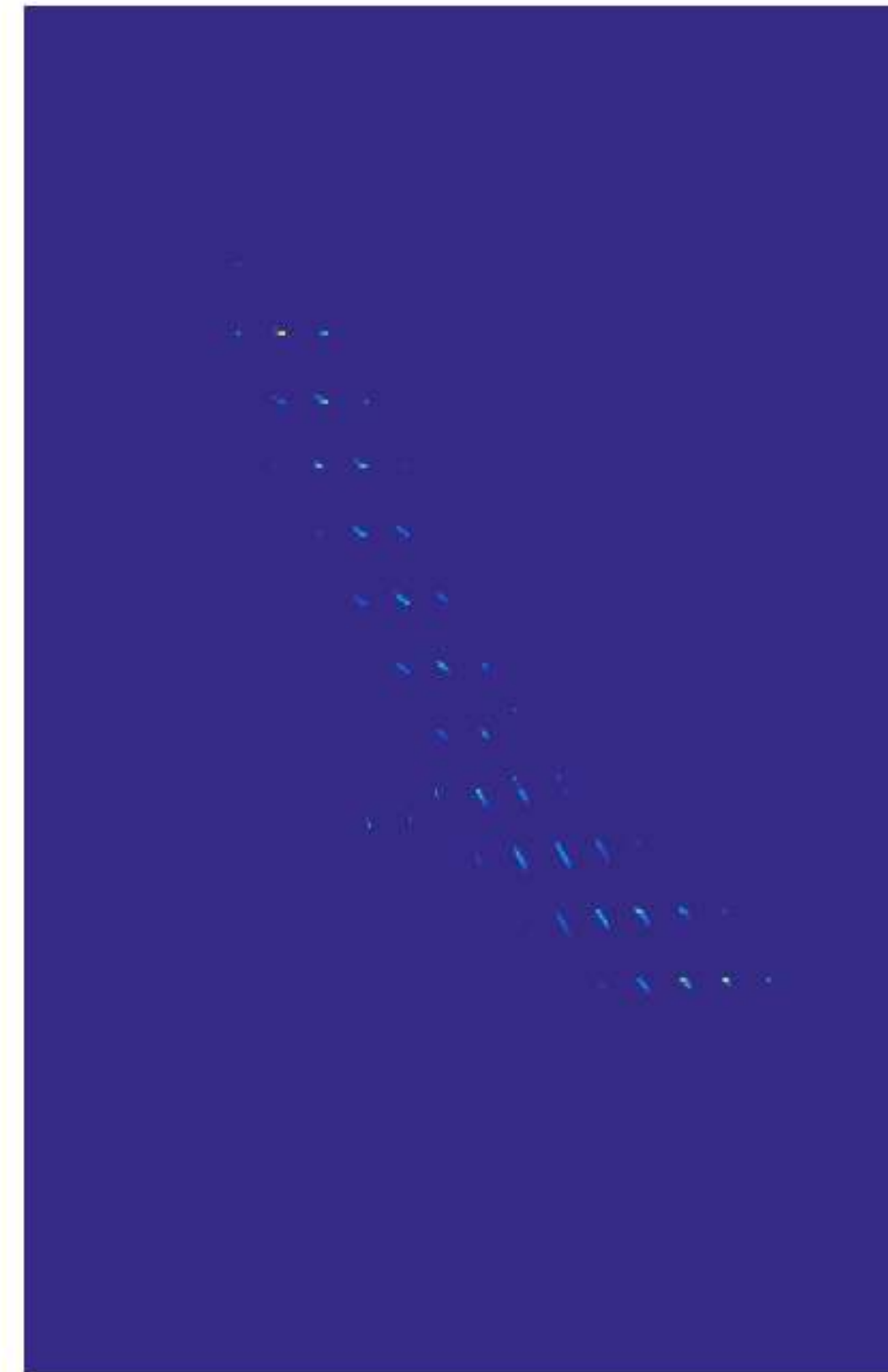


Very Fast EMD Approx. Solver



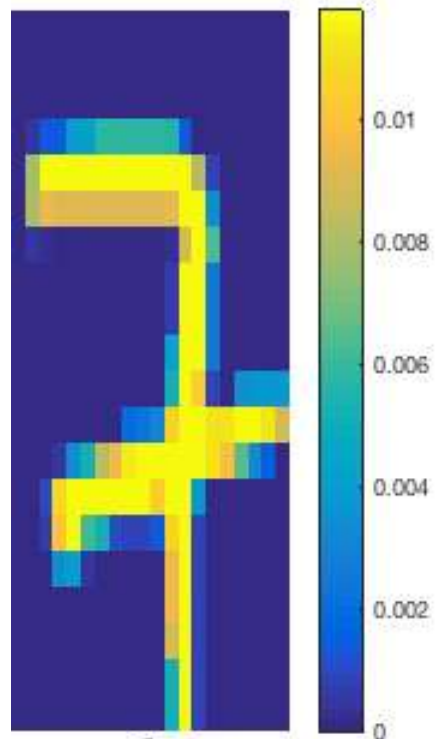
$$P_1 = D(u_1)KD(v_1)$$

$$\|P_1 - a\|_1 + \|P_1^T - b\|_1 = 1.2691$$

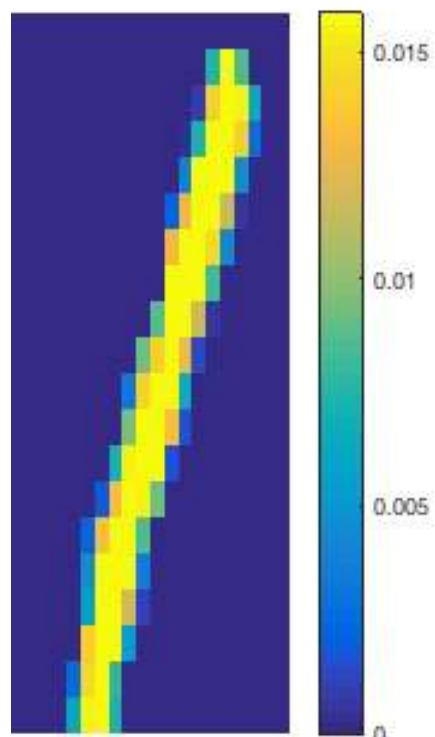


Very Fast EMD Approx. Solver

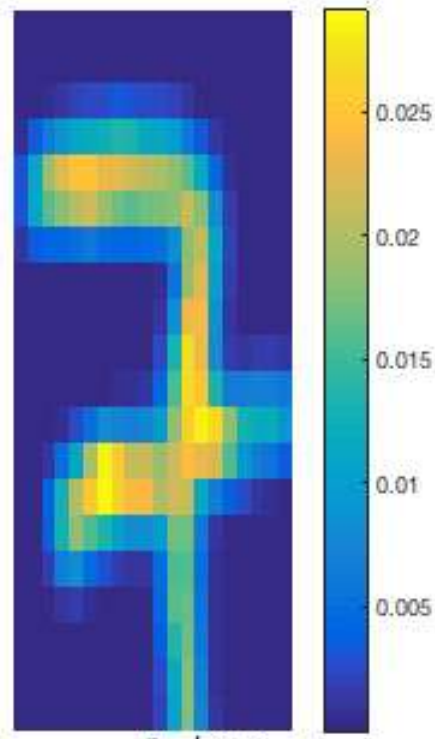
a



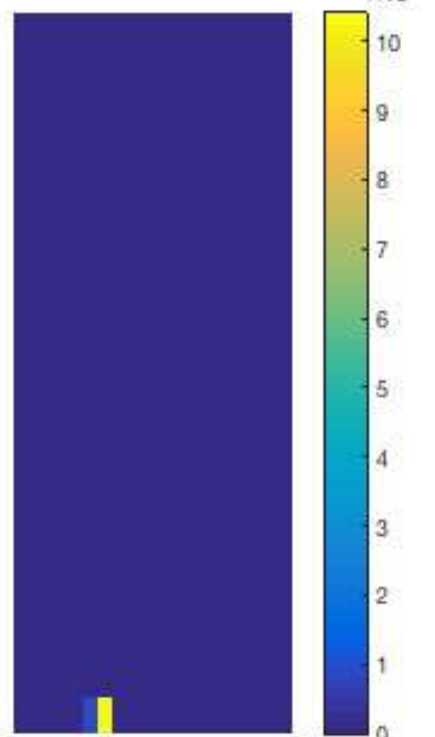
b



Ku_1

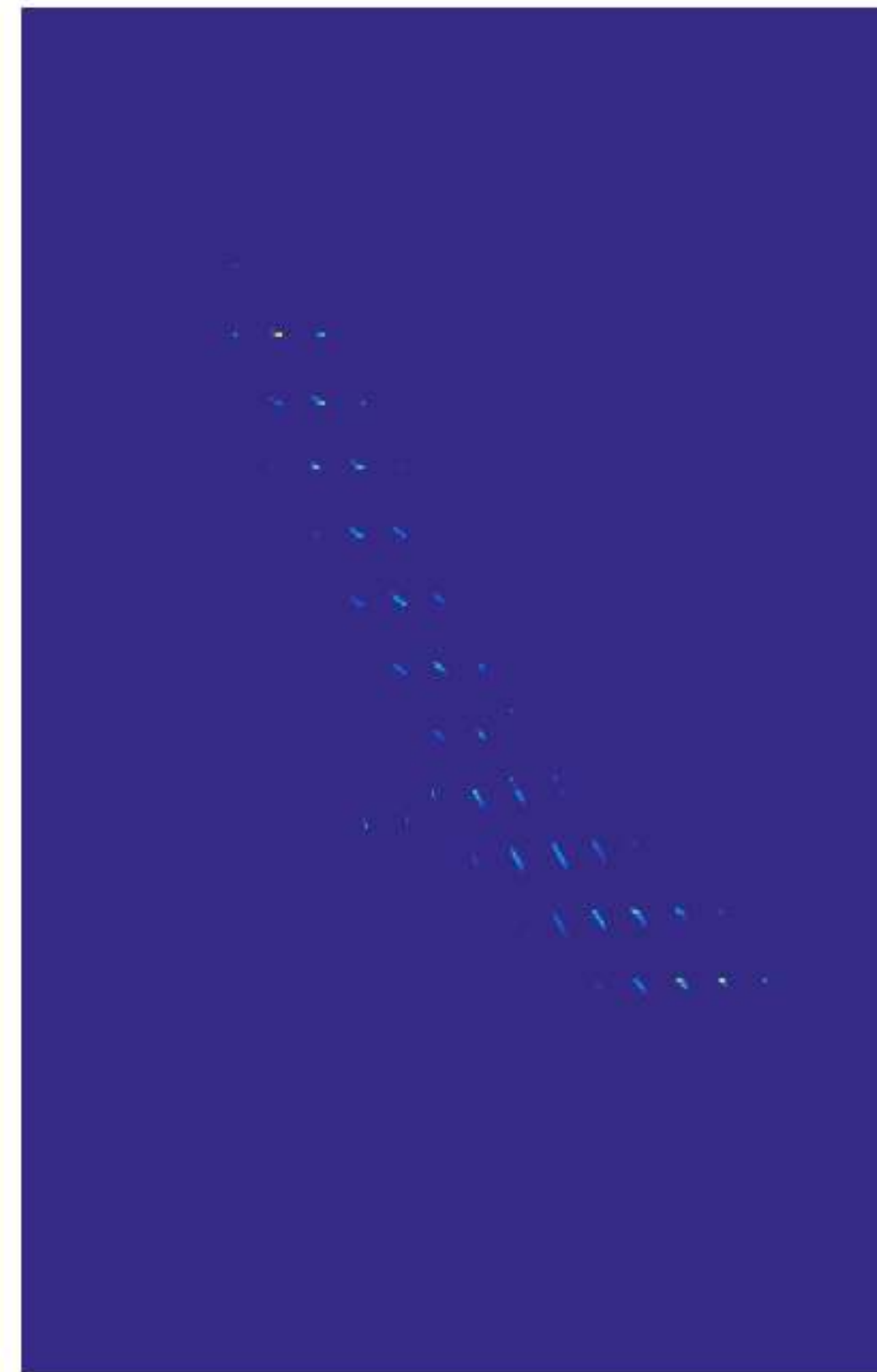


$v_2 \leftarrow b/Ku_2$

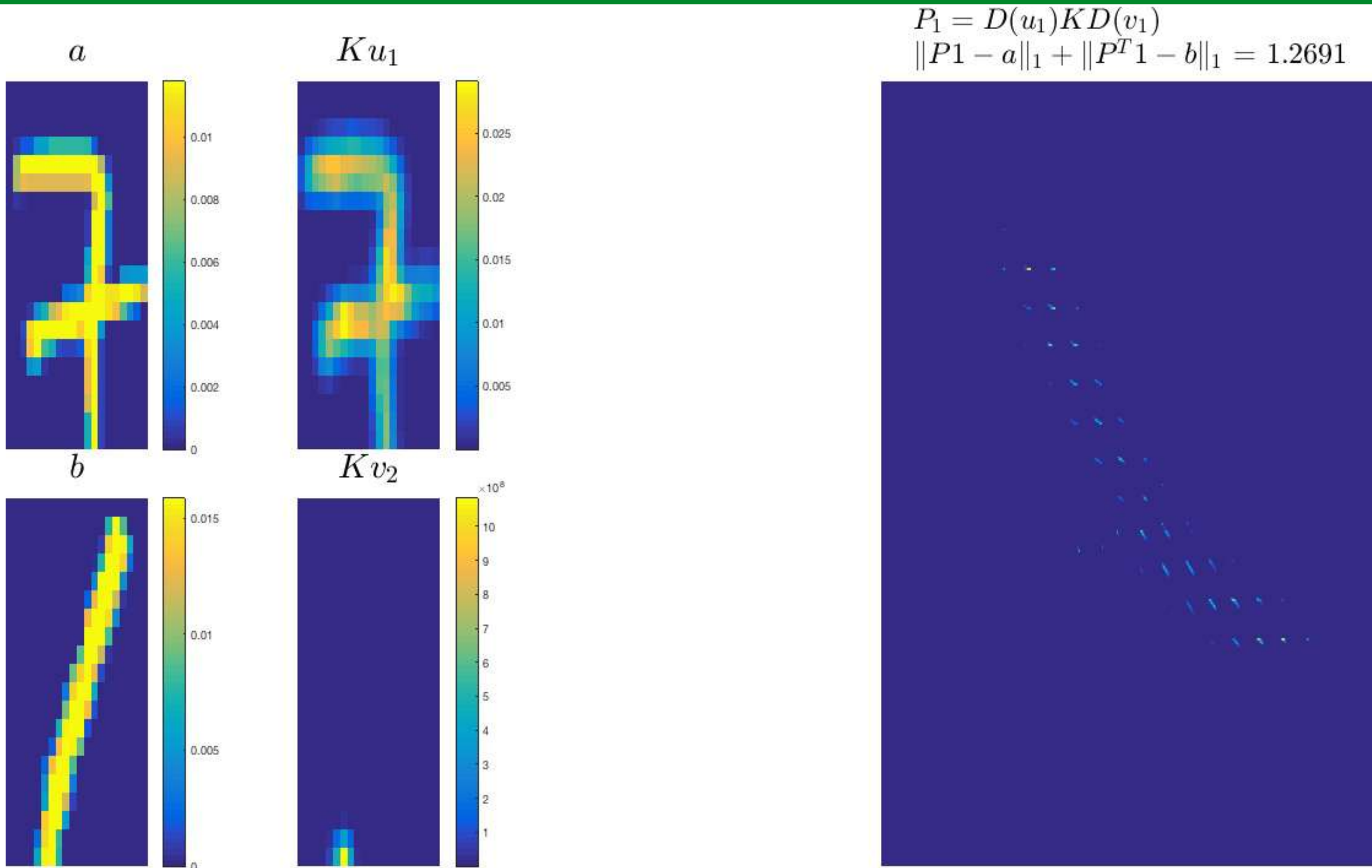


$$P_1 = D(u_1)KD(v_1)$$

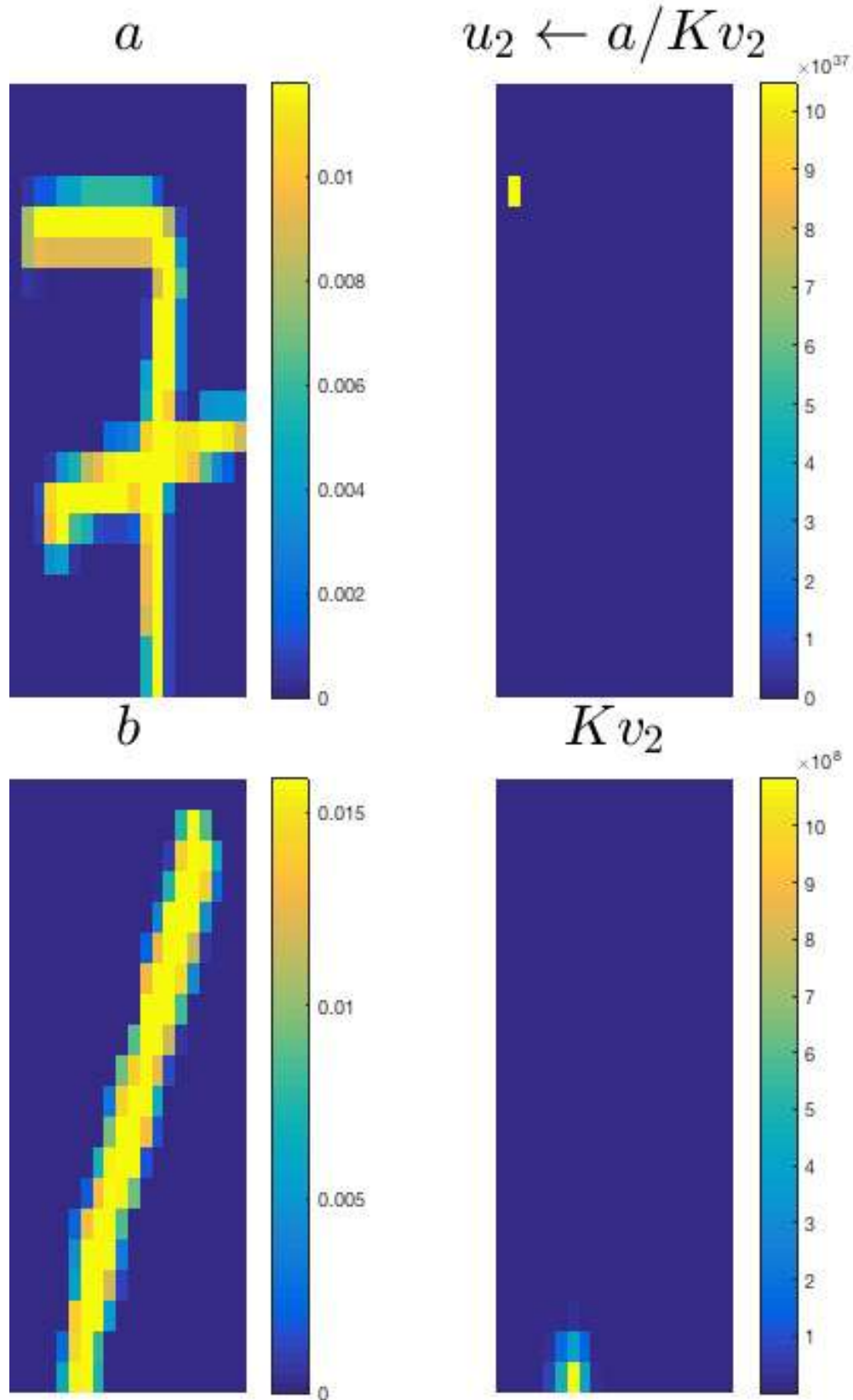
$$\|P_1 - a\|_1 + \|P_1^T - b\|_1 = 1.2691$$



Very Fast EMD Approx. Solver

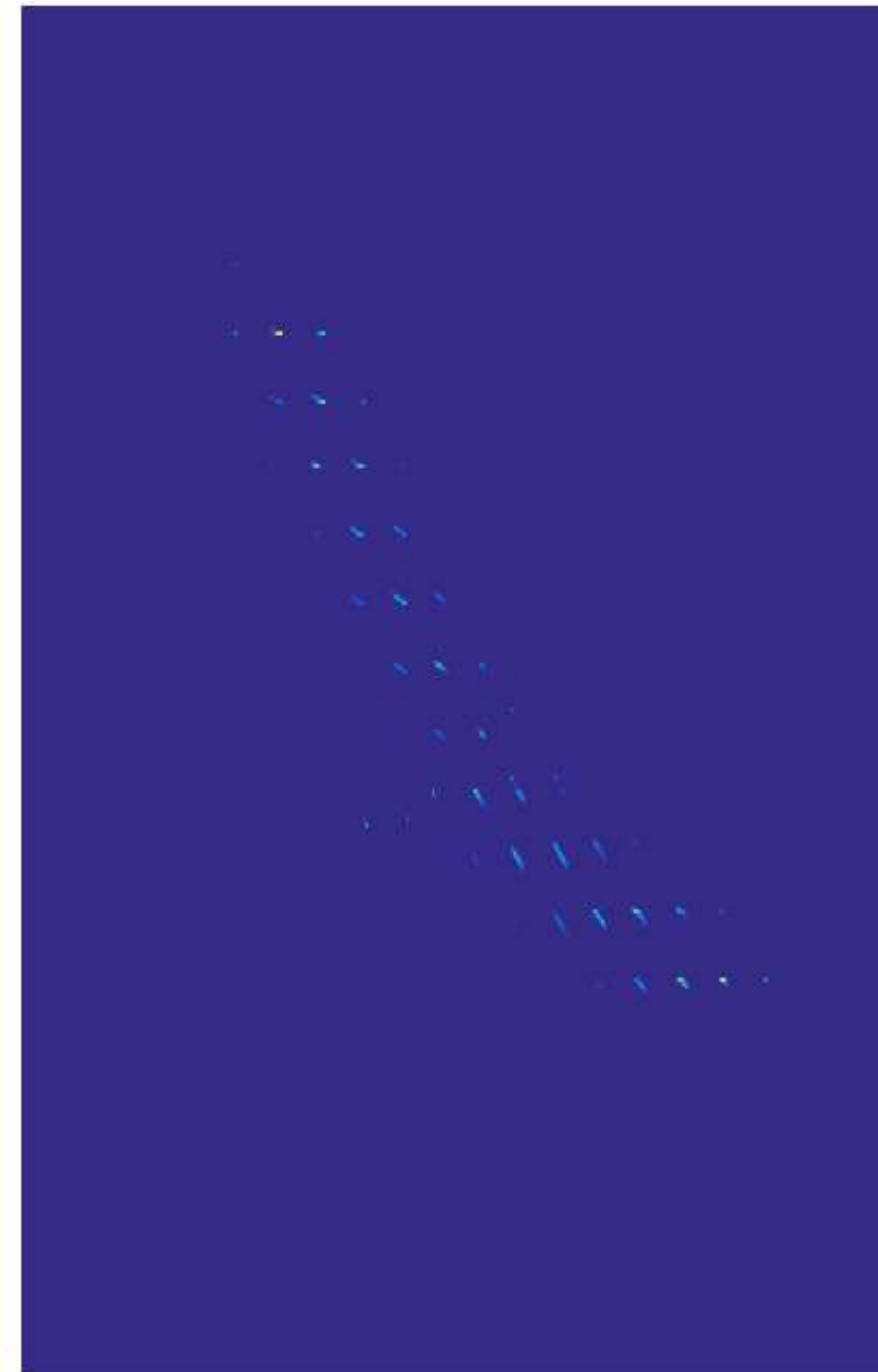


Very Fast EMD Approx. Solver

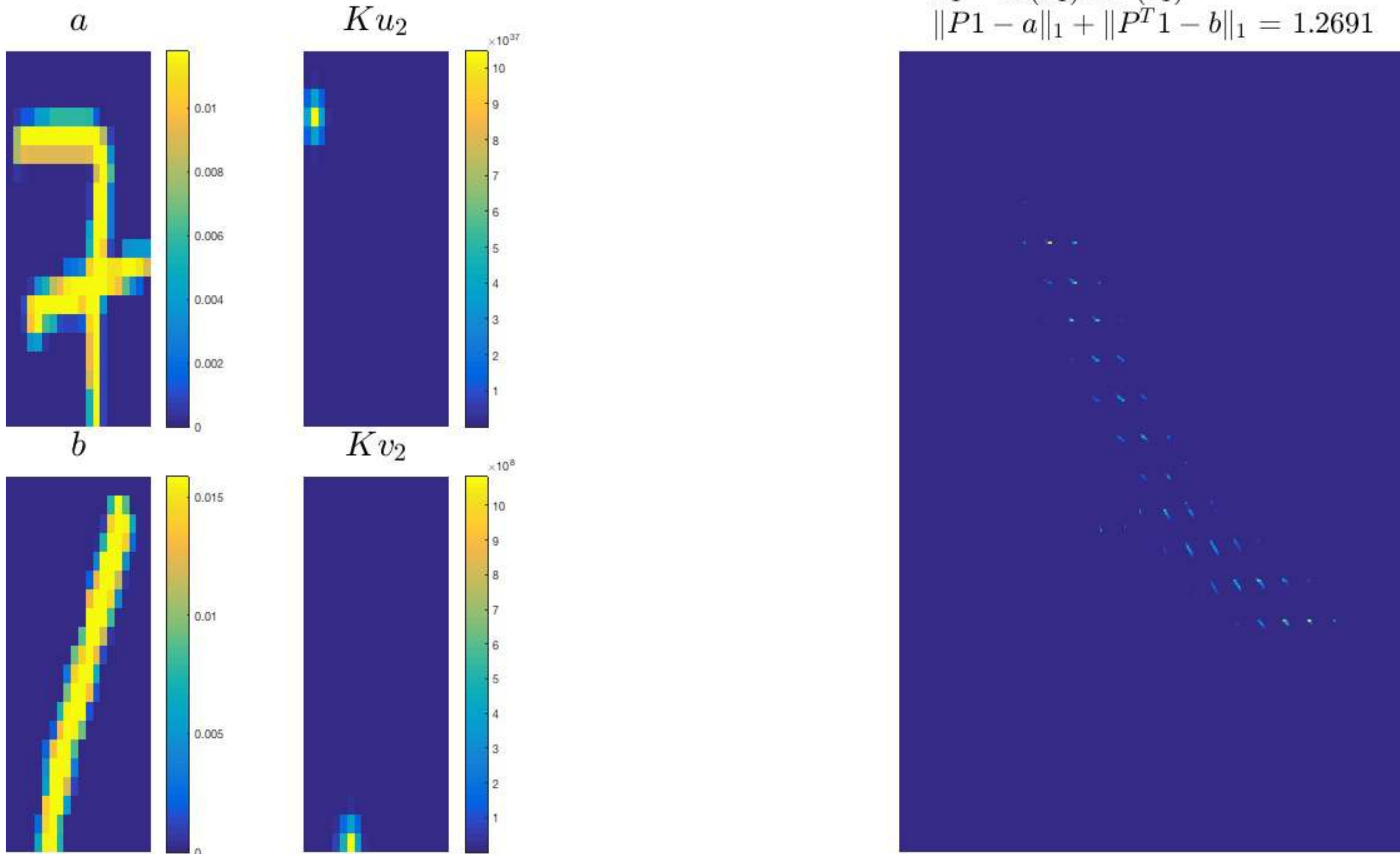


$$P_1 = D(u_1) K D(v_1)$$

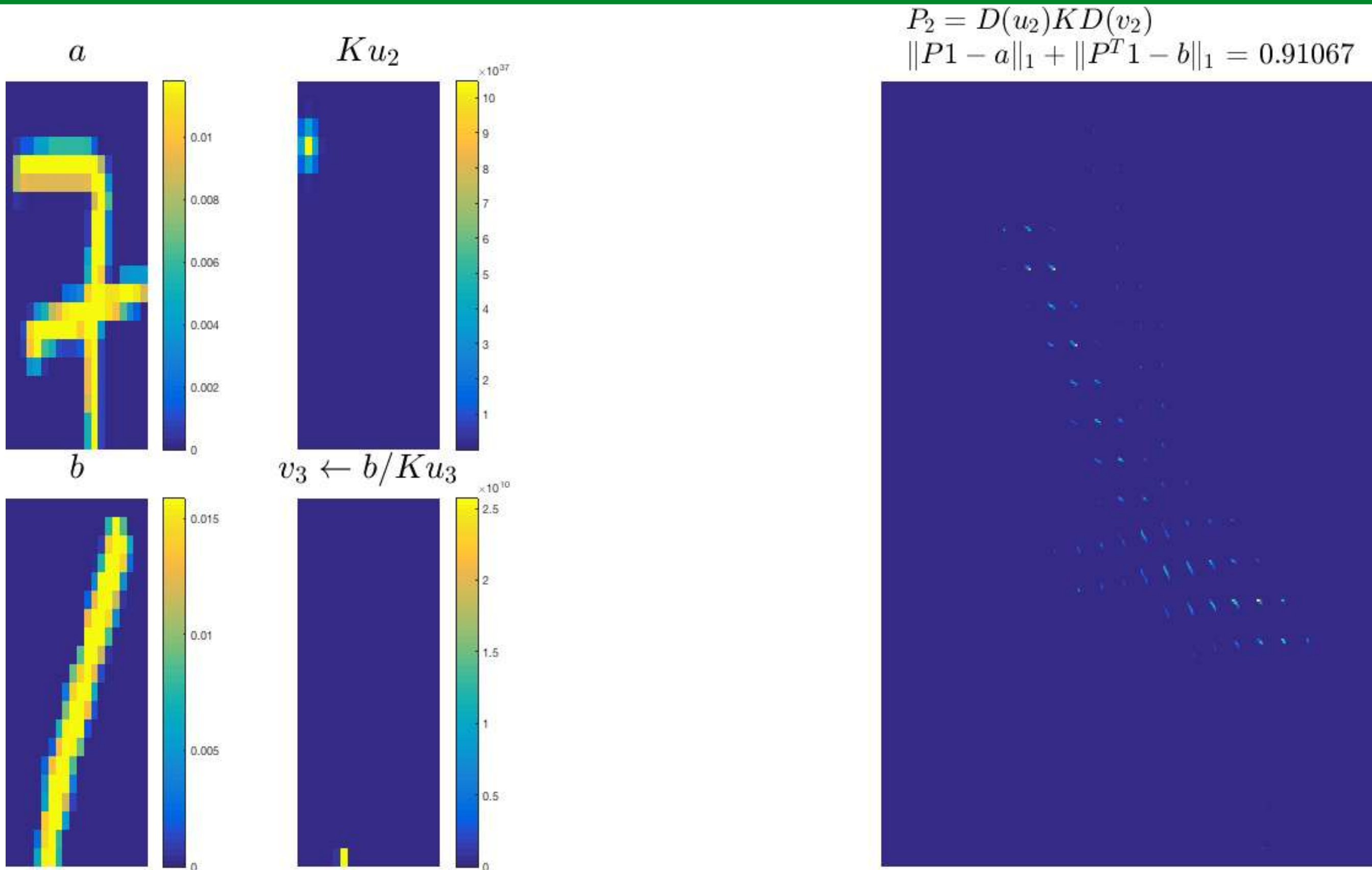
$$\|P_1 1 - a\|_1 + \|P_1^T 1 - b\|_1 = 1.2691$$



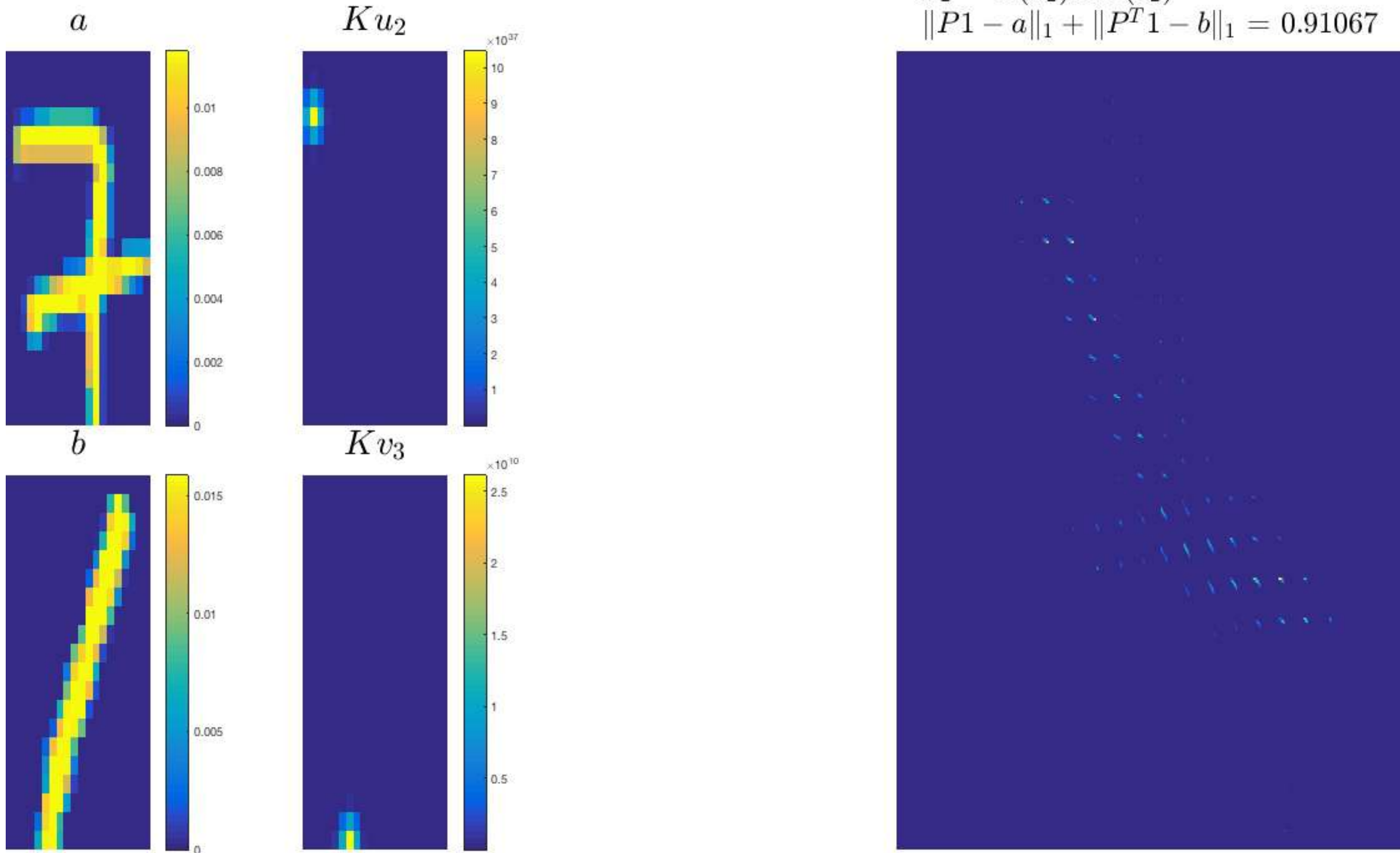
Very Fast EMD Approx. Solver



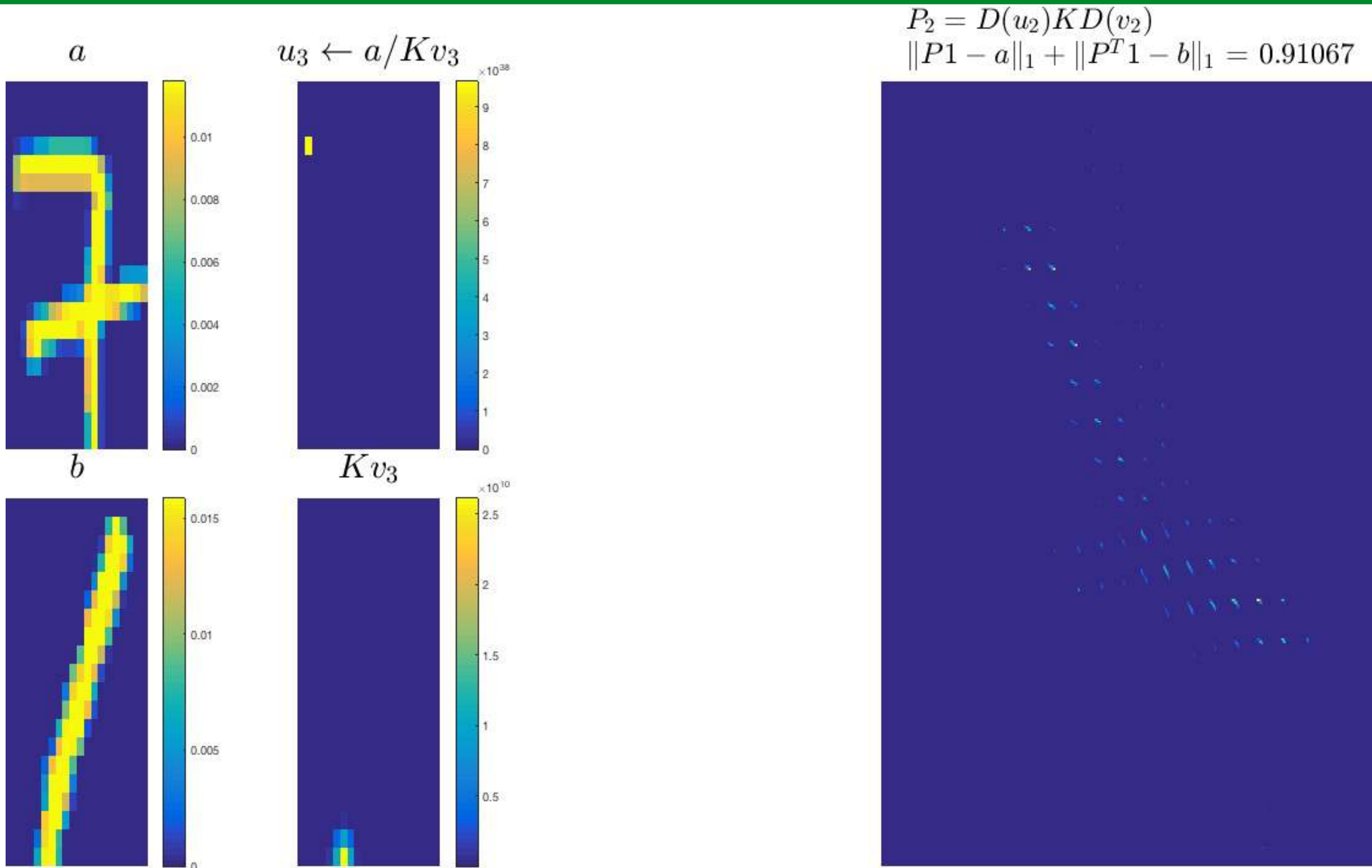
Very Fast EMD Approx. Solver



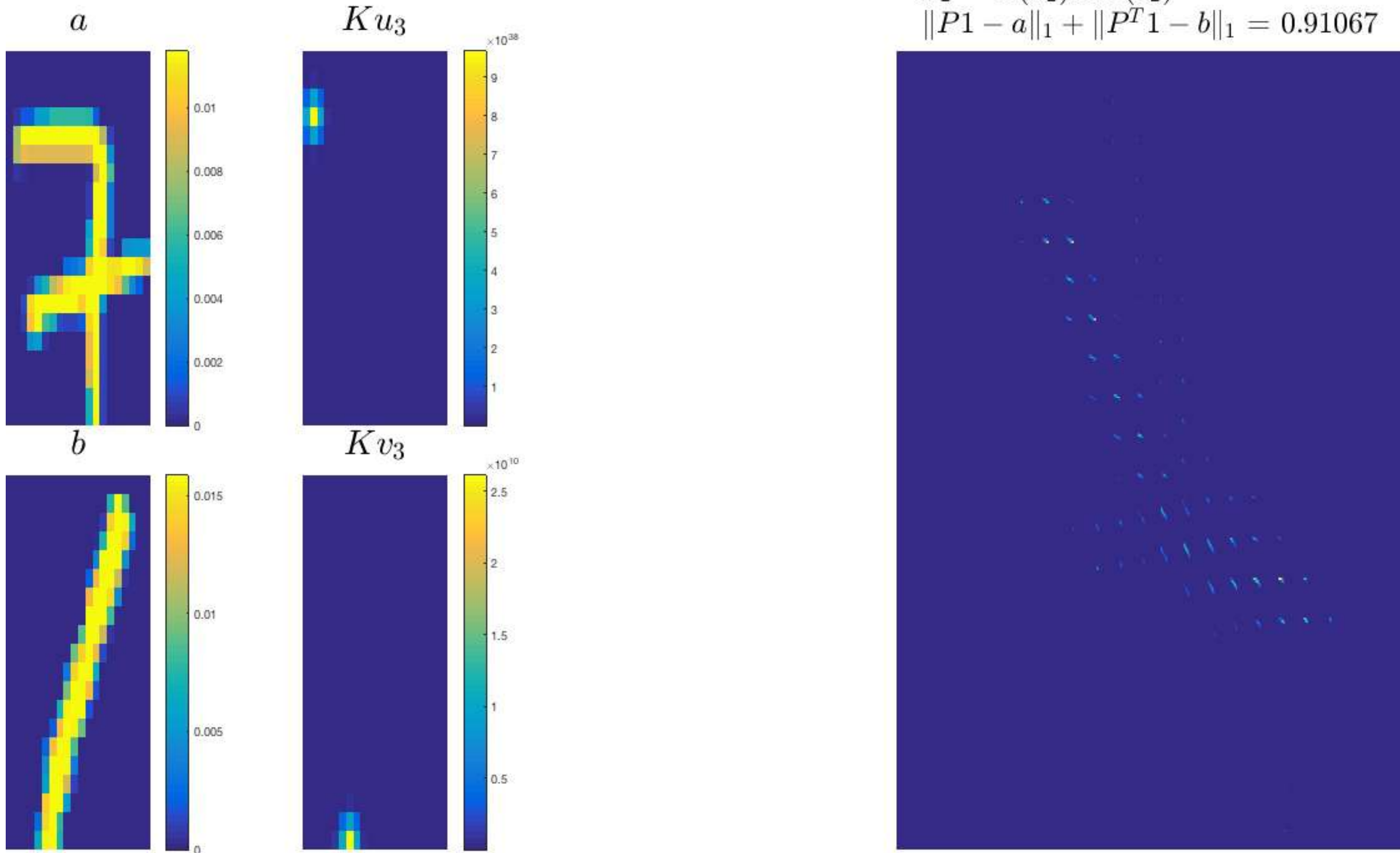
Very Fast EMD Approx. Solver



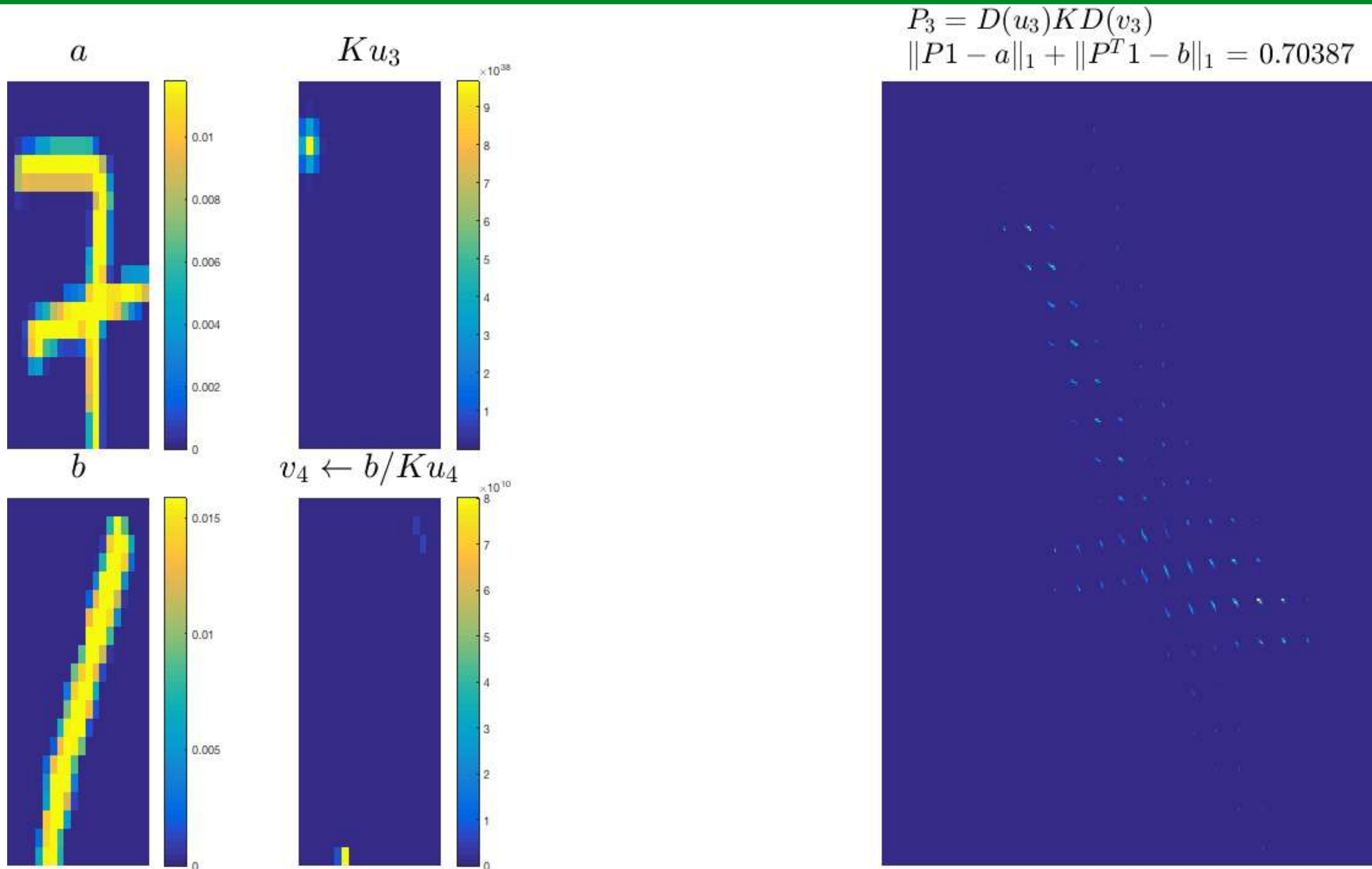
Very Fast EMD Approx. Solver



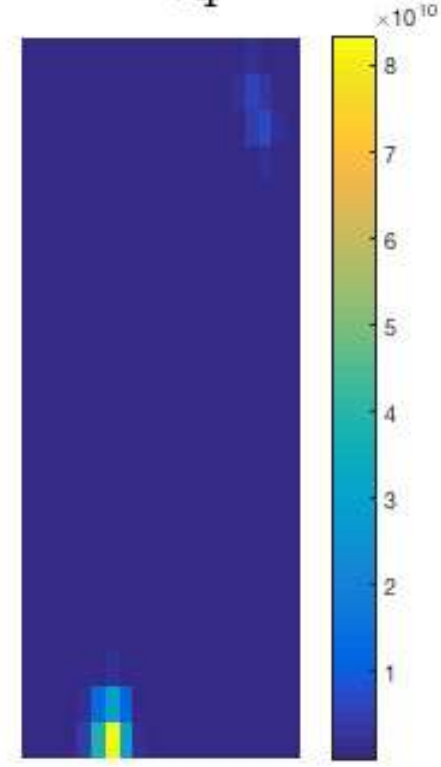
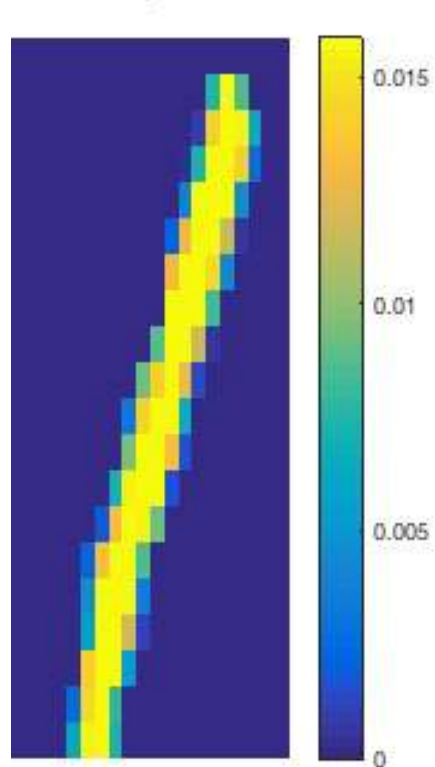
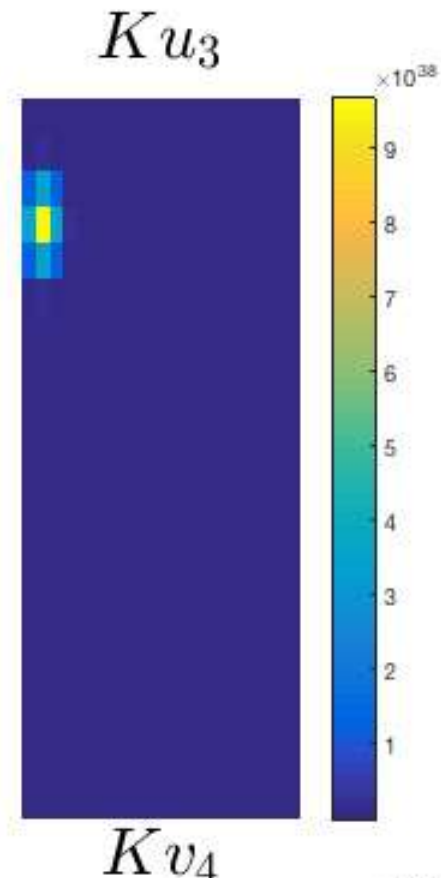
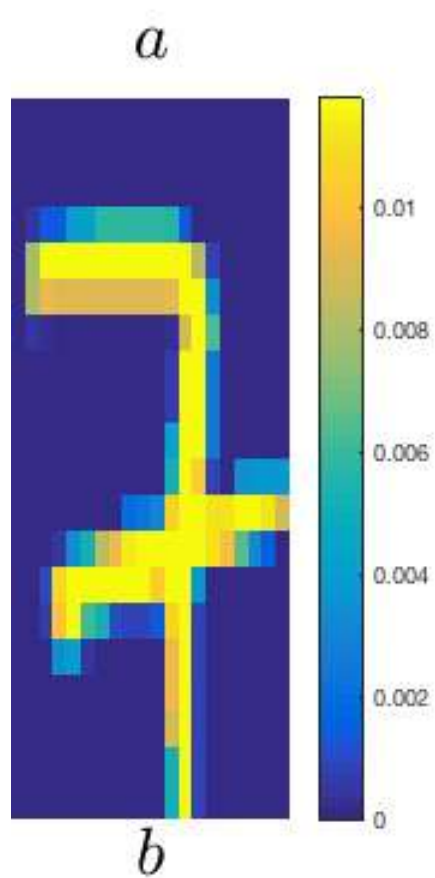
Very Fast EMD Approx. Solver



Very Fast EMD Approx. Solver

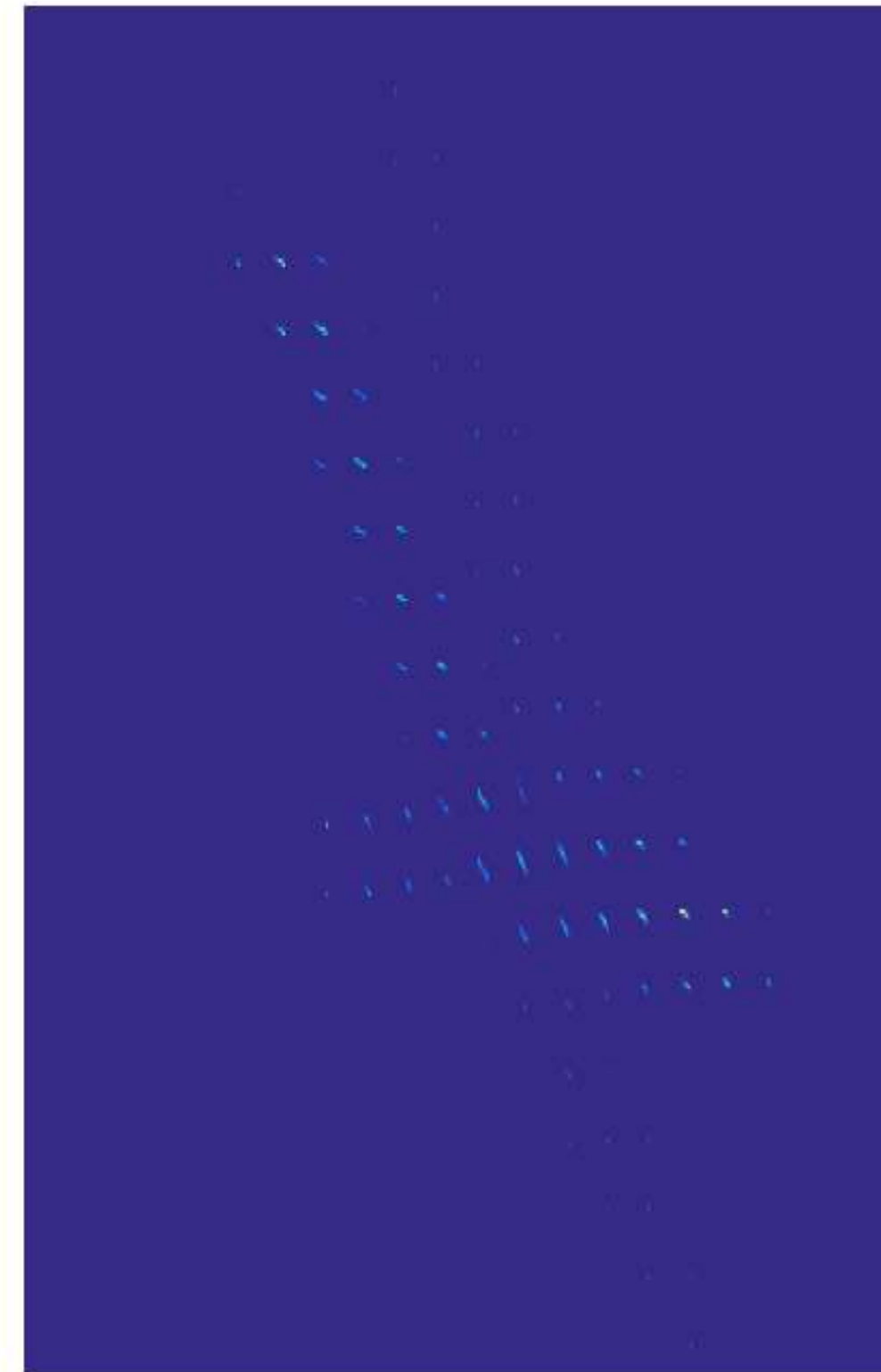


Very Fast EMD Approx. Solver

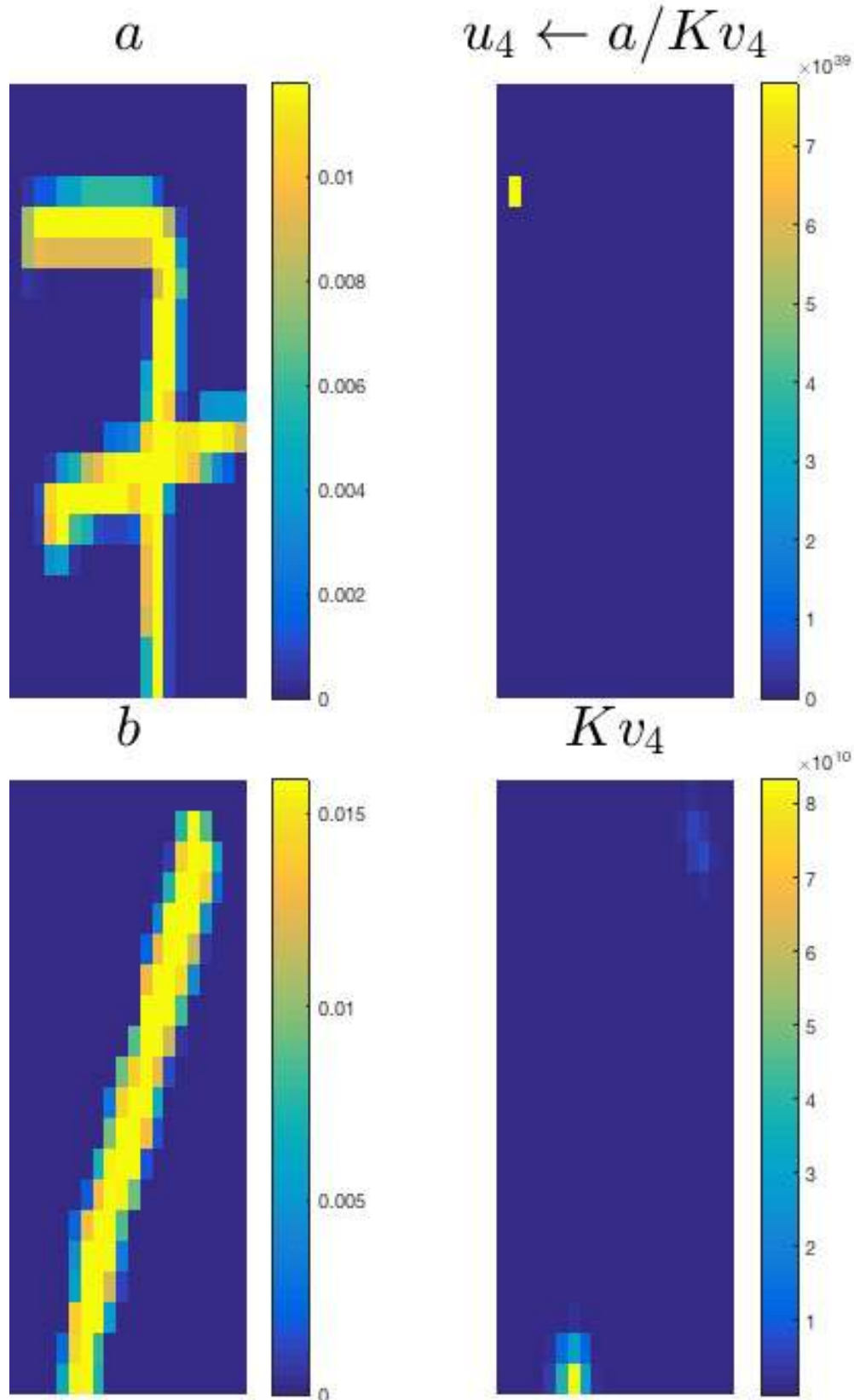


$$P_3 = D(u_3)KD(v_3)$$

$$\|P1 - a\|_1 + \|P^T 1 - b\|_1 = 0.70387$$

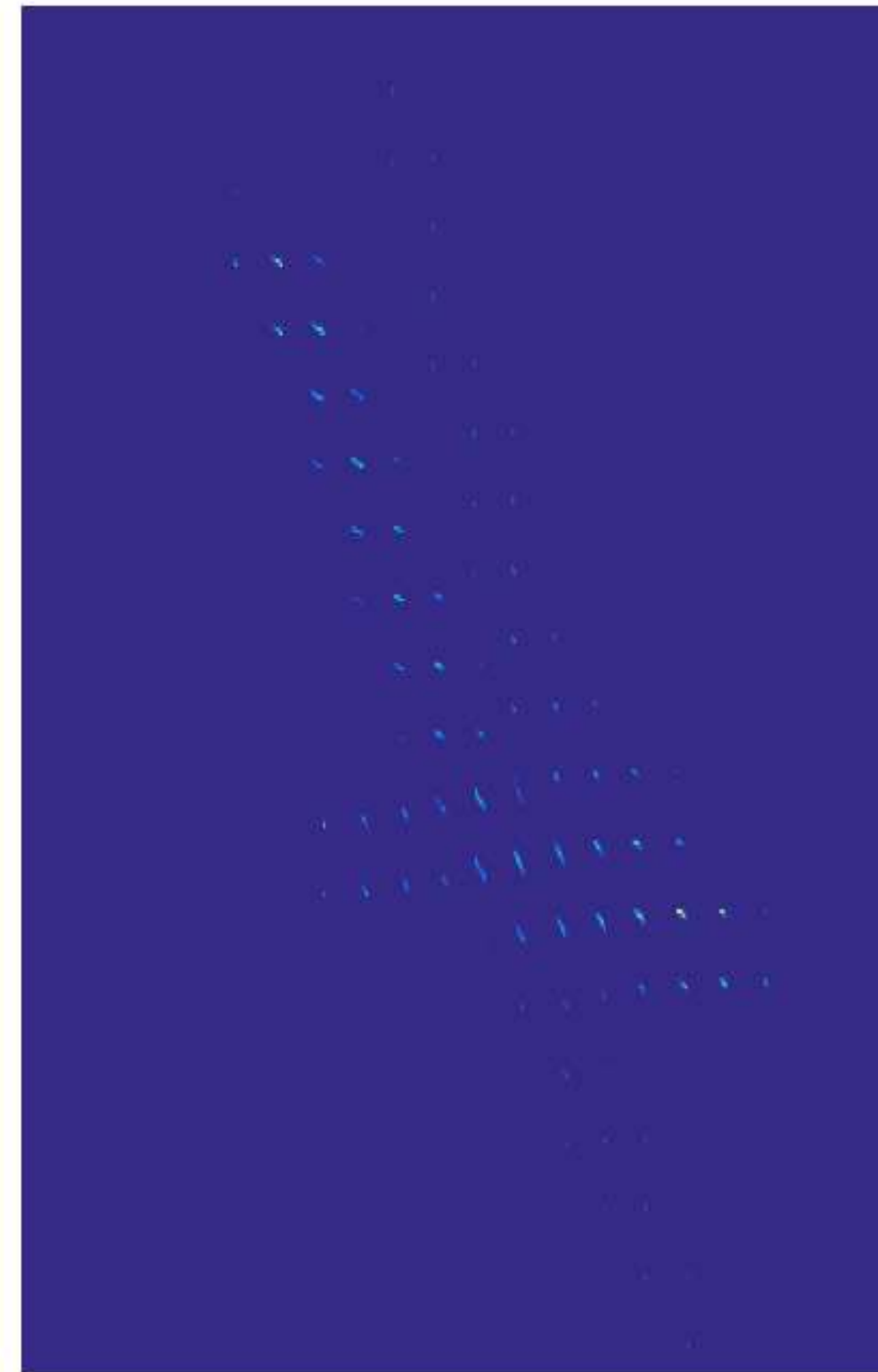


Very Fast EMD Approx. Solver

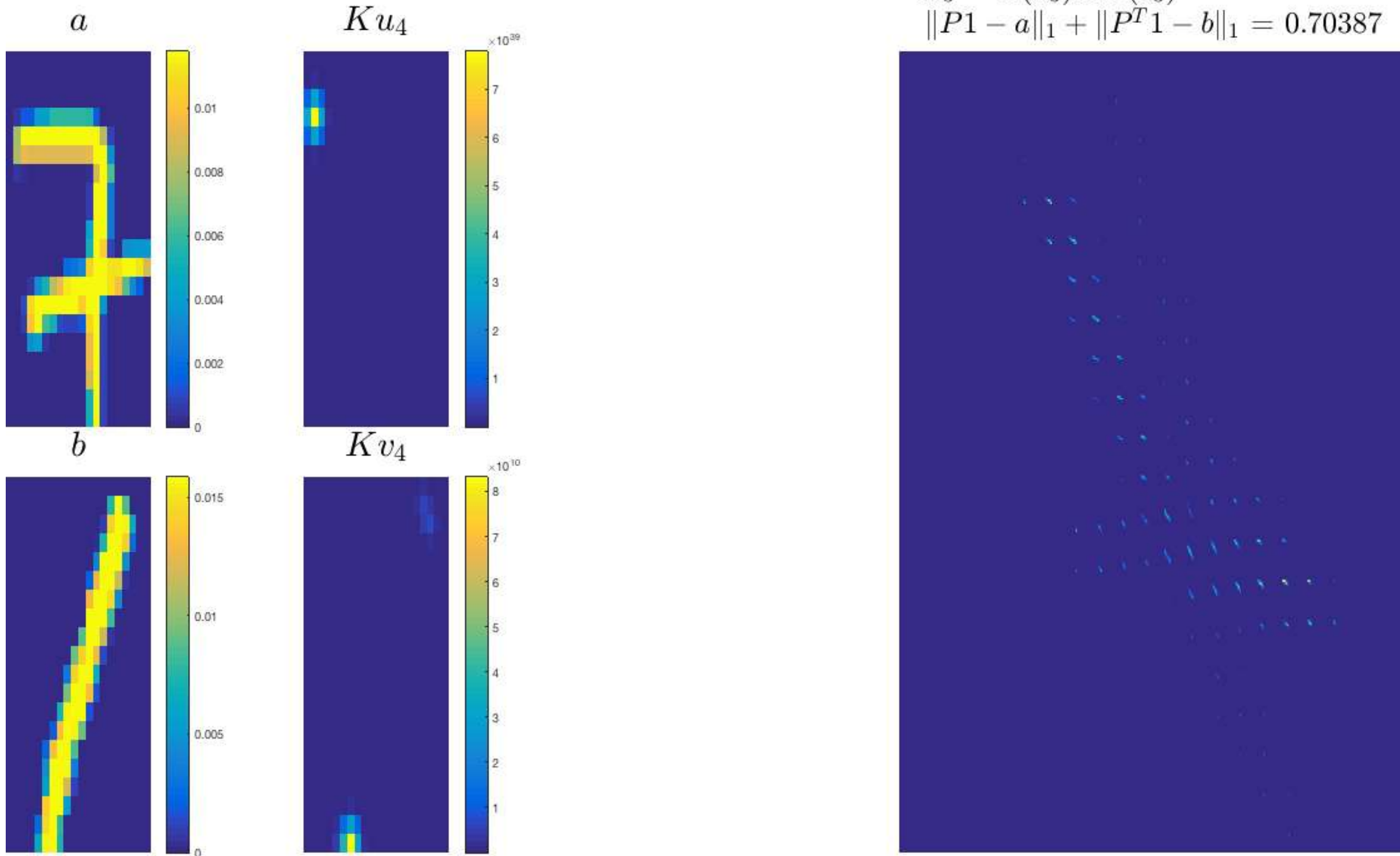


$$P_3 = D(u_3) K D(v_3)$$

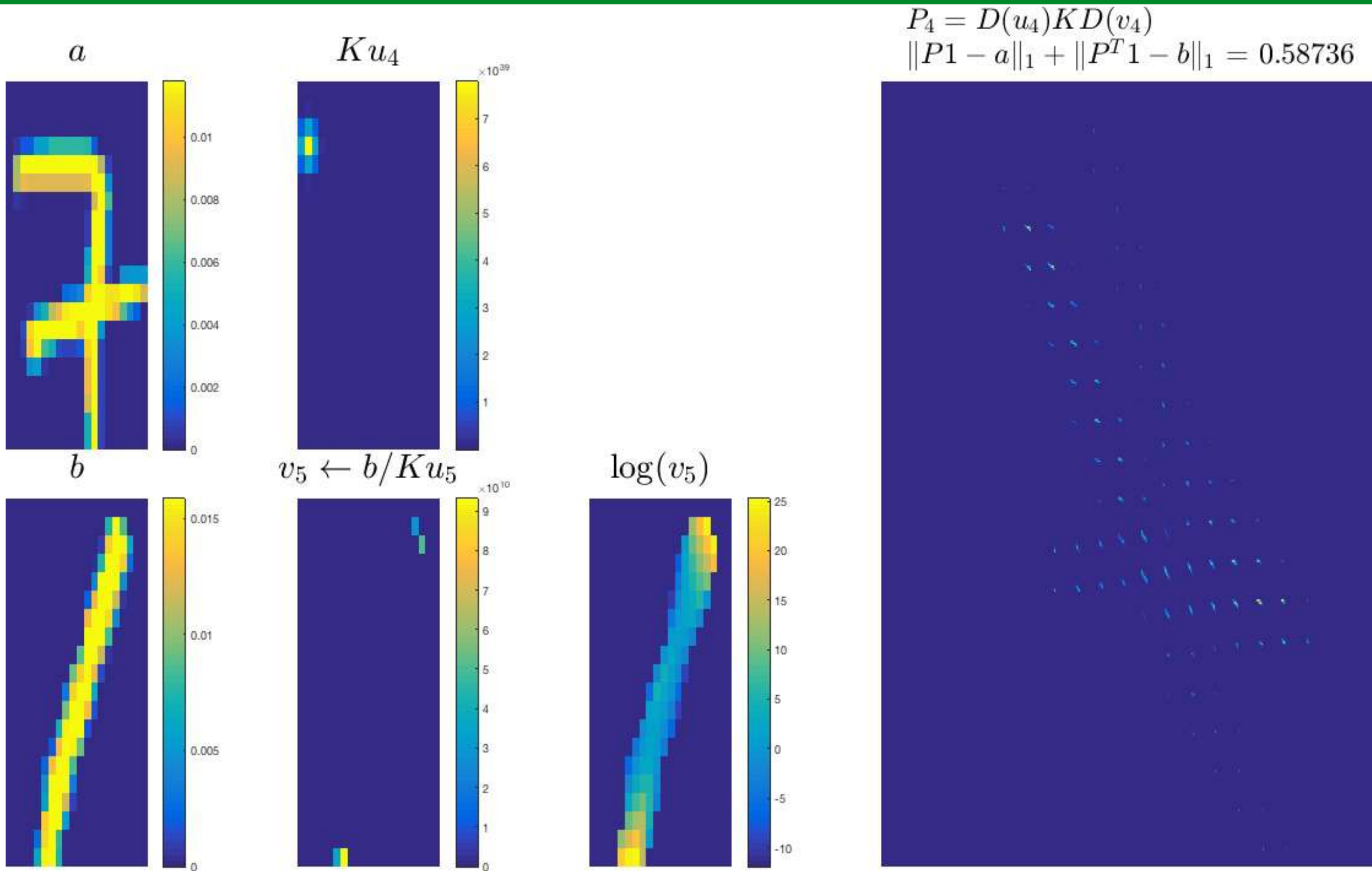
$$\|P1 - a\|_1 + \|P^T 1 - b\|_1 = 0.70387$$



Very Fast EMD Approx. Solver

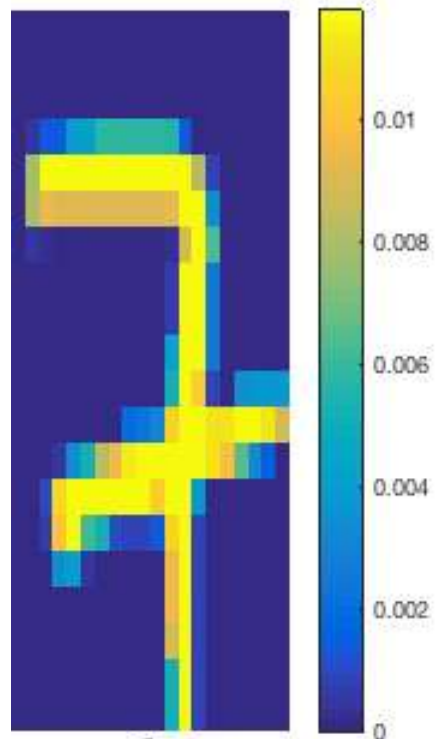


Very Fast EMD Approx. Solver

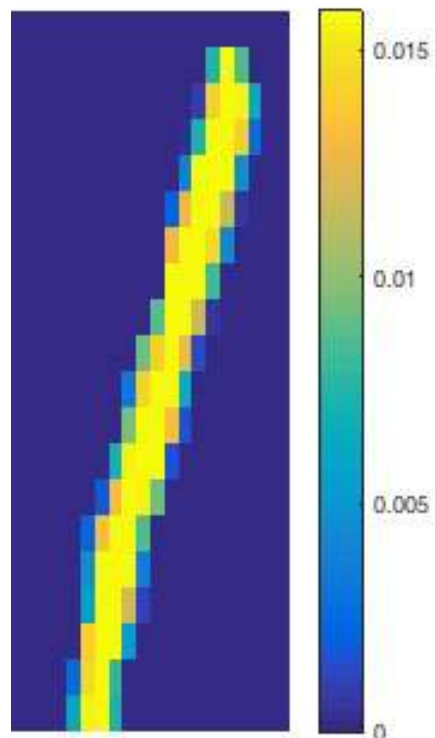


Very Fast EMD Approx. Solver

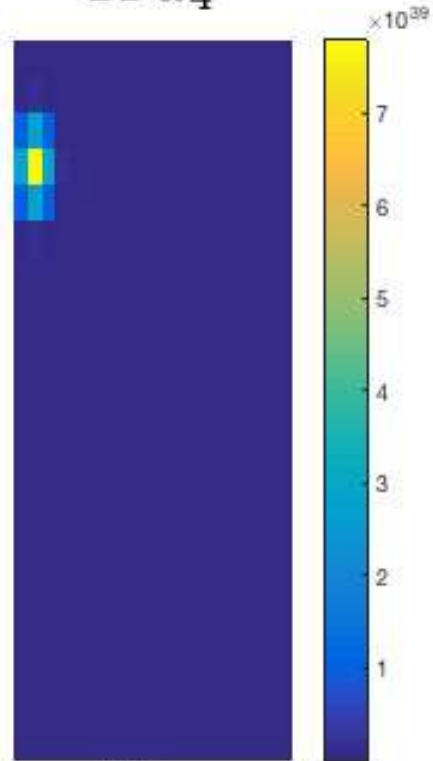
a



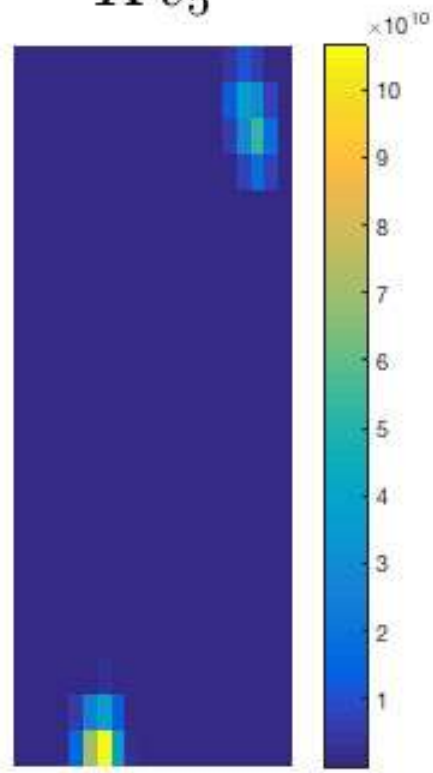
b



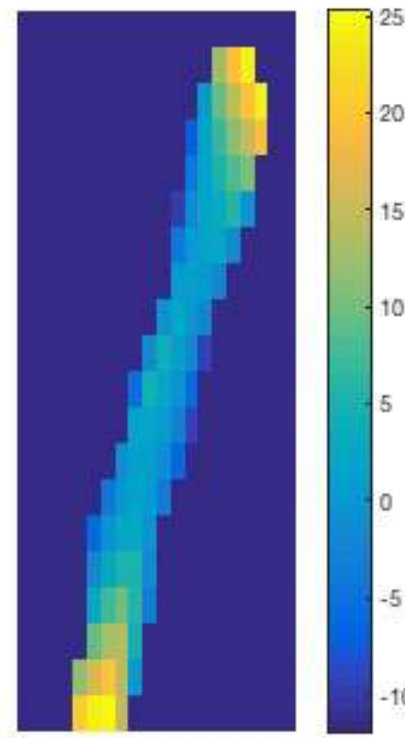
Ku_4



Kv_5

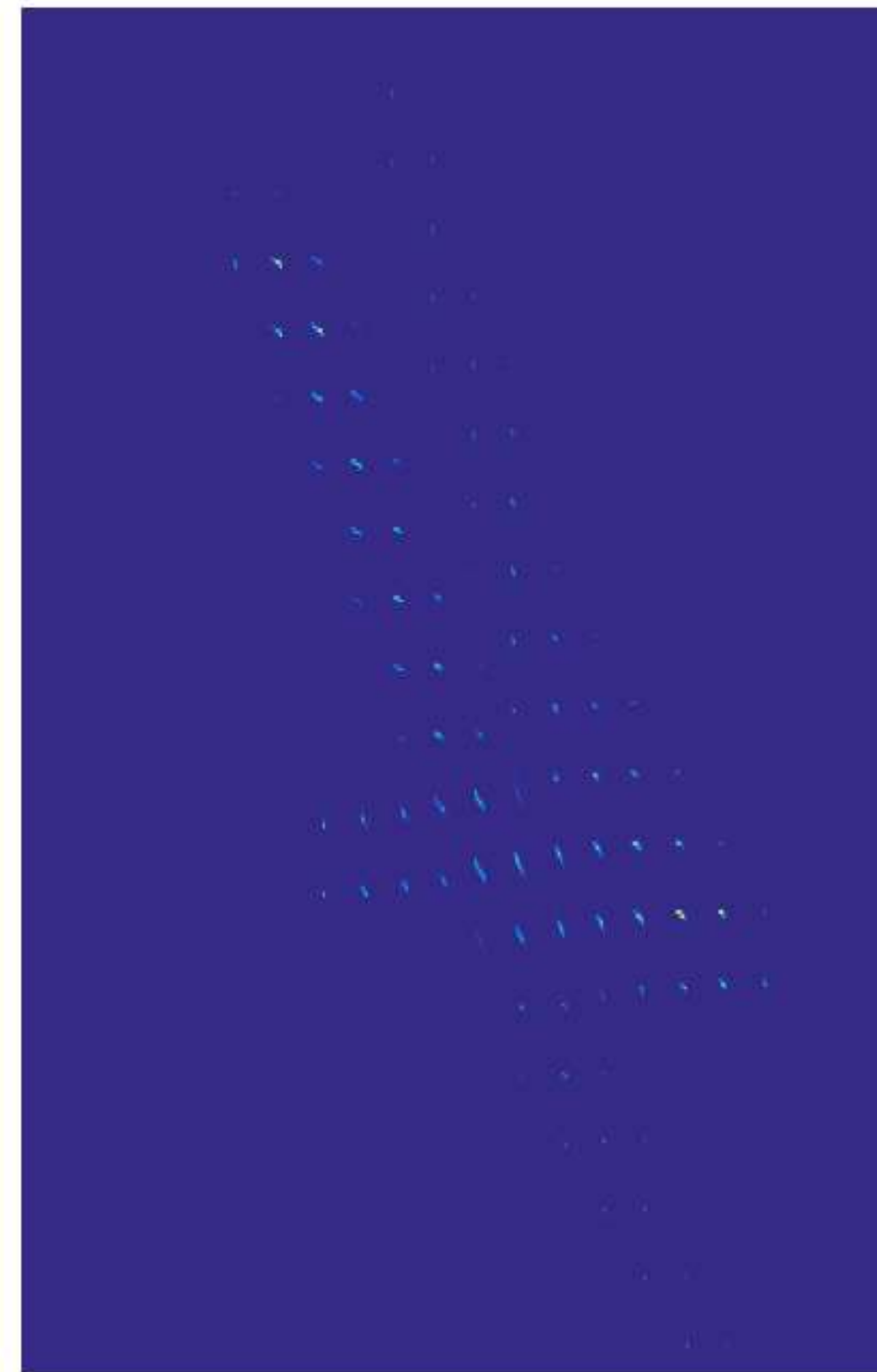


$\log(v_5)$

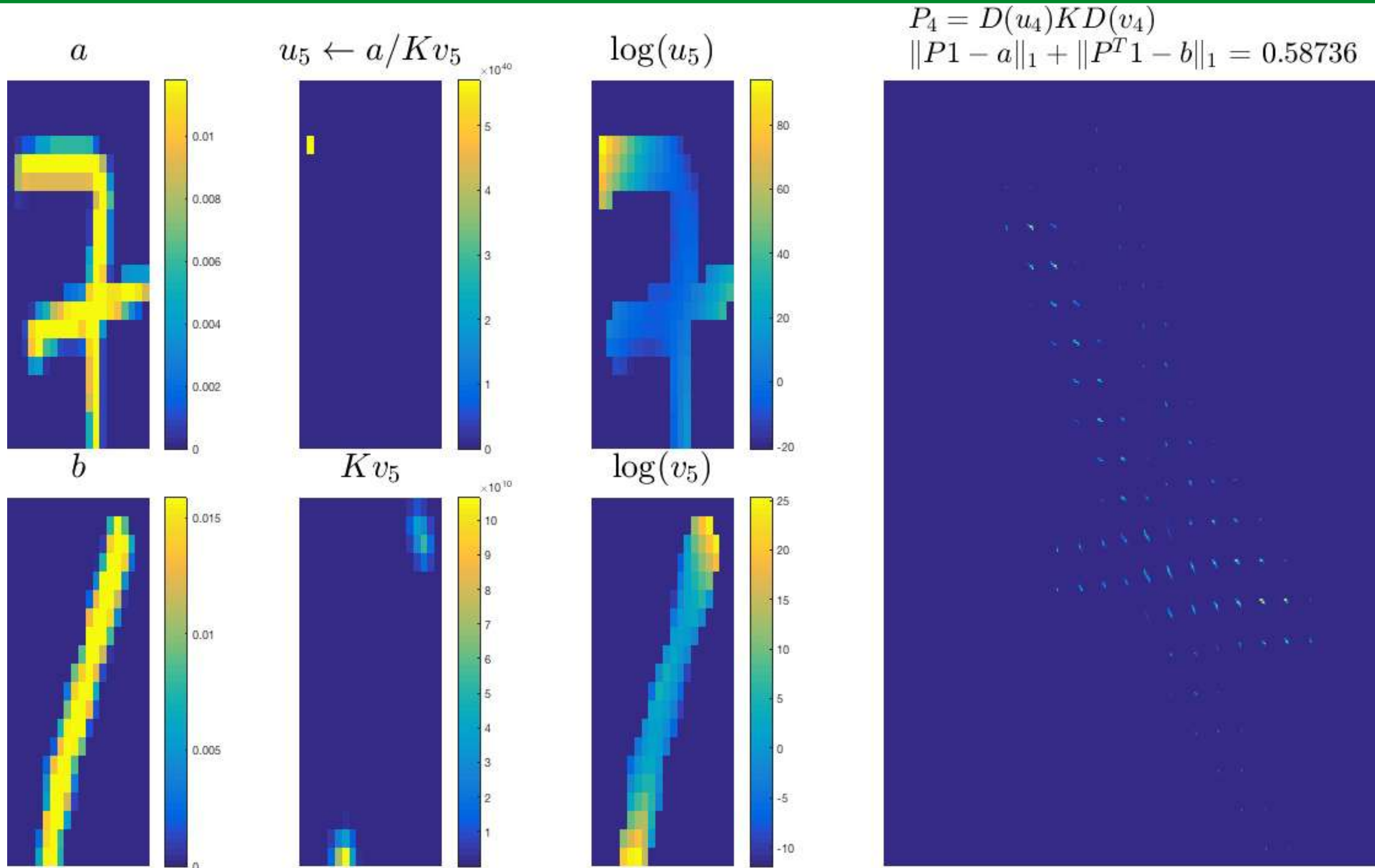


$$P_4 = D(u_4)KD(v_4)$$

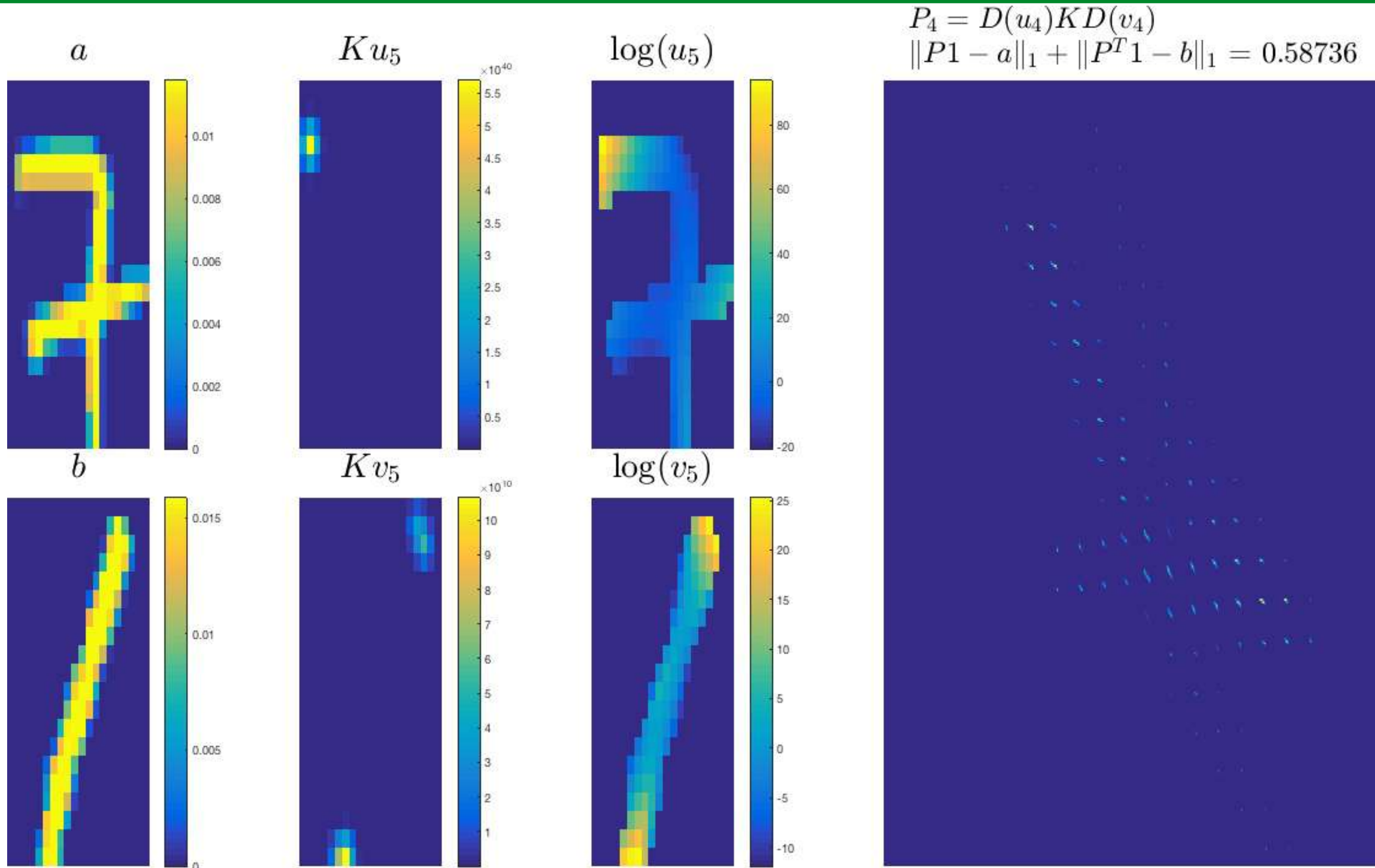
$$\|P1 - a\|_1 + \|P^T 1 - b\|_1 = 0.58736$$



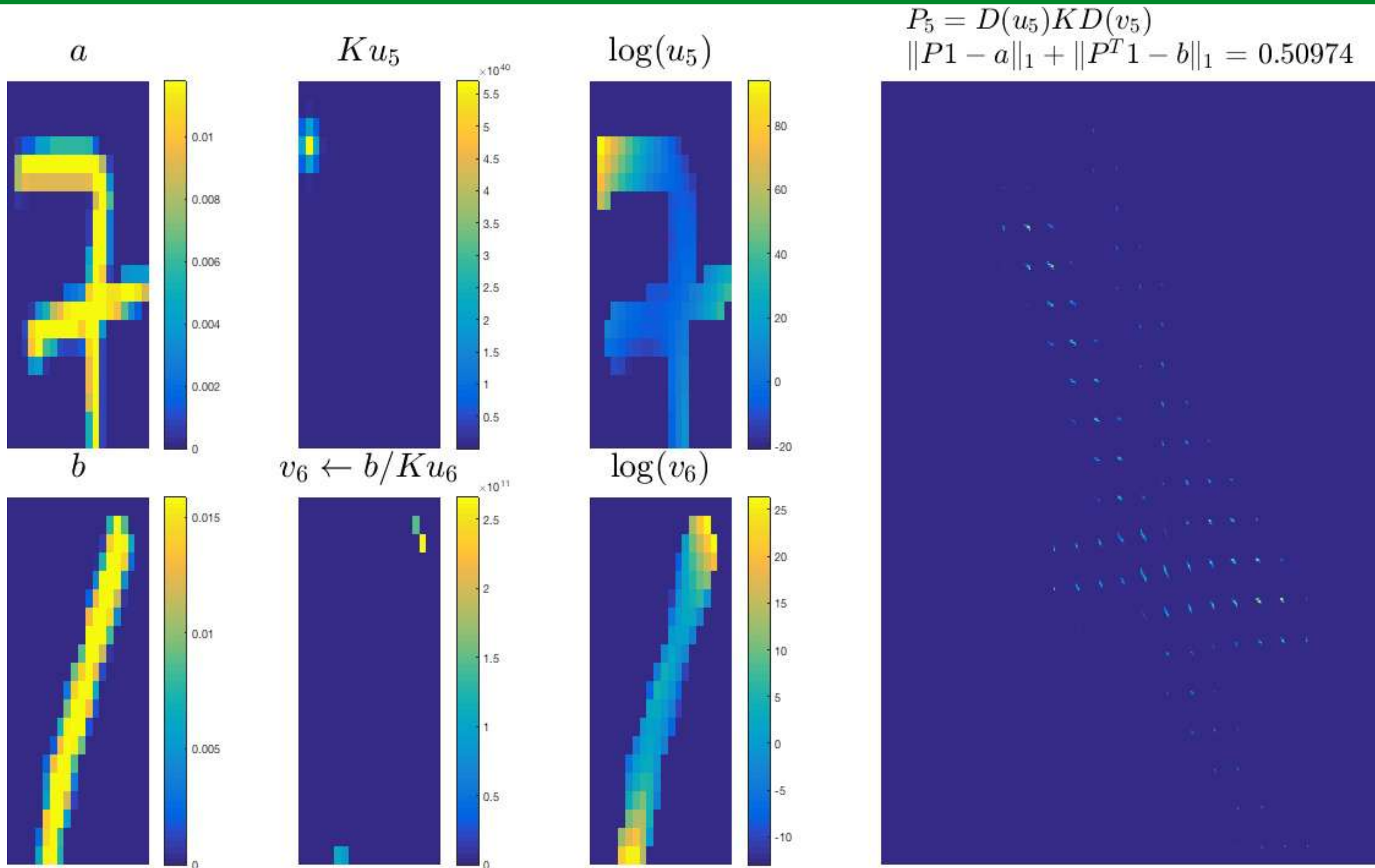
Very Fast EMD Approx. Solver



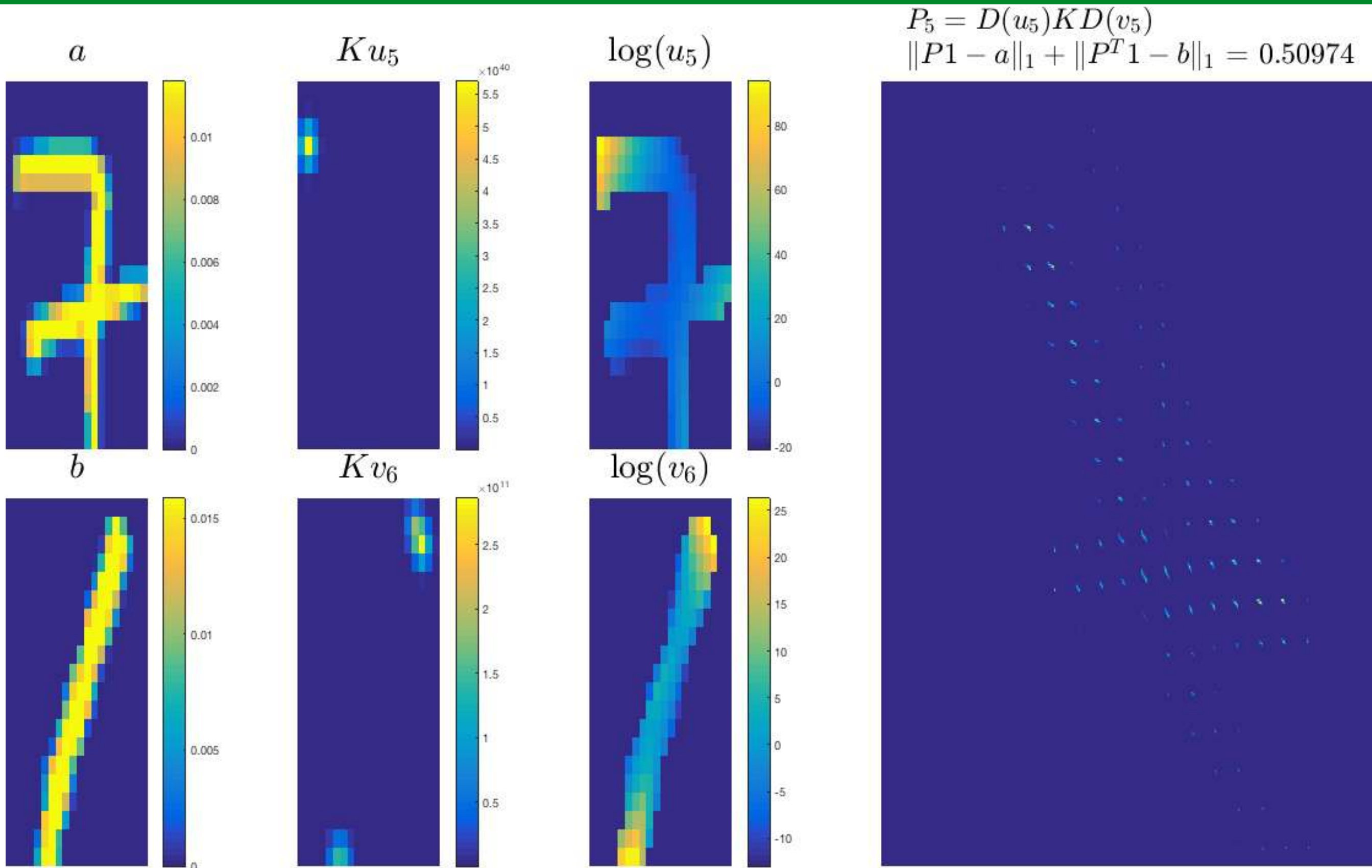
Very Fast EMD Approx. Solver



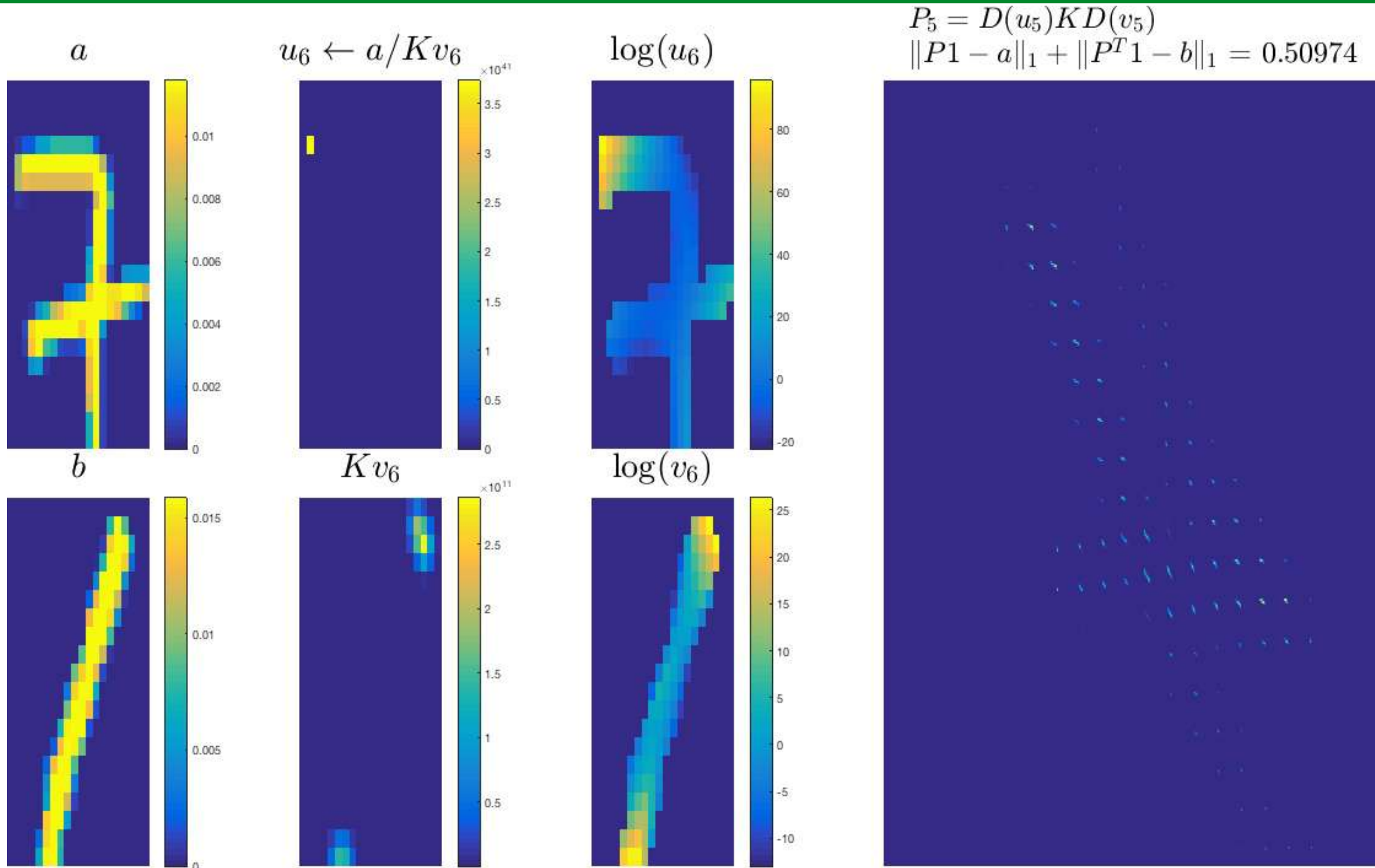
Very Fast EMD Approx. Solver



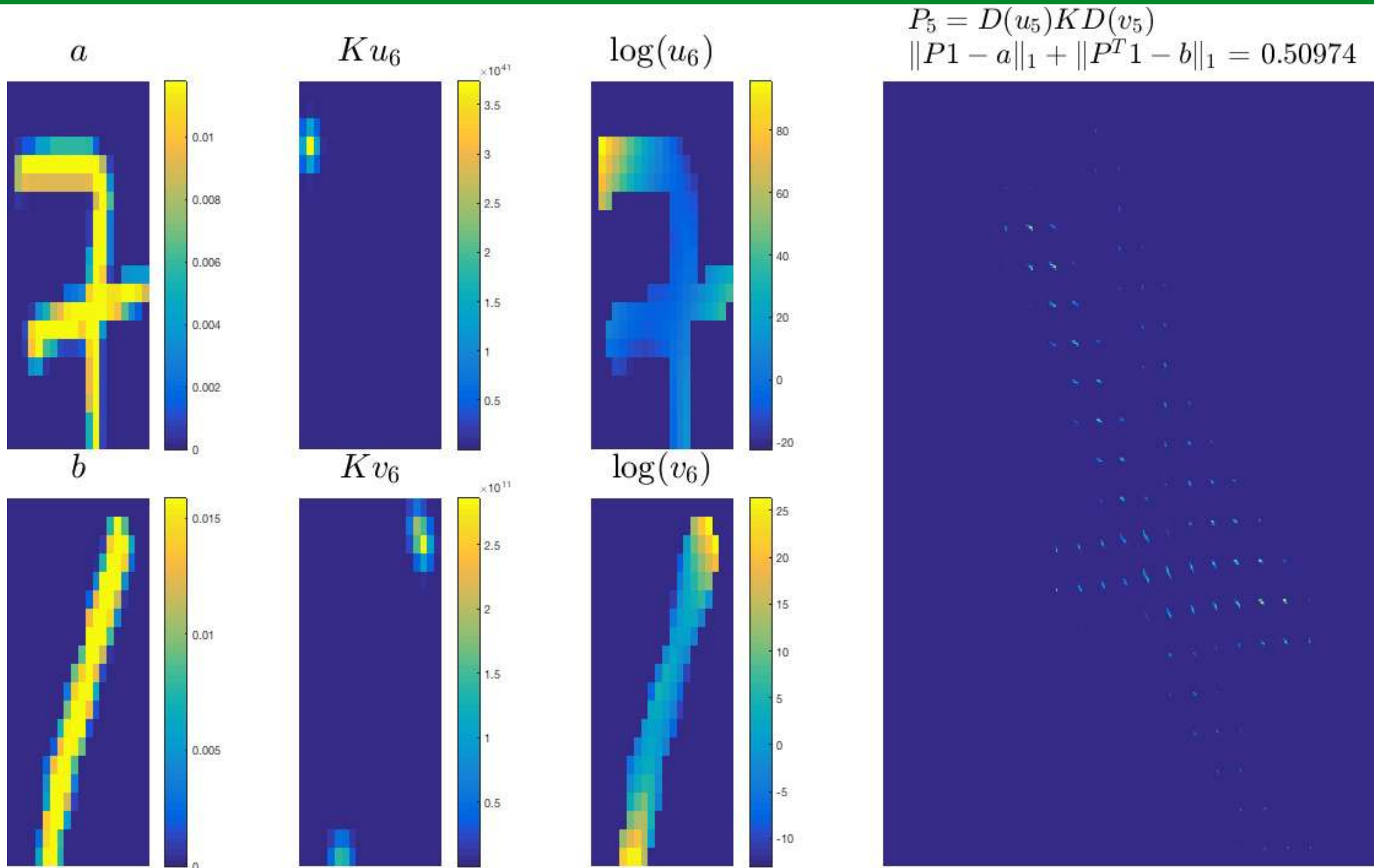
Very Fast EMD Approx. Solver



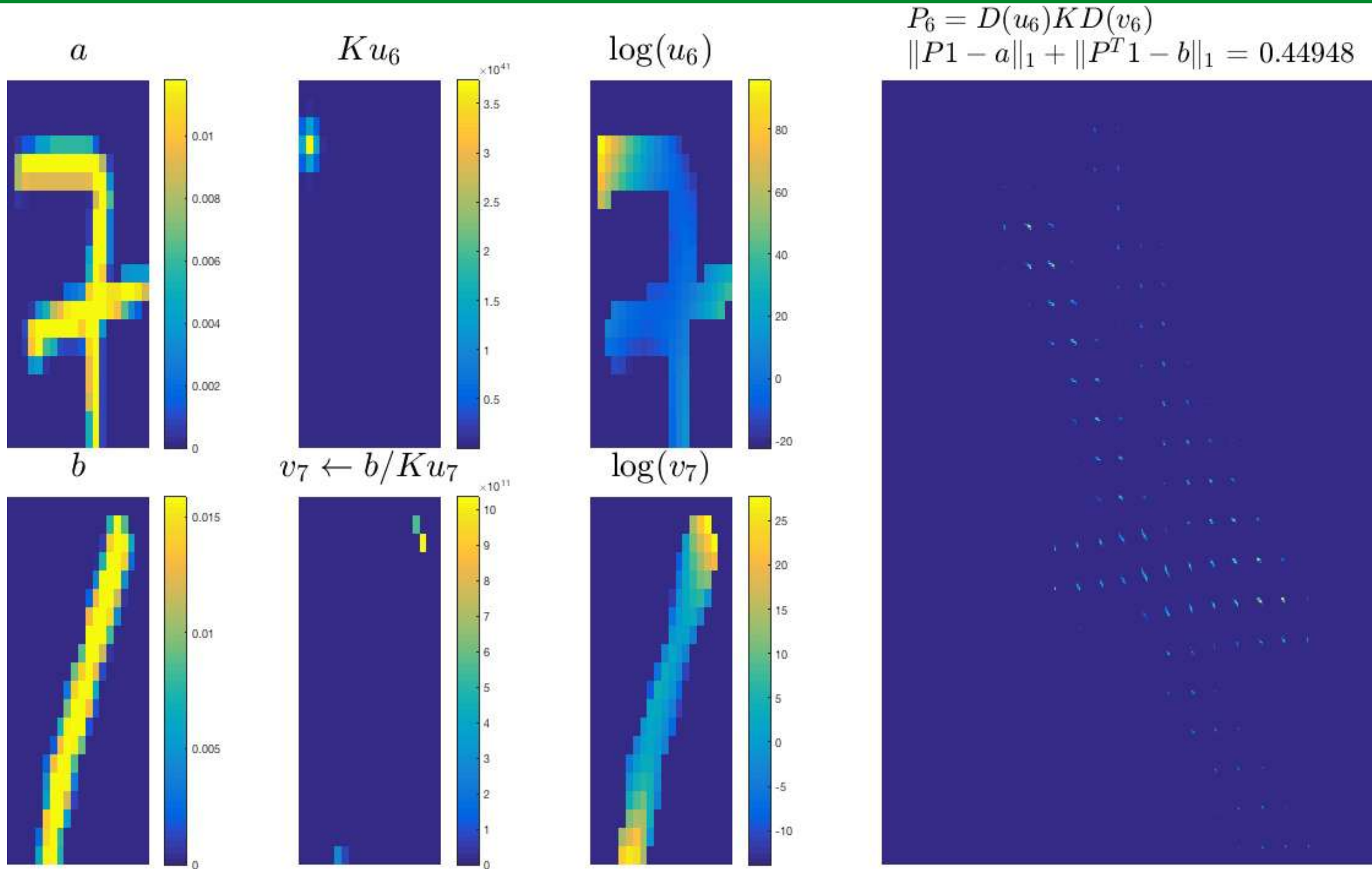
Very Fast EMD Approx. Solver



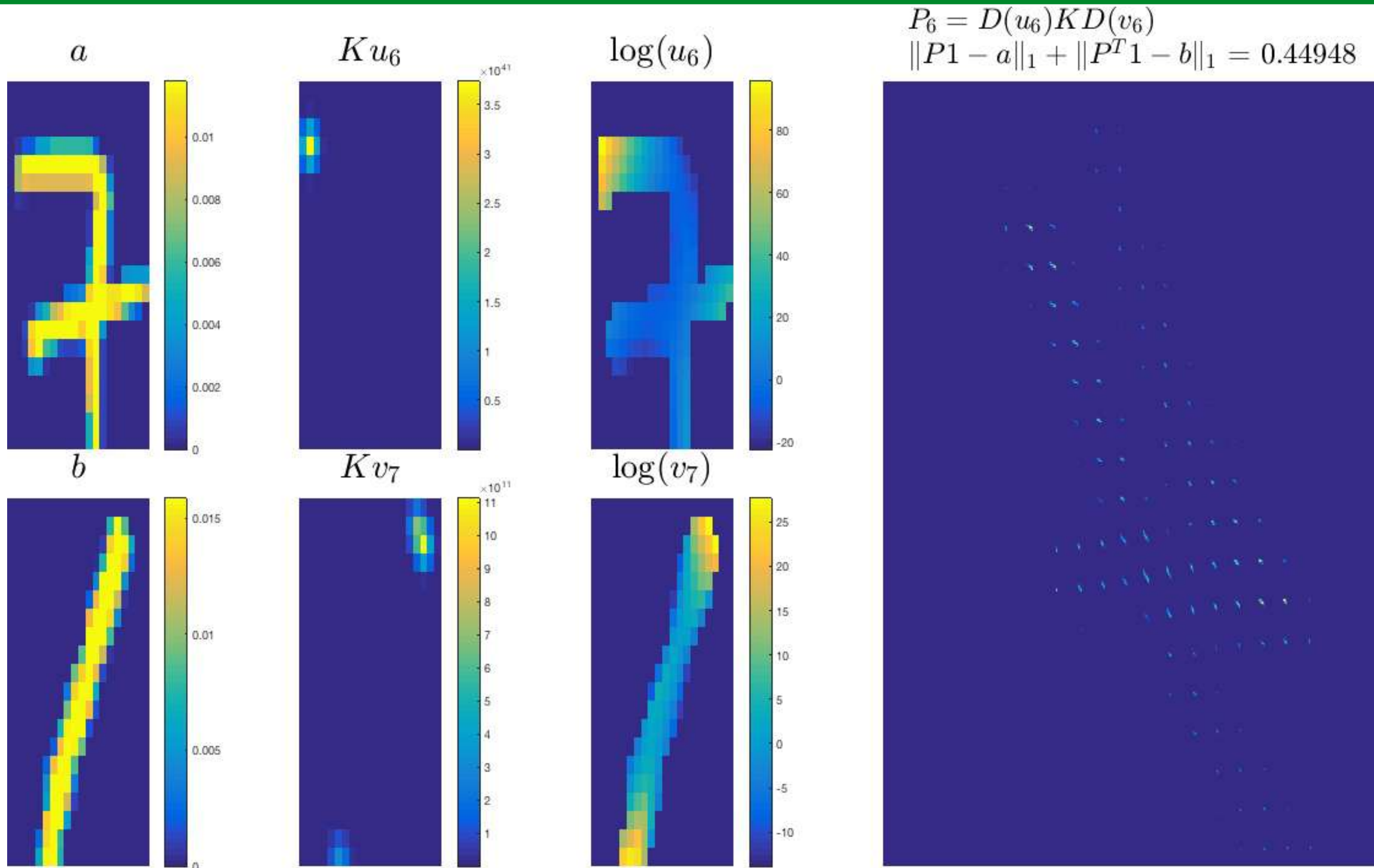
Very Fast EMD Approx. Solver



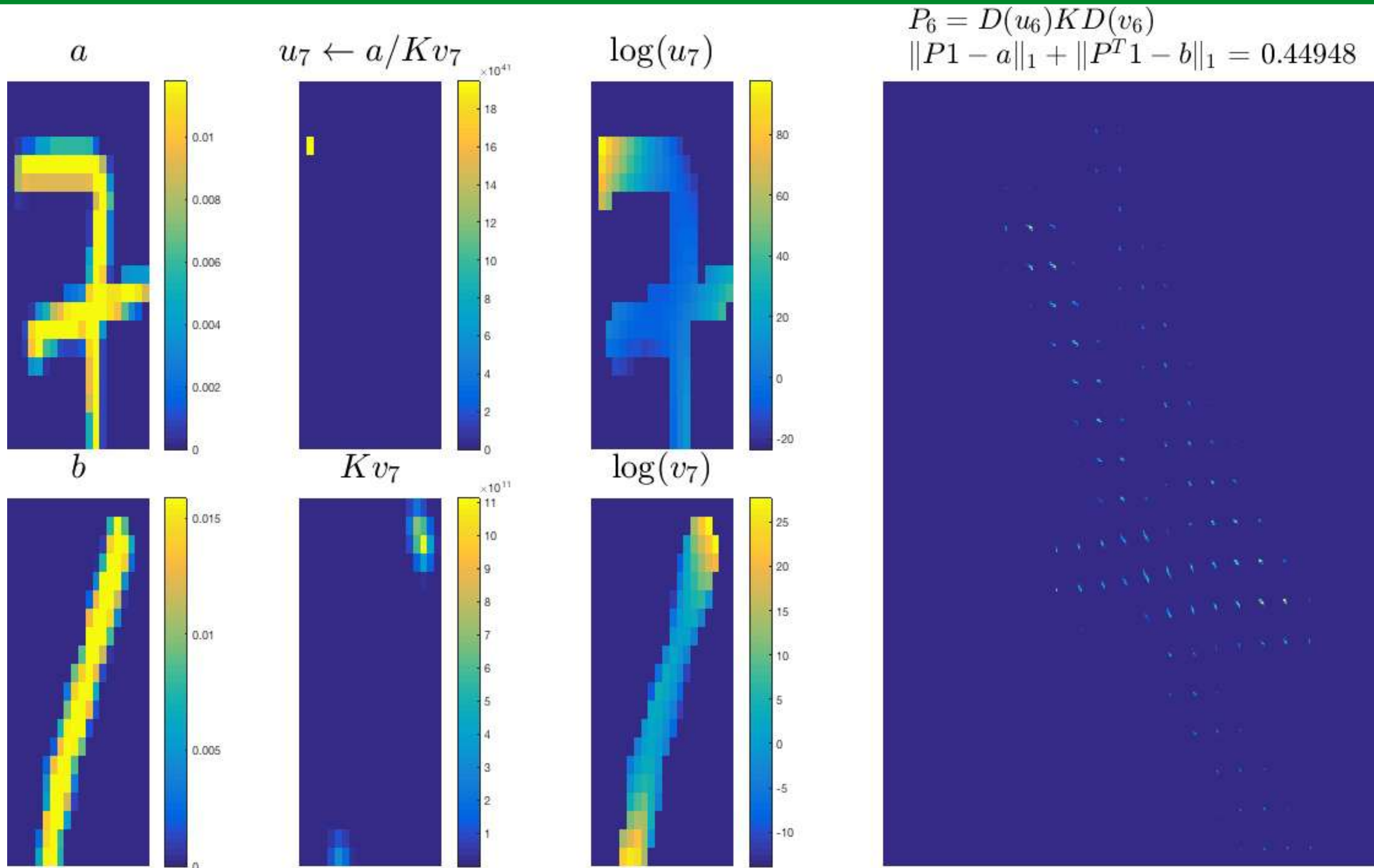
Very Fast EMD Approx. Solver



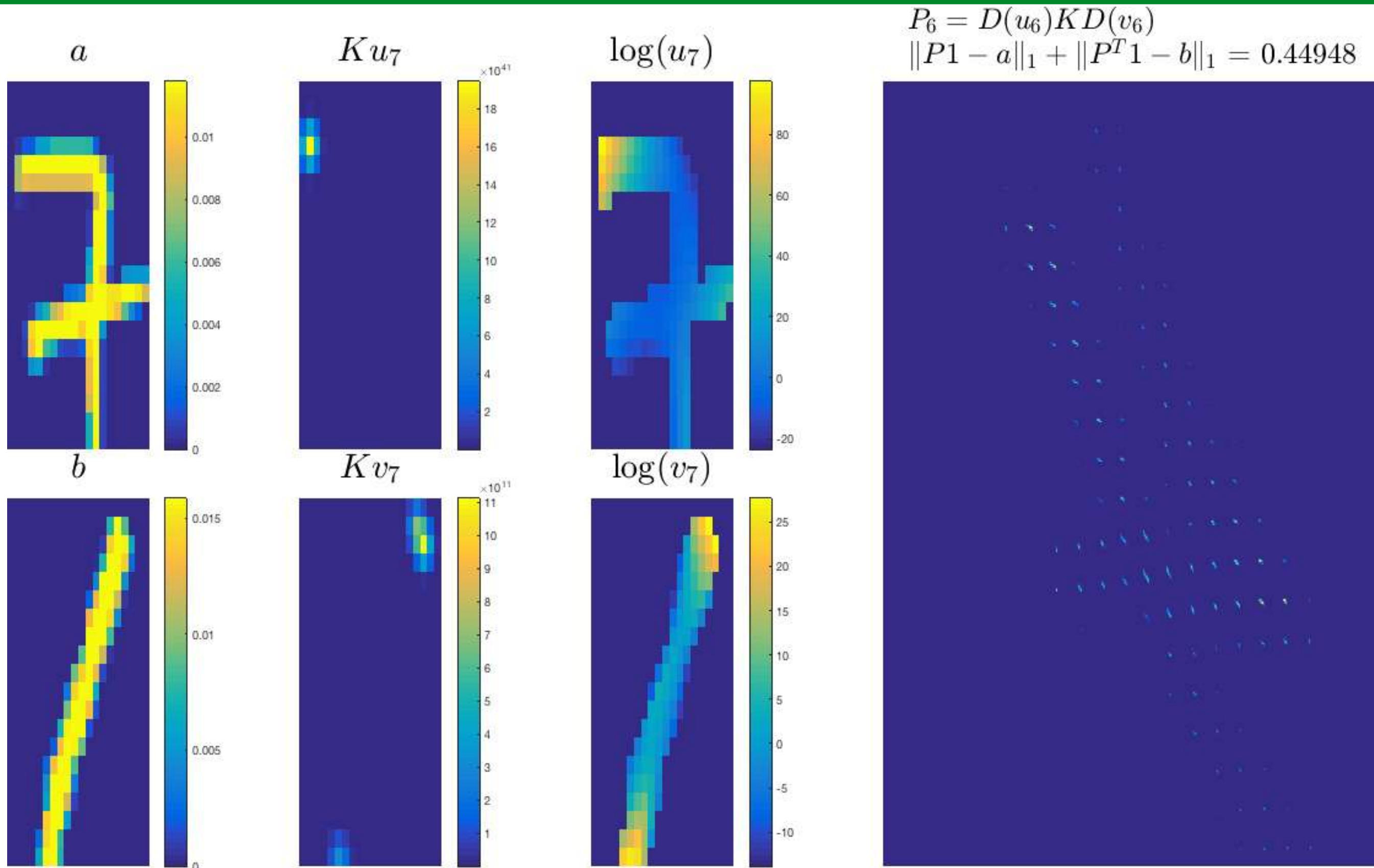
Very Fast EMD Approx. Solver



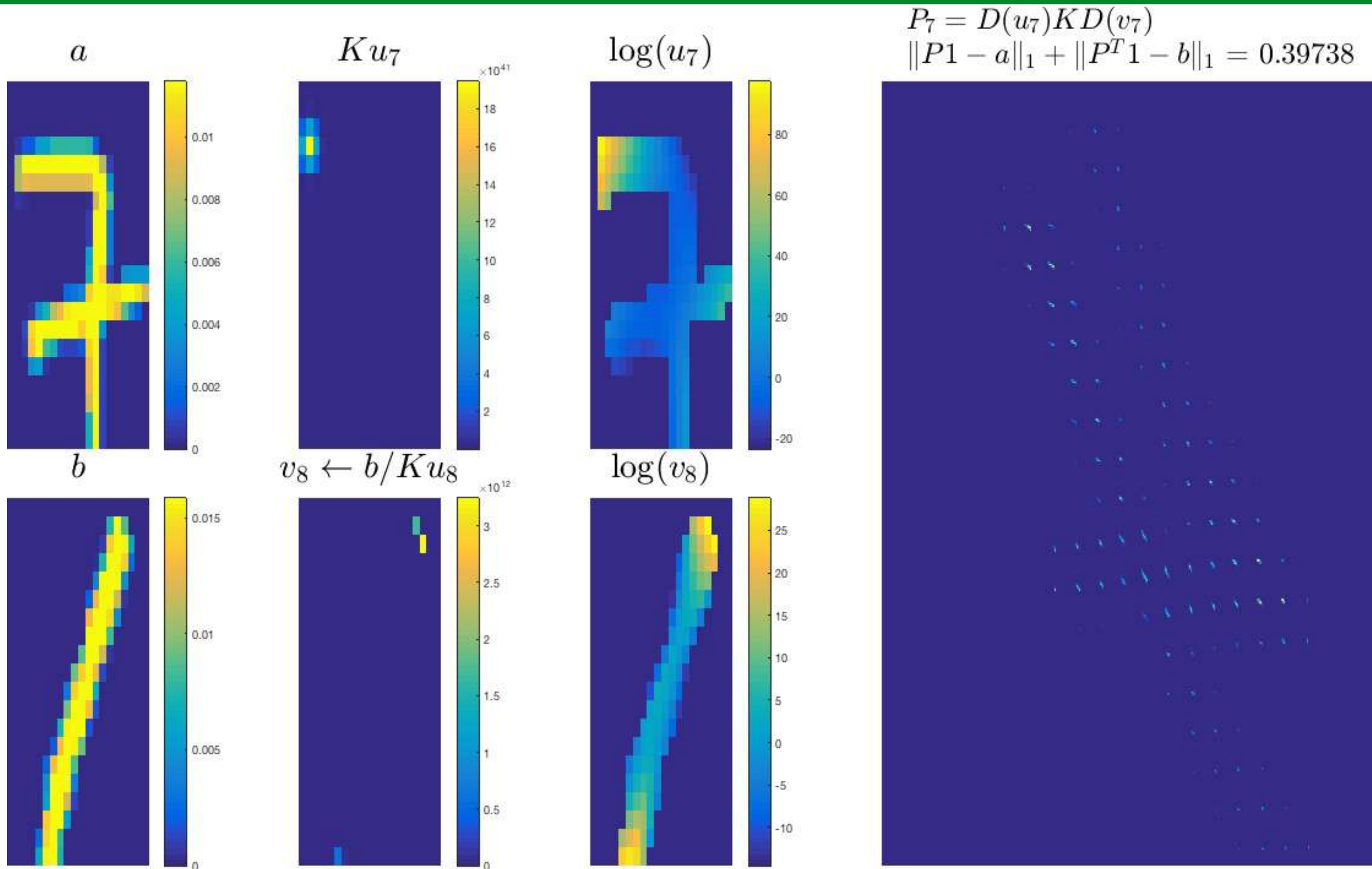
Very Fast EMD Approx. Solver



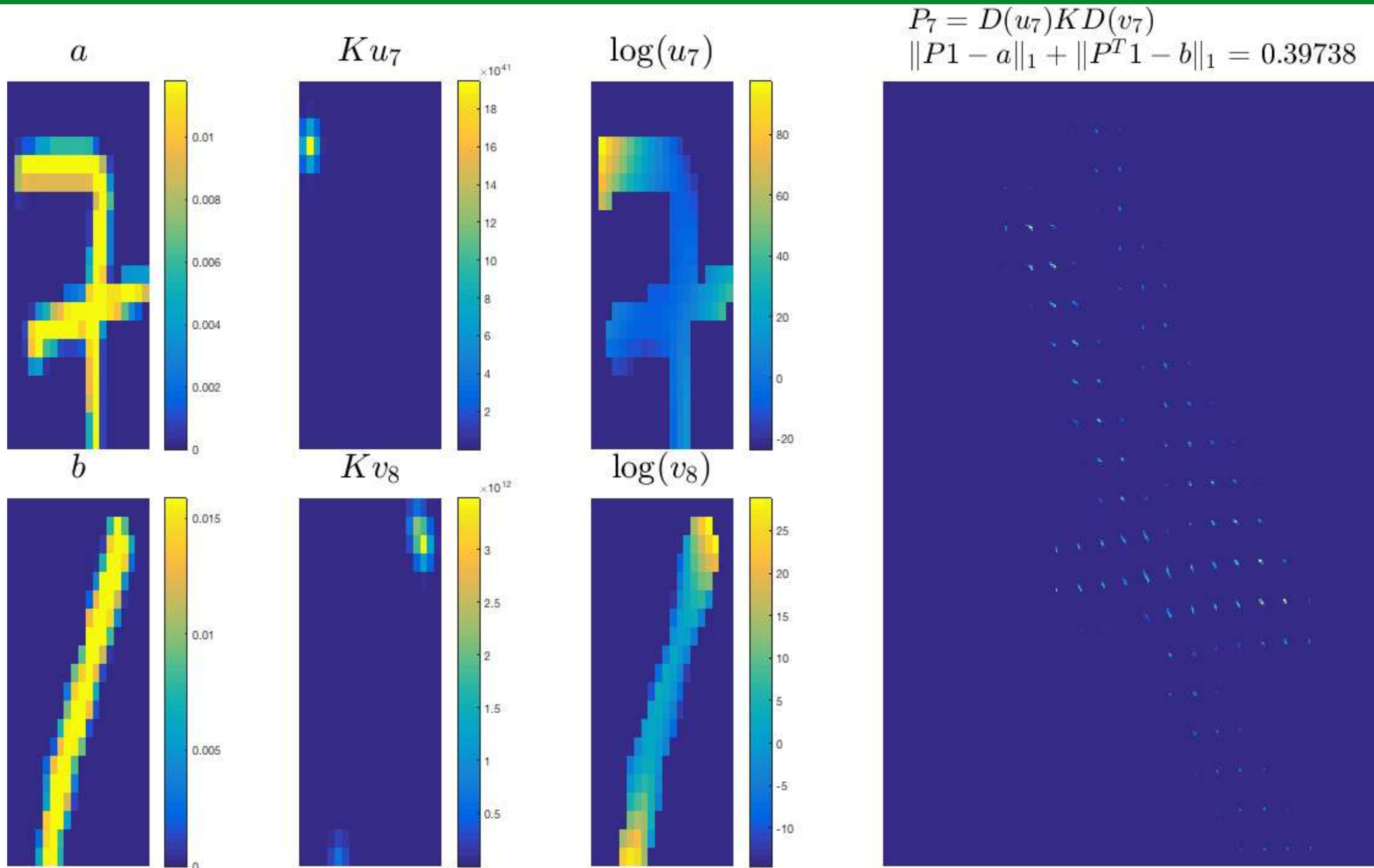
Very Fast EMD Approx. Solver



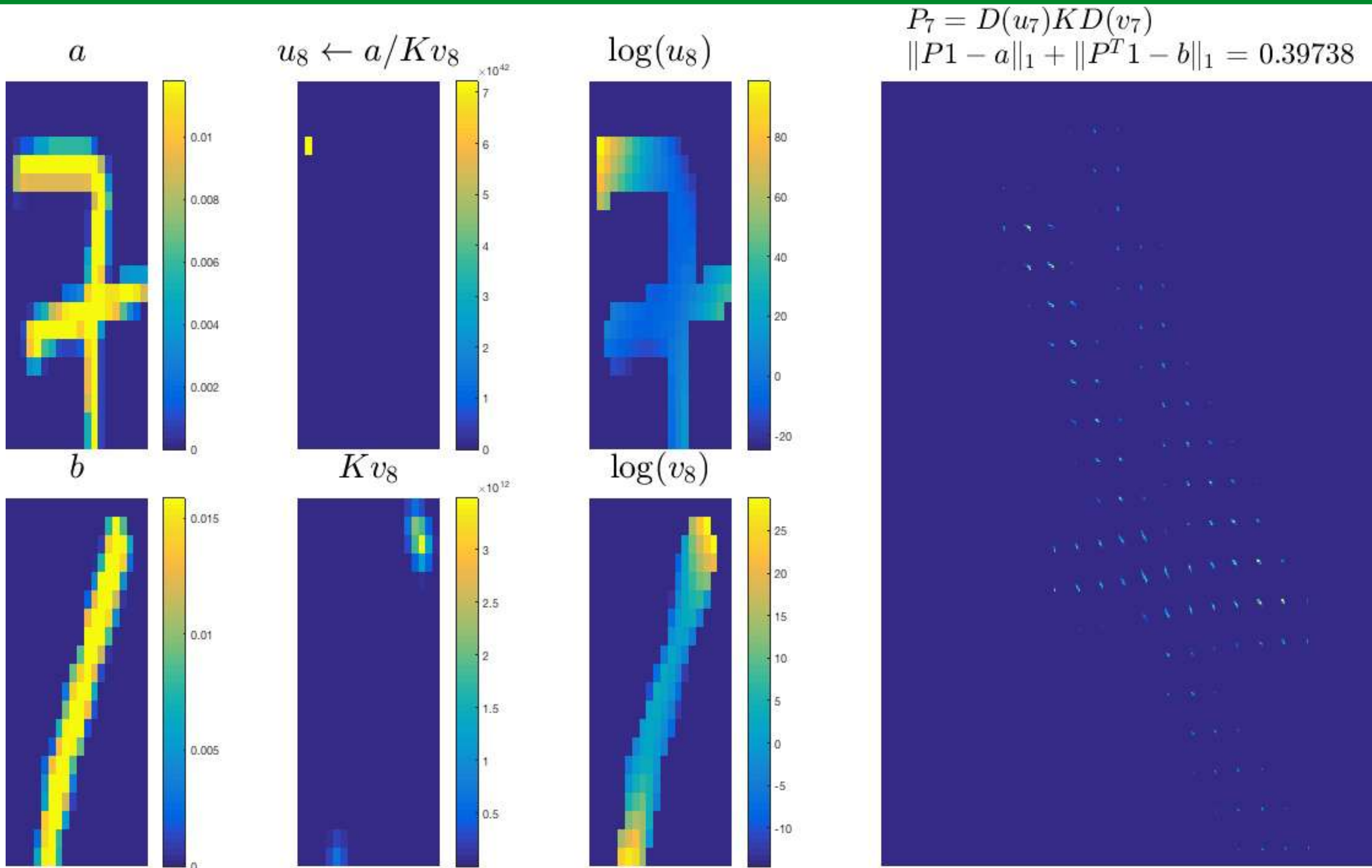
Very Fast EMD Approx. Solver



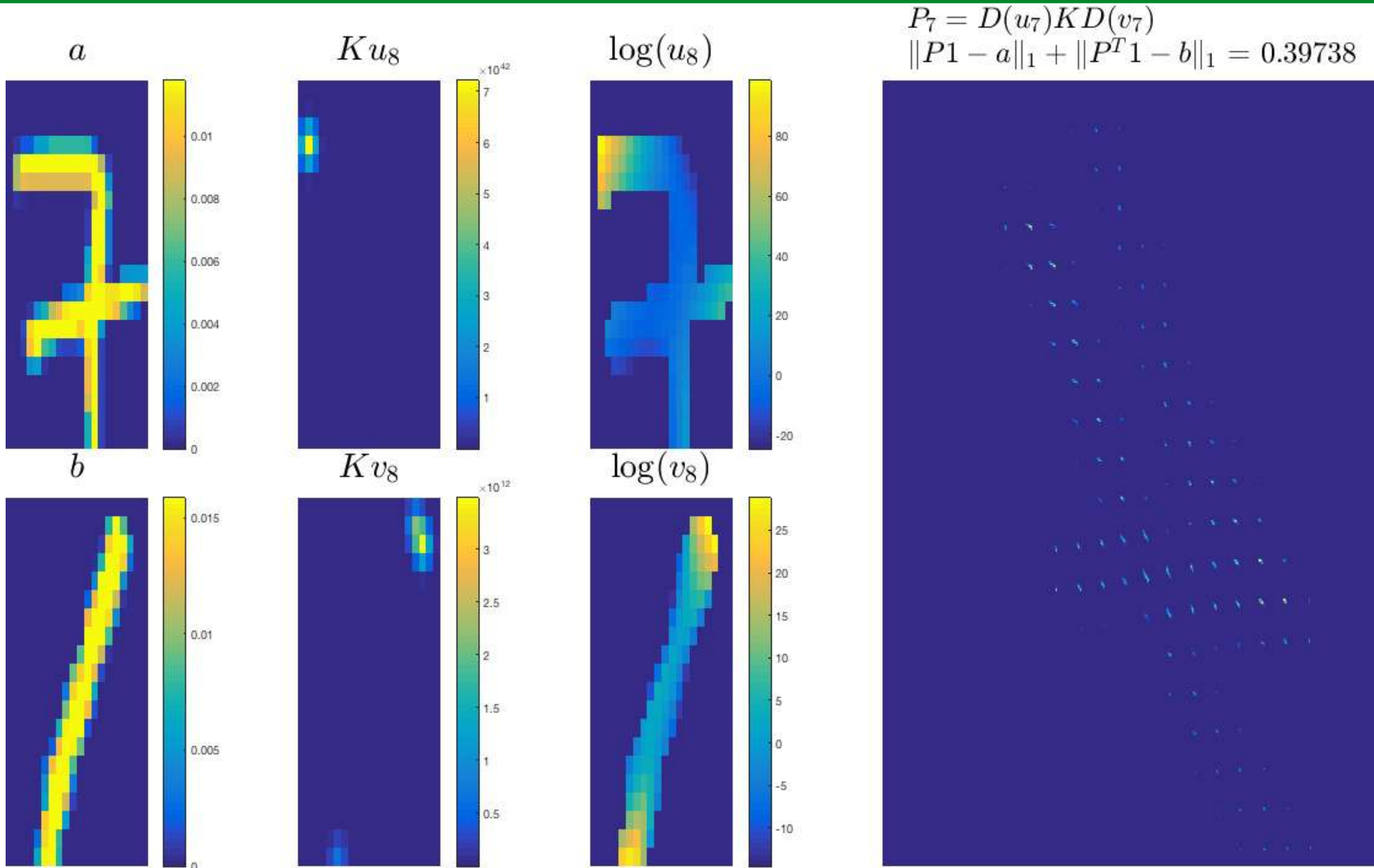
Very Fast EMD Approx. Solver



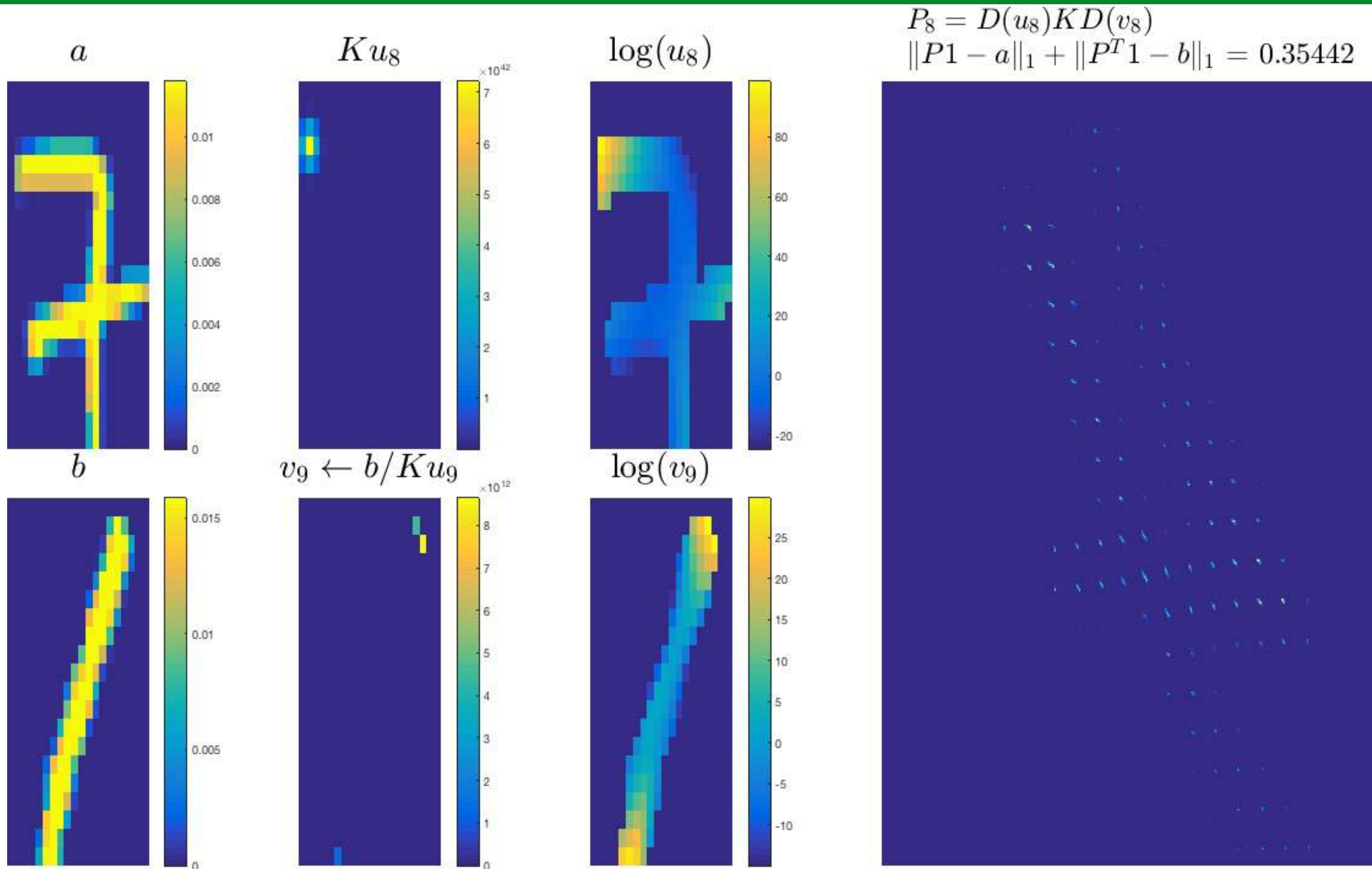
Very Fast EMD Approx. Solver



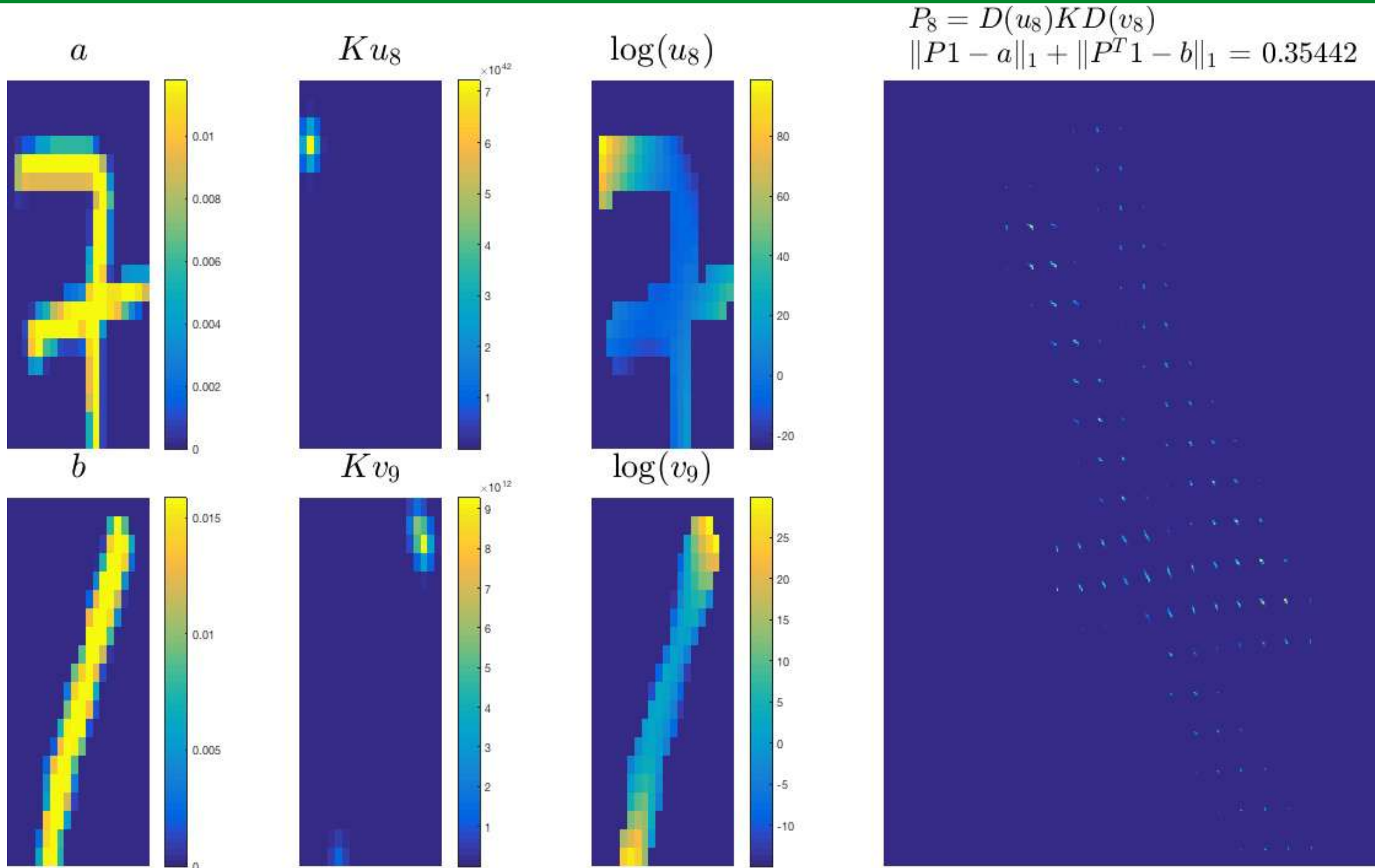
Very Fast EMD Approx. Solver



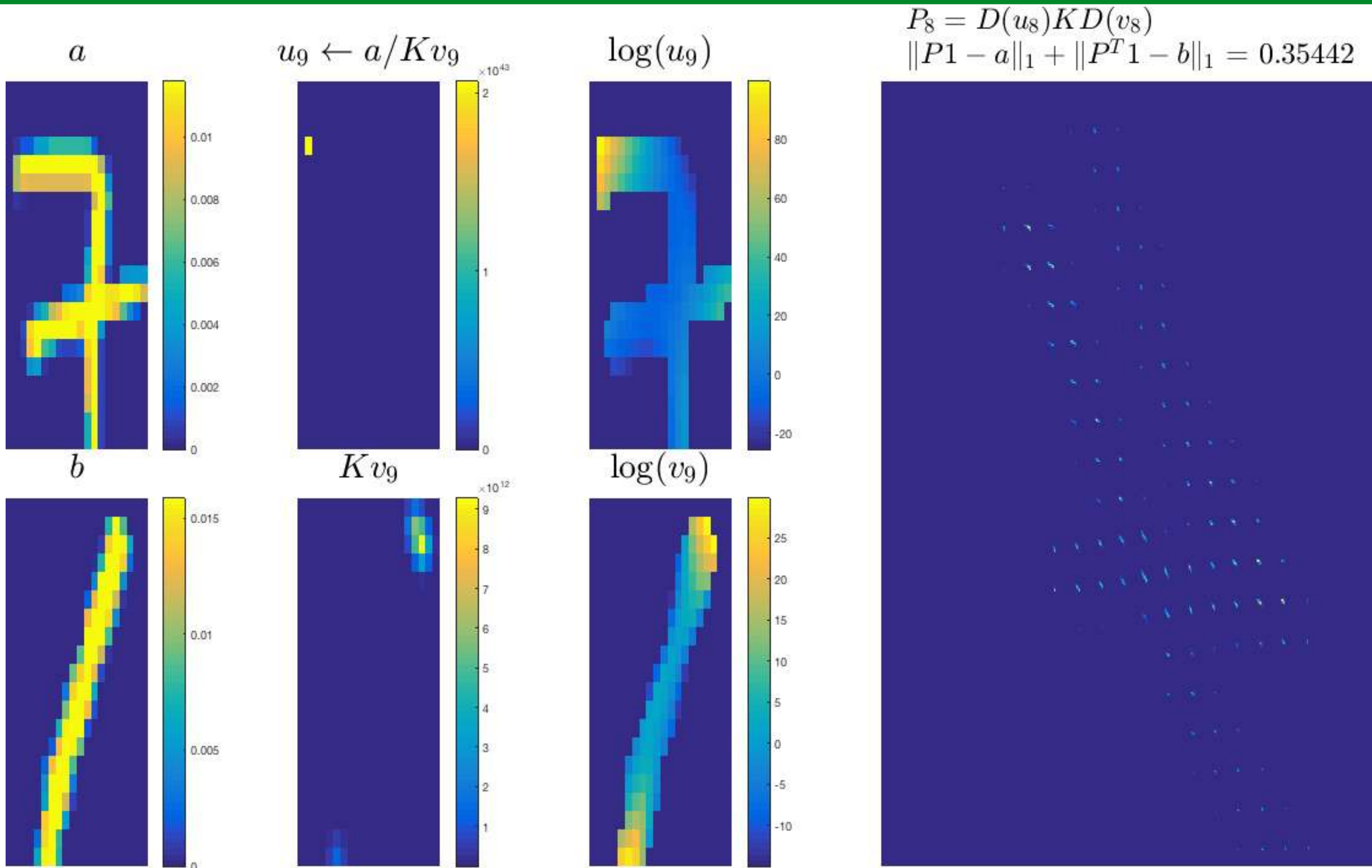
Very Fast EMD Approx. Solver



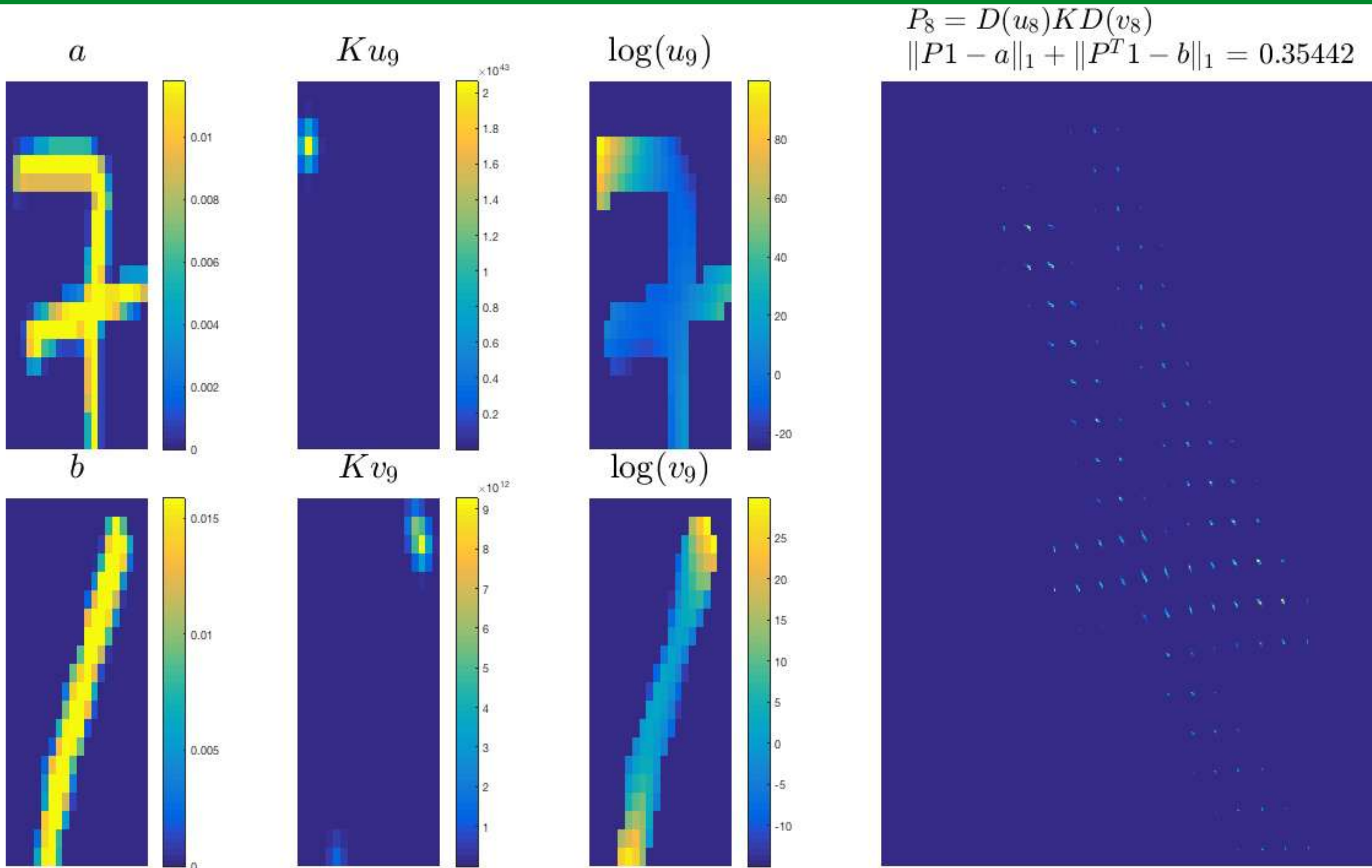
Very Fast EMD Approx. Solver



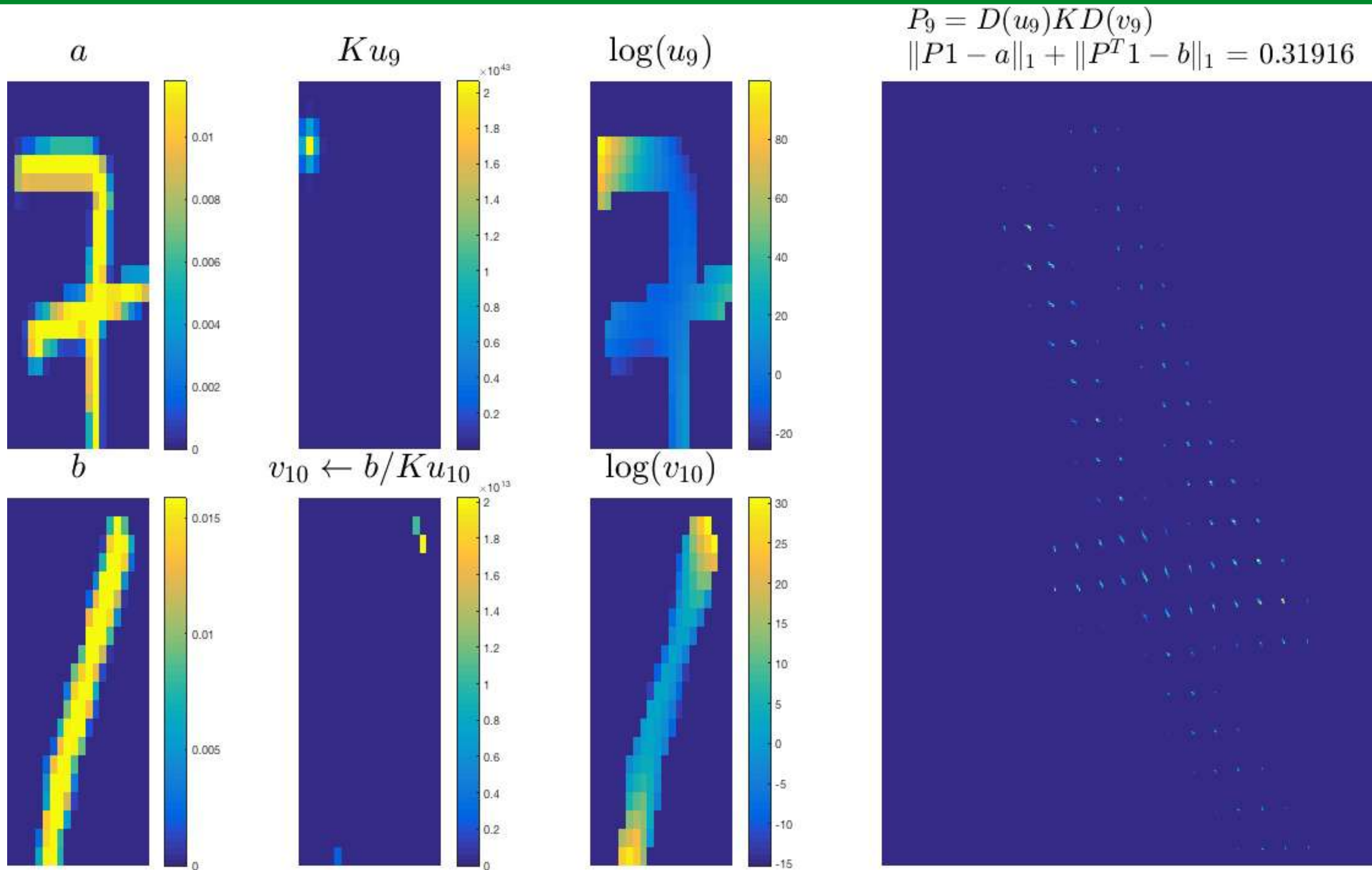
Very Fast EMD Approx. Solver



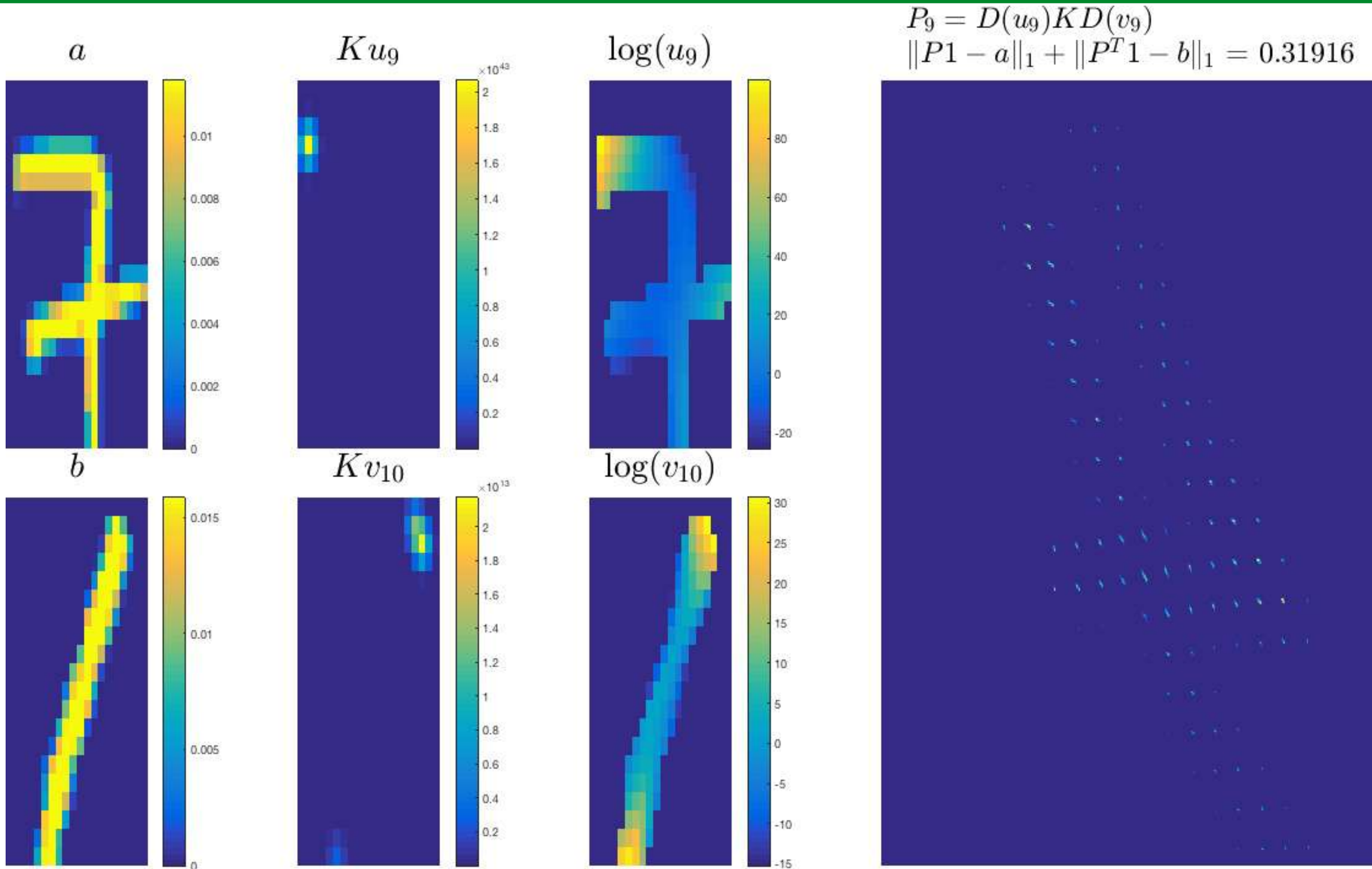
Very Fast EMD Approx. Solver



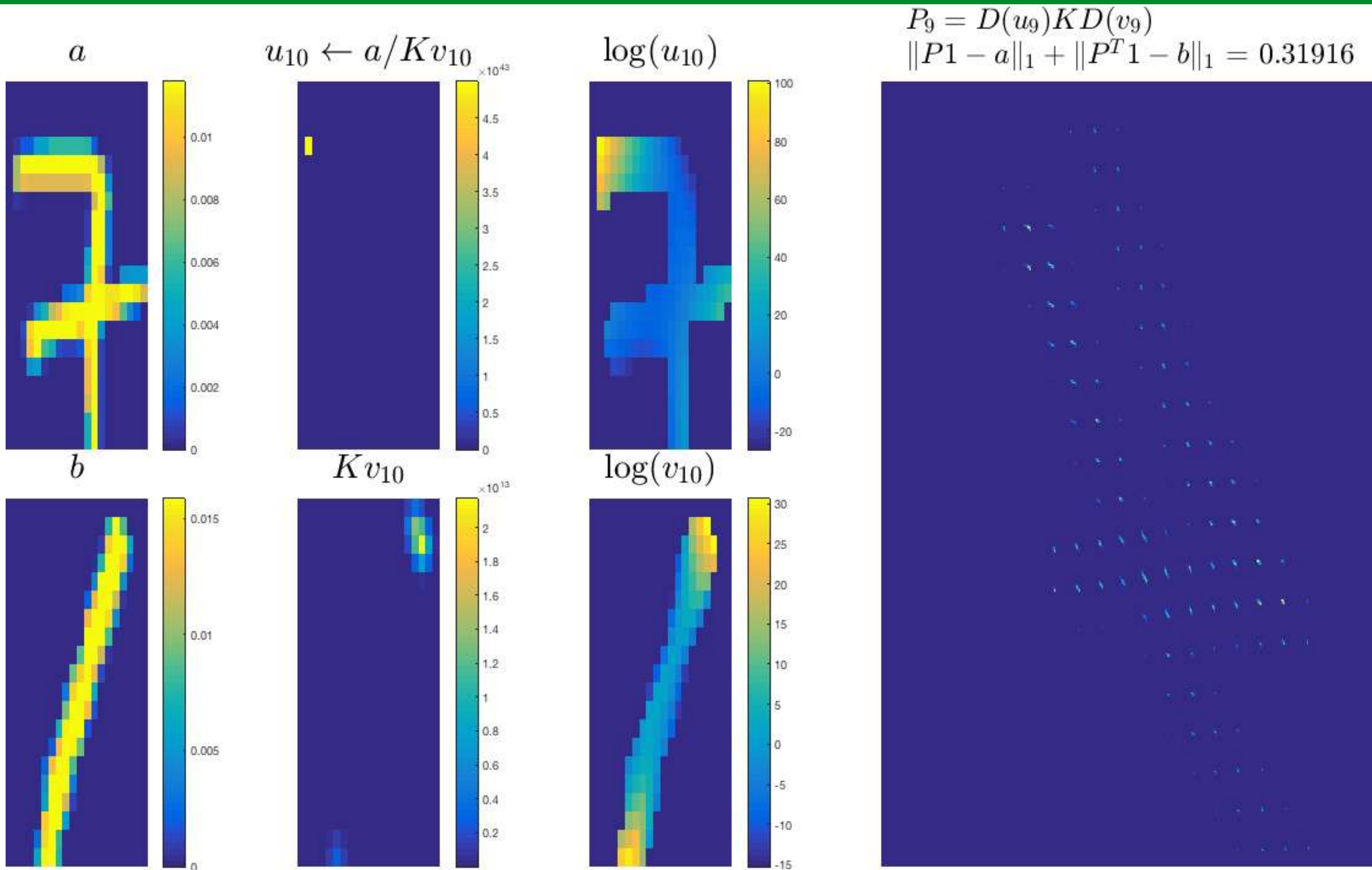
Very Fast EMD Approx. Solver



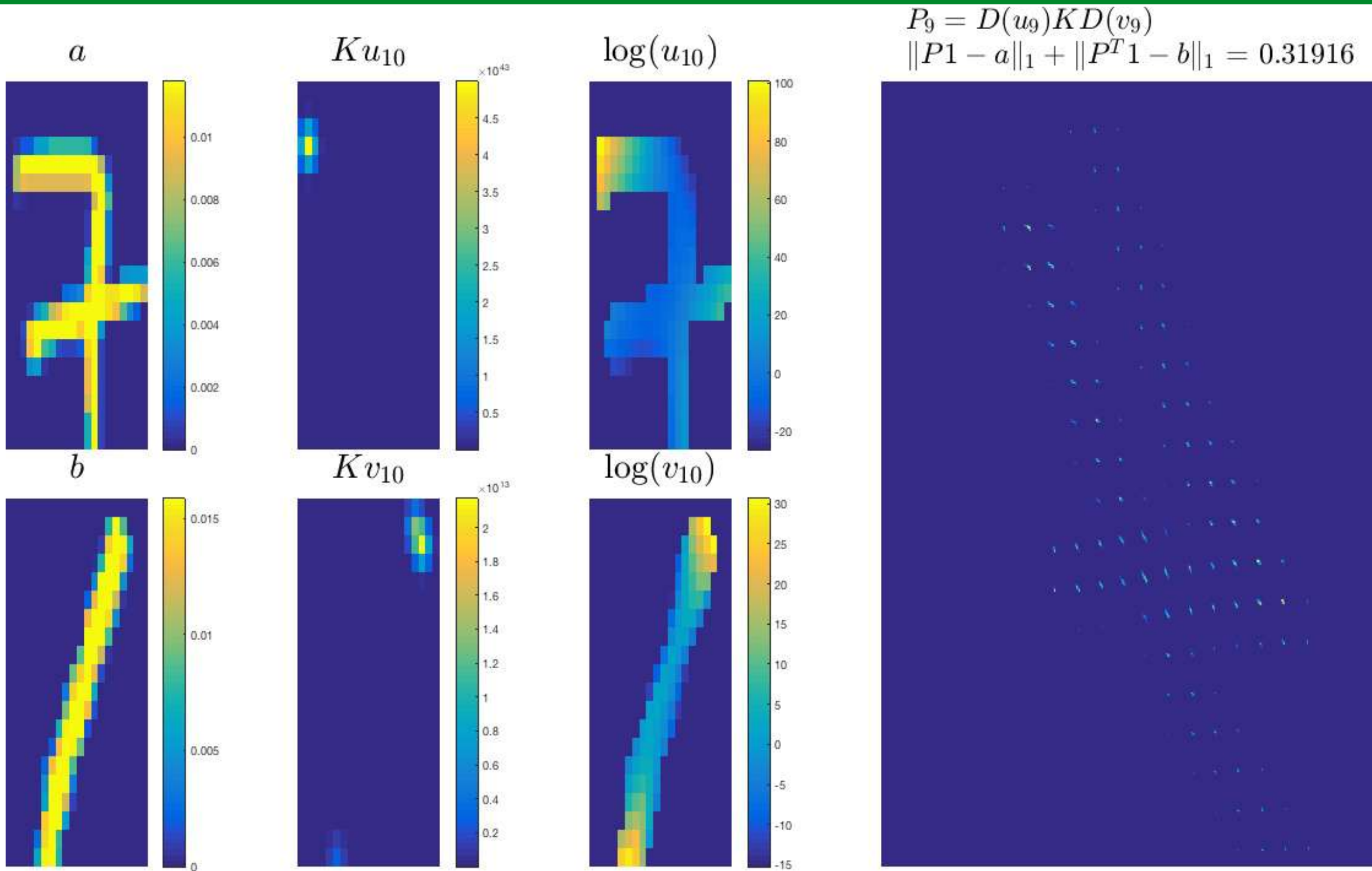
Very Fast EMD Approx. Solver



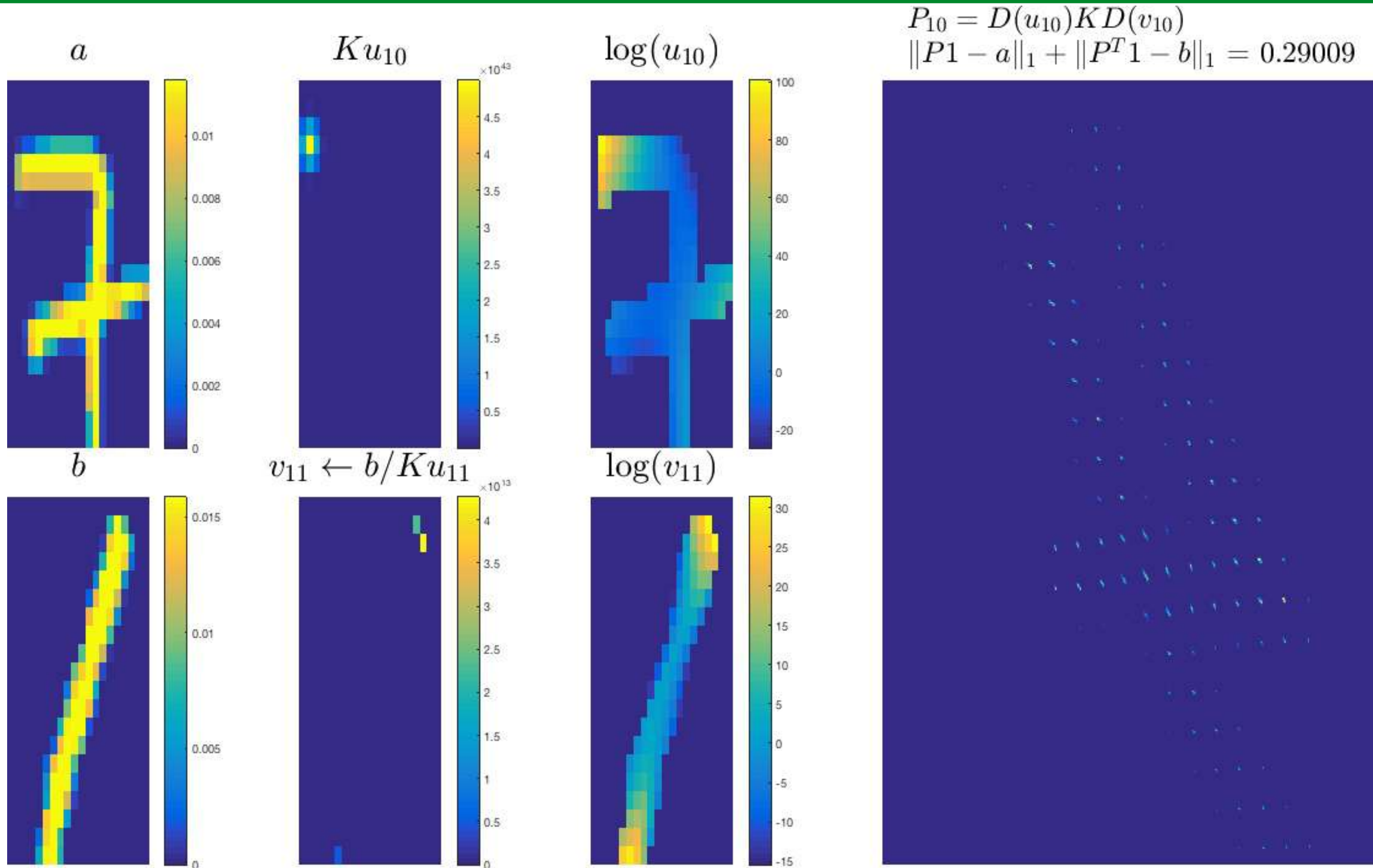
Very Fast EMD Approx. Solver



Very Fast EMD Approx. Solver



Very Fast EMD Approx. Solver



Sinkhorn as a Dual Algorithm

Def. Regularized Wasserstein, $\gamma \geq 0$

$$W_\gamma(\mu, \nu) \stackrel{\text{def}}{=} \min_{P \in U(a, b)} \langle P, M_{\mathbf{x}\mathbf{y}} \rangle - \gamma E(P)$$

REGULARIZED DISCRETE PRIMAL

$$W_\gamma(\mu, \nu) = \max_{\alpha, \beta} \alpha^T a + \beta^T b - \gamma (e^{\alpha/\gamma})^T K (e^{\beta/\gamma})$$

where $K = \left[e^{-\frac{D^p(\mathbf{x}_i, \mathbf{y}_j)}{\gamma}} \right]_{ij}$

REGULARIZED DISCRETE DUAL

Sinkhorn = *Block Coordinate Ascent* on Dual

Block Coordinate Ascent, *a.k.a Sinkhorn*

$$W_\gamma(\boldsymbol{\mu}, \boldsymbol{\nu}) = \max_{\boldsymbol{\alpha}, \boldsymbol{\beta}} \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} (e^{\boldsymbol{\beta}/\gamma})$$

REGULARIZED DISCRETE DUAL

$$\mathcal{E}(\boldsymbol{\alpha}, \boldsymbol{\beta}) = \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} e^{\boldsymbol{\beta}/\gamma}$$

Block Coordinate Ascent, *a.k.a Sinkhorn*

$$W_\gamma(\boldsymbol{\mu}, \boldsymbol{\nu}) = \max_{\boldsymbol{\alpha}, \boldsymbol{\beta}} \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} (e^{\boldsymbol{\beta}/\gamma})$$

REGULARIZED DISCRETE DUAL

$$\mathcal{E}(\boldsymbol{\alpha}, \boldsymbol{\beta}) = \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} e^{\boldsymbol{\beta}/\gamma}$$

$$\nabla_{\boldsymbol{\alpha}} \mathcal{E} = \boldsymbol{a} - e^{\boldsymbol{\alpha}/\gamma} \odot \boldsymbol{K} e^{\boldsymbol{\beta}/\gamma}$$

$$\nabla_{\boldsymbol{\beta}} \mathcal{E} = \boldsymbol{b} - e^{\boldsymbol{\beta}/\gamma} \odot \boldsymbol{K}^T e^{\boldsymbol{\alpha}/\gamma}$$

Block Coordinate Ascent, *a.k.a Sinkhorn*

$$W_\gamma(\boldsymbol{\mu}, \boldsymbol{\nu}) = \max_{\boldsymbol{\alpha}, \boldsymbol{\beta}} \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} (e^{\boldsymbol{\beta}/\gamma})$$

REGULARIZED DISCRETE DUAL

$$\mathcal{E}(\boldsymbol{\alpha}, \boldsymbol{\beta}) = \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} e^{\boldsymbol{\beta}/\gamma}$$

$$\nabla_{\boldsymbol{\alpha}} \mathcal{E} = \boldsymbol{a} - e^{\boldsymbol{\alpha}/\gamma} \odot \boldsymbol{K} e^{\boldsymbol{\beta}/\gamma}$$

$$\boldsymbol{\alpha} \leftarrow \gamma \left(\log \boldsymbol{a} - \log \boldsymbol{K} (e^{\boldsymbol{\beta}/\gamma}) \right)$$

$$\nabla_{\boldsymbol{\beta}} \mathcal{E} = \boldsymbol{b} - e^{\boldsymbol{\beta}/\gamma} \odot \boldsymbol{K}^T e^{\boldsymbol{\alpha}/\gamma}$$

$$\boldsymbol{\beta} \leftarrow \gamma \left(\log \boldsymbol{b} - \log \boldsymbol{K}^T (e^{\boldsymbol{\alpha}/\gamma}) \right)$$

Block Coordinate Ascent, *a.k.a Sinkhorn*

$$W_\gamma(\mu, \nu) = \max_{\alpha, \beta} \alpha^T a + \beta^T b - \gamma (e^{\alpha/\gamma})^T K (e^{\beta/\gamma})$$

REGULARIZED DISCRETE DUAL

Block Coordinate Ascent, *a.k.a* Sinkhorn

$$W_\gamma(\boldsymbol{\mu}, \boldsymbol{\nu}) = \max_{\boldsymbol{\alpha}, \boldsymbol{\beta}} \boldsymbol{\alpha}^T \boldsymbol{a} + \boldsymbol{\beta}^T \boldsymbol{b} - \gamma (e^{\boldsymbol{\alpha}/\gamma})^T \boldsymbol{K} (e^{\boldsymbol{\beta}/\gamma})$$

REGULARIZED DISCRETE DUAL

$$(\boldsymbol{u}, \boldsymbol{v}) \stackrel{\text{def}}{=} (e^{\boldsymbol{\alpha}/\gamma}, e^{\boldsymbol{\beta}/\gamma})$$

$$\boldsymbol{\alpha} \leftarrow \gamma \left(\log \boldsymbol{a} - \log \boldsymbol{K} (e^{\boldsymbol{\beta}/\gamma}) \right)$$

$$\boldsymbol{u} \leftarrow \frac{\boldsymbol{a}}{\boldsymbol{K} \boldsymbol{v}}$$

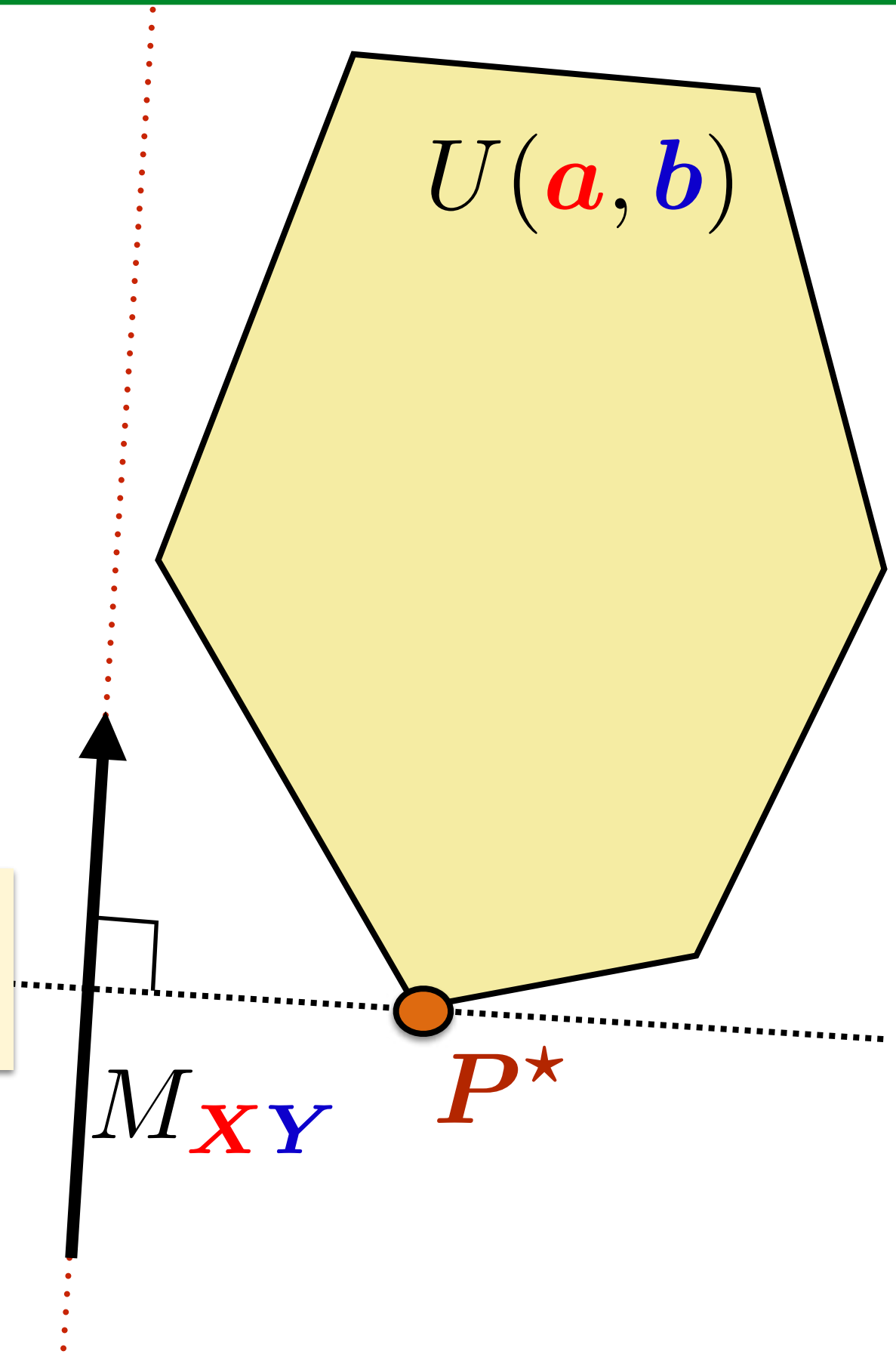
$$\boldsymbol{\beta} \leftarrow \gamma \left(\log \boldsymbol{b} - \log \boldsymbol{K}^T (e^{\boldsymbol{\alpha}/\gamma}) \right)$$

$$\boldsymbol{v} \leftarrow \frac{\boldsymbol{b}}{\boldsymbol{K}^T \boldsymbol{u}}$$

Sinkhorn, W and *Energy Distance*

$$\mu = \sum_{i=1}^n a_i \delta_{x_i} \quad \nu = \sum_{j=1}^m b_j \delta_{y_j}$$

$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{x}\mathbf{y}} \rangle$$

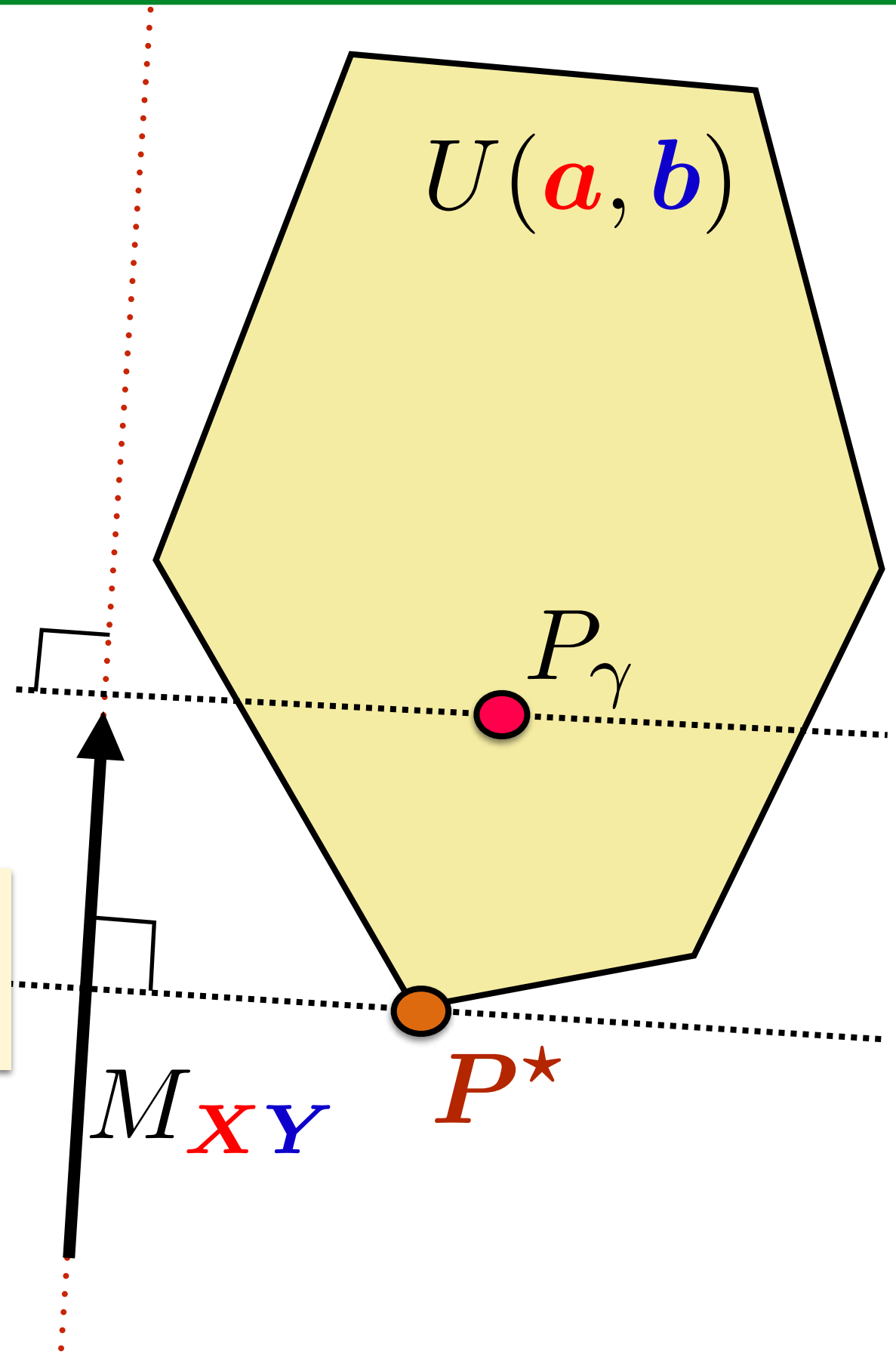


Sinkhorn, W and *Energy Distance*

$$\mu = \sum_{i=1}^n a_i \delta_{x_i} \quad \nu = \sum_{j=1}^m b_j \delta_{y_j}$$

$$W_\gamma(\mu, \nu) = \langle P_\gamma, M_{\mathbf{x}\mathbf{y}} \rangle$$

$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{x}\mathbf{y}} \rangle$$



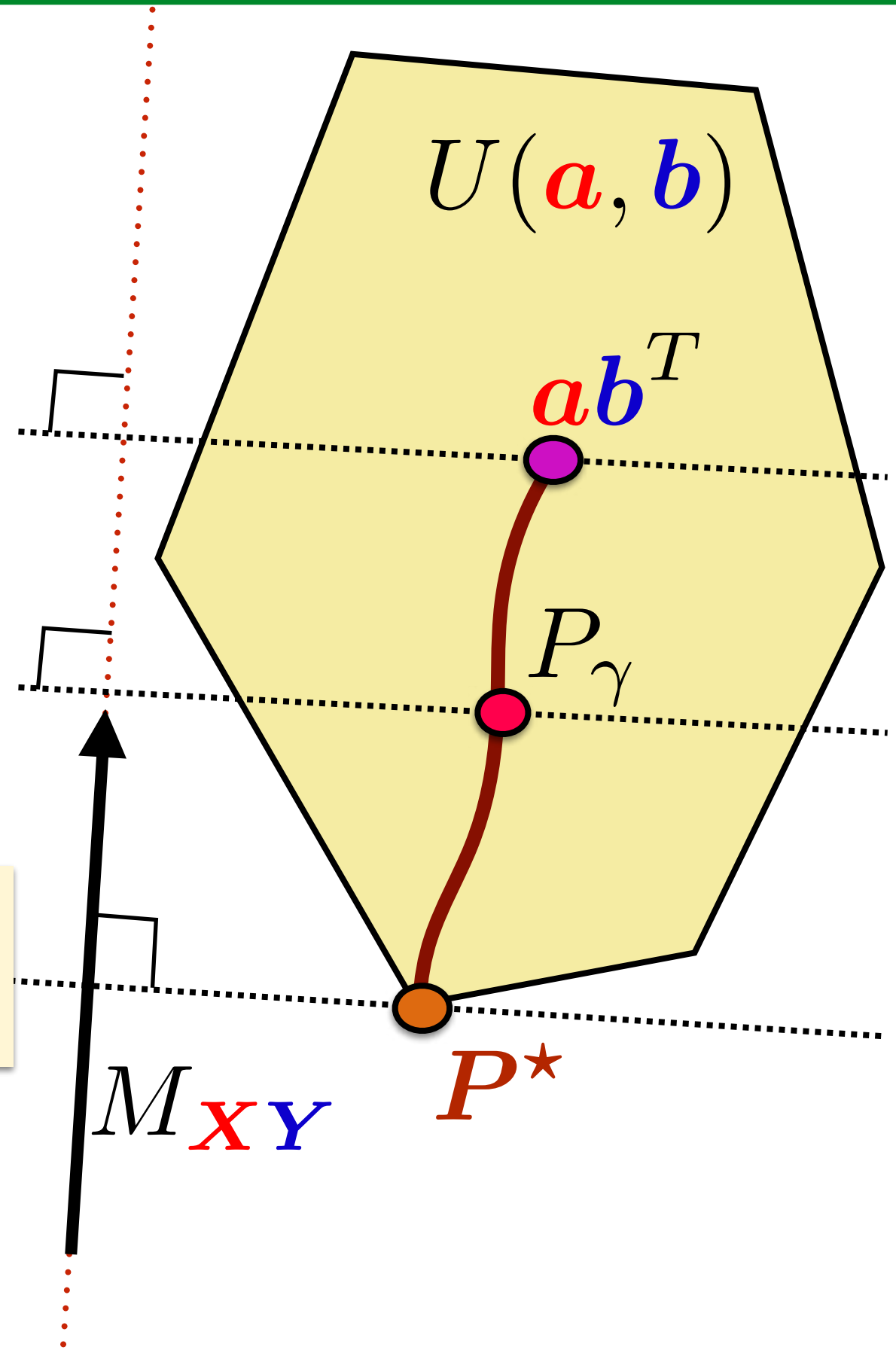
Sinkhorn, W and *Energy Distance*

$$\mu = \sum_{i=1}^n a_i \delta_{x_i} \quad \nu = \sum_{j=1}^m b_j \delta_{y_j}$$

$$\mathcal{E}(\mu, \nu) = \langle ab^T, M_{\mathbf{x}\mathbf{y}} \rangle$$

$$W_\gamma(\mu, \nu) = \langle P_\gamma, M_{\mathbf{x}\mathbf{y}} \rangle$$

$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{x}\mathbf{y}} \rangle$$



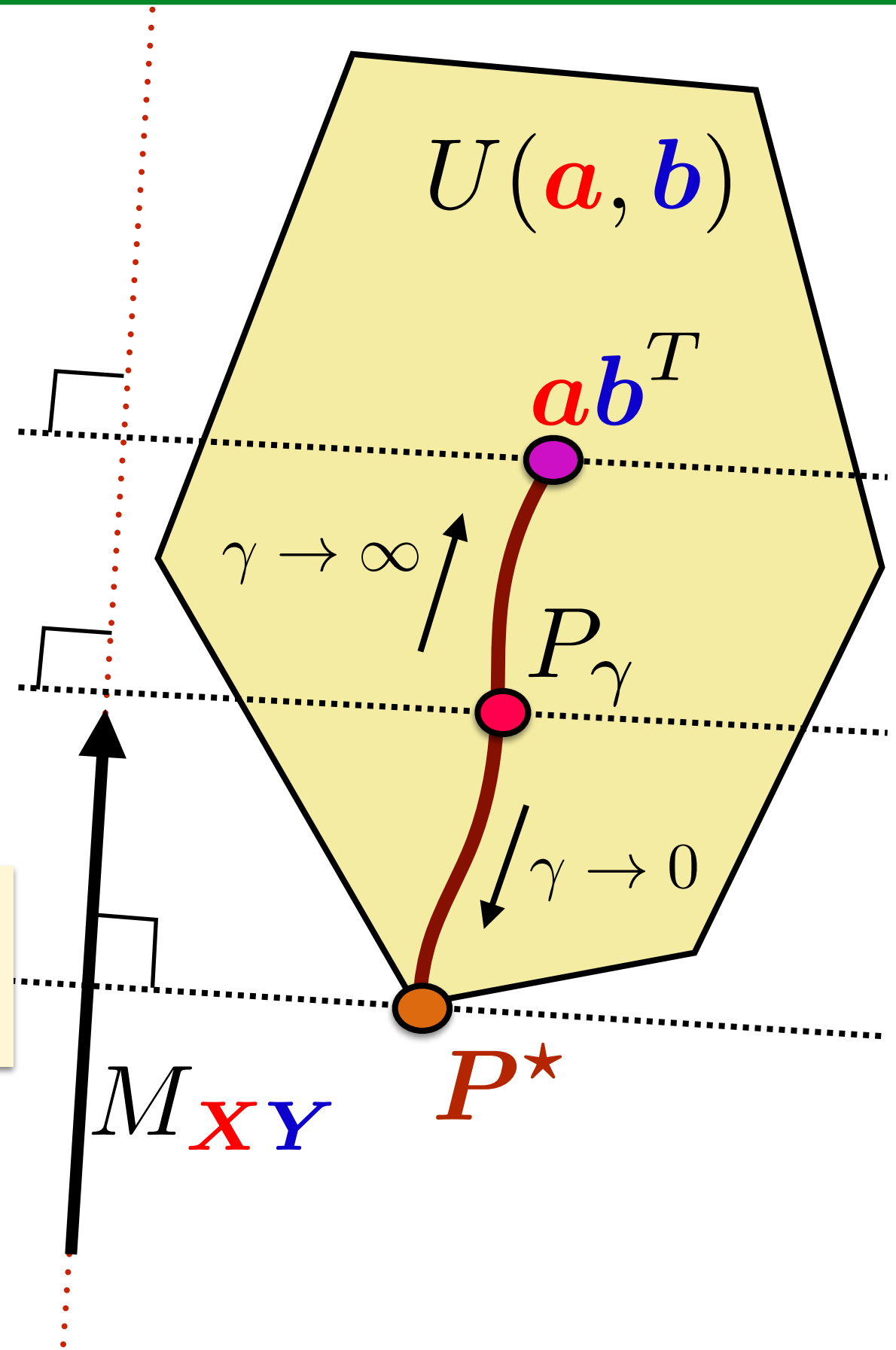
Sinkhorn, W and *Energy Distance*

$$\mu = \sum_{i=1}^n a_i \delta_{x_i} \quad \nu = \sum_{j=1}^m b_j \delta_{y_j}$$

$$\mathcal{E}(\mu, \nu) = \langle ab^T, M_{\mathbf{XY}} \rangle$$

$$W_\gamma(\mu, \nu) = \langle P_\gamma, M_{\mathbf{XY}} \rangle$$

$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{XY}} \rangle$$



Sinkhorn, W and *Energy Distance*

$$\mathcal{E}(\mu, \nu) = \langle ab^T, M_{\mathbf{X}\mathbf{Y}} \rangle$$

$$\text{ED}(\mu, \nu) = \mathcal{E}(\mu, \nu) - \frac{1}{2}(\mathcal{E}(\mu, \mu) + \mathcal{E}(\nu, \nu))$$

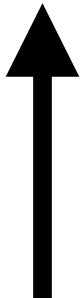
$$W_\gamma(\mu, \nu) = \langle P_\gamma, M_{\mathbf{X}\mathbf{Y}} \rangle - \gamma E(P_\gamma)$$

$$\bar{W}_\gamma(\mu, \nu) = W_\gamma(\mu, \nu) - \frac{1}{2}(W_\gamma(\mu, \mu) + W_\gamma(\nu, \nu))$$

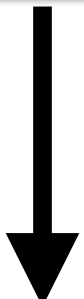
$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{X}\mathbf{Y}} \rangle$$

Sinkhorn, W and *Energy Distance*

$$\mathcal{MMD}(\mu, \nu) = \mathcal{E}(\mu, \nu) - \frac{1}{2}(\mathcal{E}(\mu, \mu) + \mathcal{E}(\nu, \nu))$$

$\gamma \rightarrow \infty$ 

$$\bar{W}_\gamma(\mu, \nu) = W_\gamma(\mu, \nu) - \frac{1}{2}(W_\gamma(\mu, \mu) + W_\gamma(\nu, \nu))$$

$\gamma \rightarrow 0$ 

$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{x}\mathbf{y}} \rangle$$

How to compare them?

i.i.d samples $\mathbf{x}_1, \dots, \mathbf{x}_n \sim \mu$, $\mathbf{y}_1, \dots, \mathbf{y}_m \sim \nu$,

$$\hat{\mu}_n \stackrel{\text{def}}{=} \frac{1}{n} \sum_i \delta_{\mathbf{x}_i}, \hat{\nu}_m \stackrel{\text{def}}{=} \frac{1}{m} \sum_j \delta_{\mathbf{y}_j}$$

Computational properties

Effort to compute/approximate $\Delta(\hat{\mu}_n, \hat{\nu}_m)$?

Statistical properties

$$|\Delta(\mu, \nu) - \Delta(\hat{\mu}_n, \hat{\nu}_n)| \leq f(n)?$$

Sinkhorn in between W and MMD

$$MMD(\mu, \nu) = \mathcal{E}(\mu, \nu) - \frac{1}{2}(\mathcal{E}(\mu, \mu) + \mathcal{E}(\nu, \nu))$$

$(n + m)^2$ $O(1/\sqrt{n})$ [see Arthur]

$$W^p(\mu, \nu) = \langle P^*, M_{XY} \rangle$$

$$O((n + m)nm \log(n + m))$$

$$O(1/n^{1/d})$$

Sinkhorn in between W and MMD

$$\mathcal{MMD}(\mu, \nu) = \mathcal{E}(\mu, \nu) - \frac{1}{2}(\mathcal{E}(\mu, \mu) + \mathcal{E}(\nu, \nu))$$

$$(n + m)^2$$

$$O(1/\sqrt{n})$$

[see Arthur]

$$\bar{W}_\gamma(\mu, \nu) = W_\gamma(\mu, \nu) - \frac{1}{2}(W_\gamma(\mu, \mu) + W_\gamma(\nu, \nu))$$

$$O((n + m)^2)$$

$$O\left(\frac{1}{\gamma^{d/2}\sqrt{n}}\right)$$

[GCBBCP'18]

[FSVATP'18]

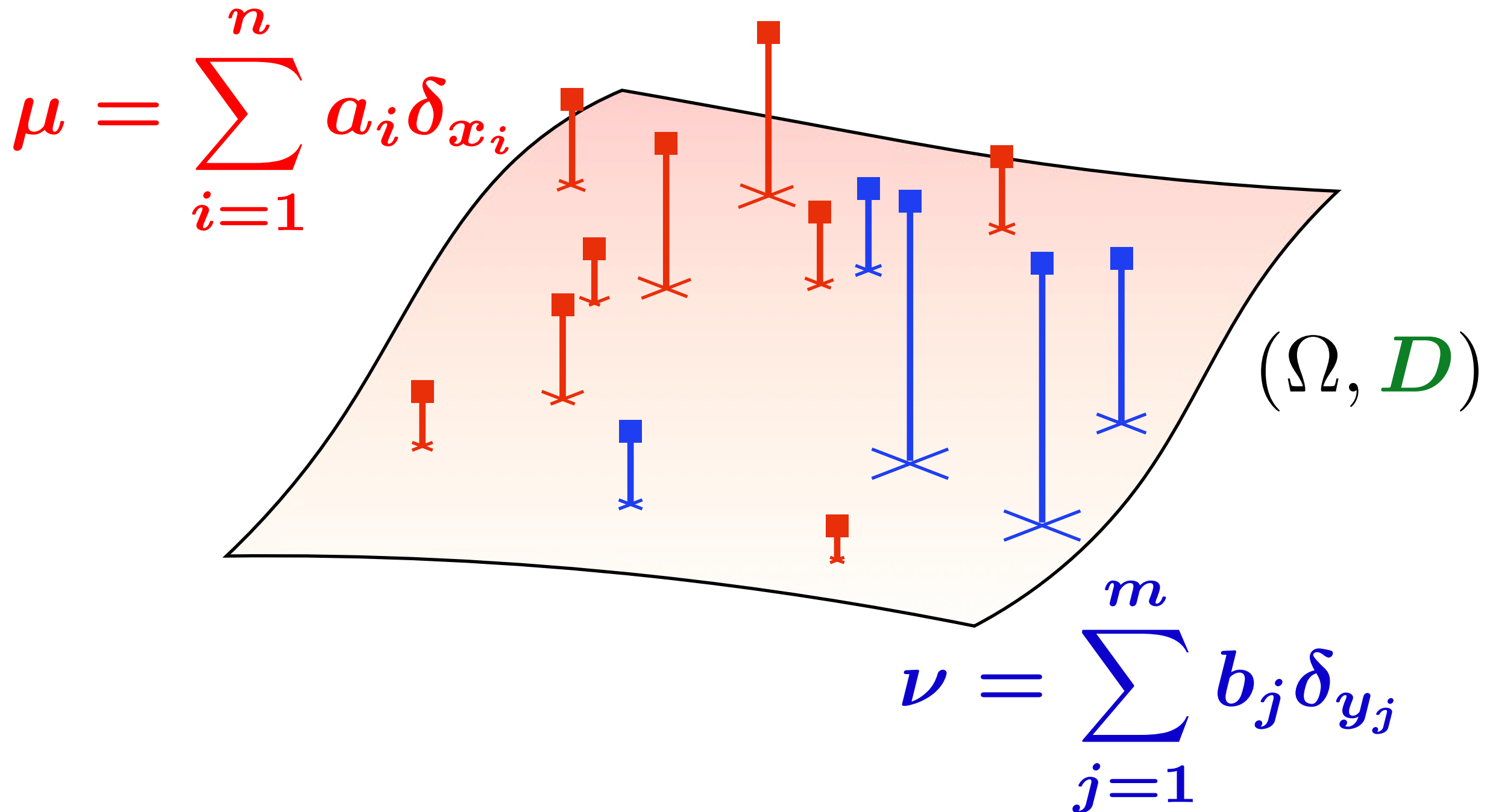
$$W^p(\mu, \nu) = \langle P^\star, M_{\mathbf{X}\mathbf{Y}} \rangle$$

$$O((n + m)nm \log(n + m))$$

$$O(1/n^{1/d})$$

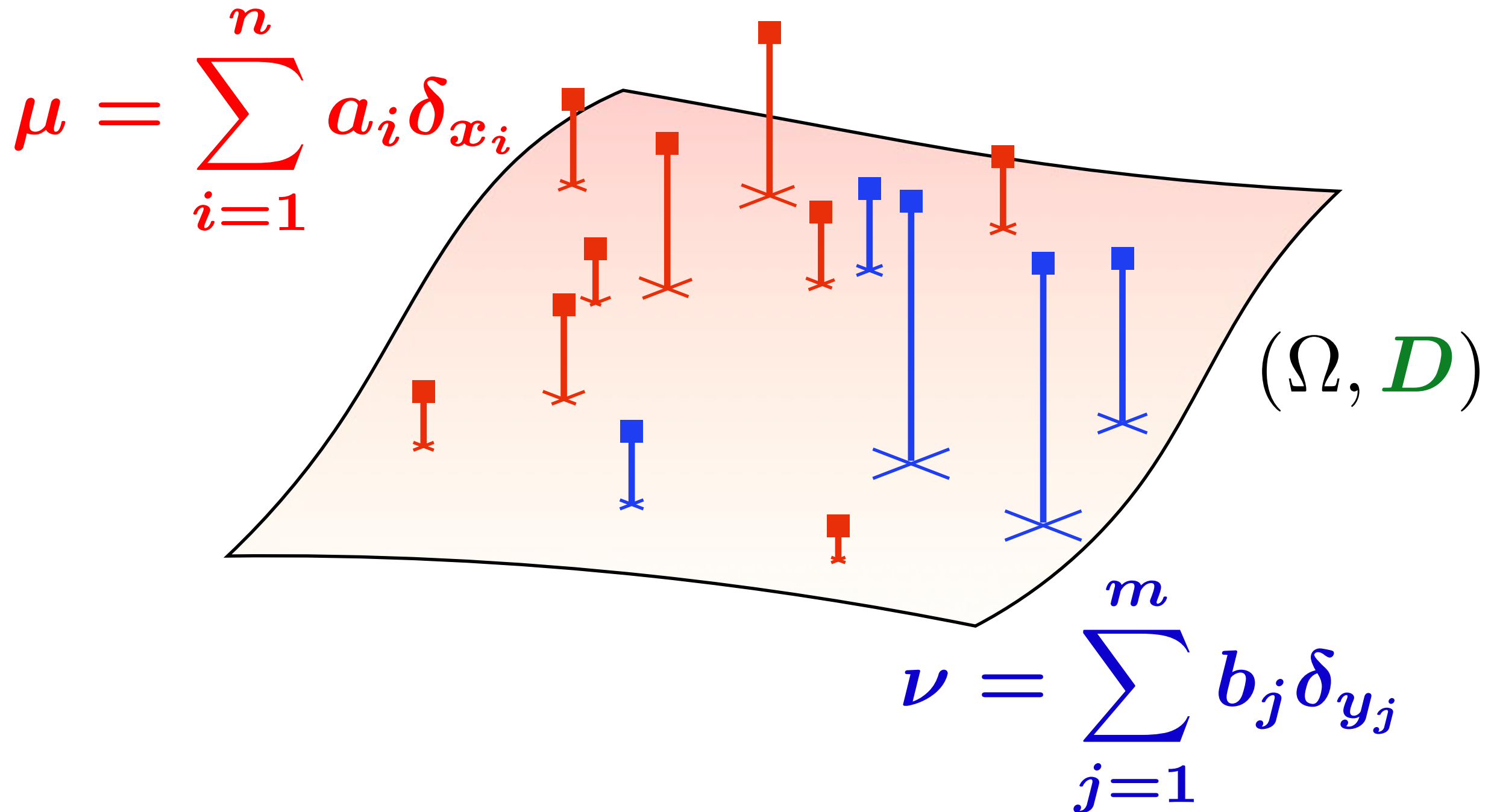
Differentiability of W

$$W((a, X), (b, Y))$$



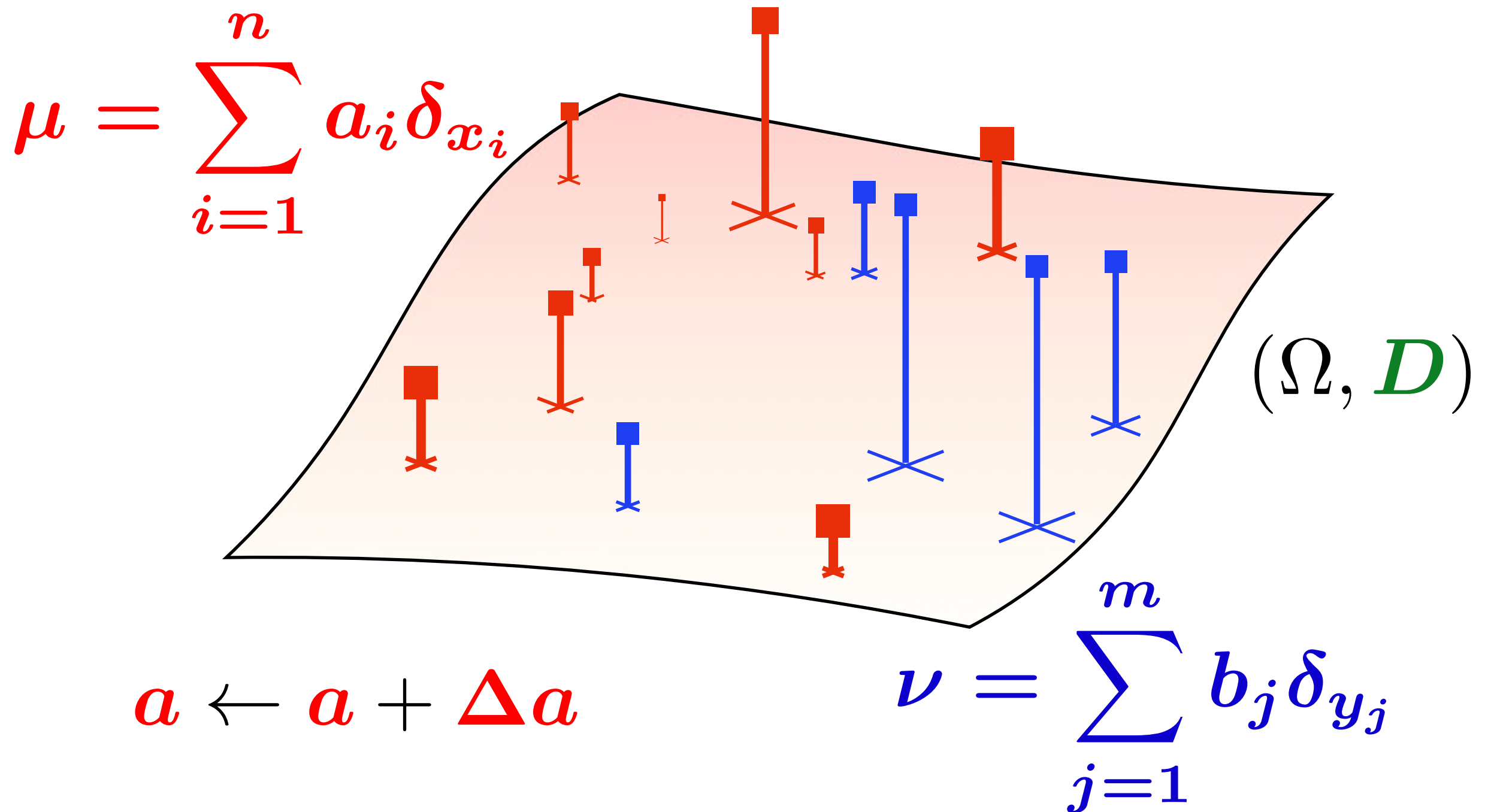
Differentiability of W

$$W((a + \Delta a, X), (b, Y)) = W((a, X), (b, Y)) + ??$$



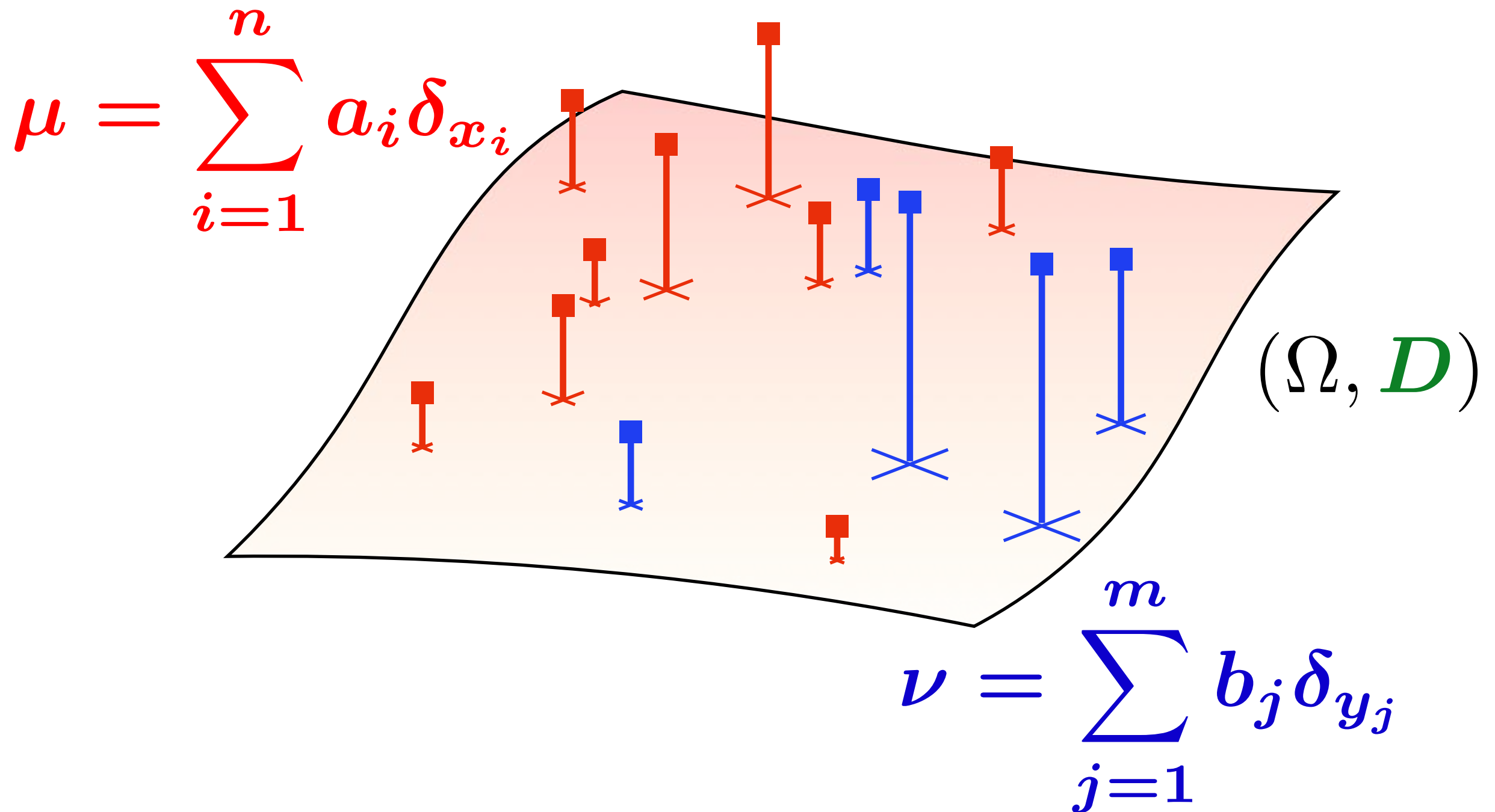
Differentiability of W

$$W((a + \Delta a, X), (b, Y)) = W((a, X), (b, Y)) + ??$$



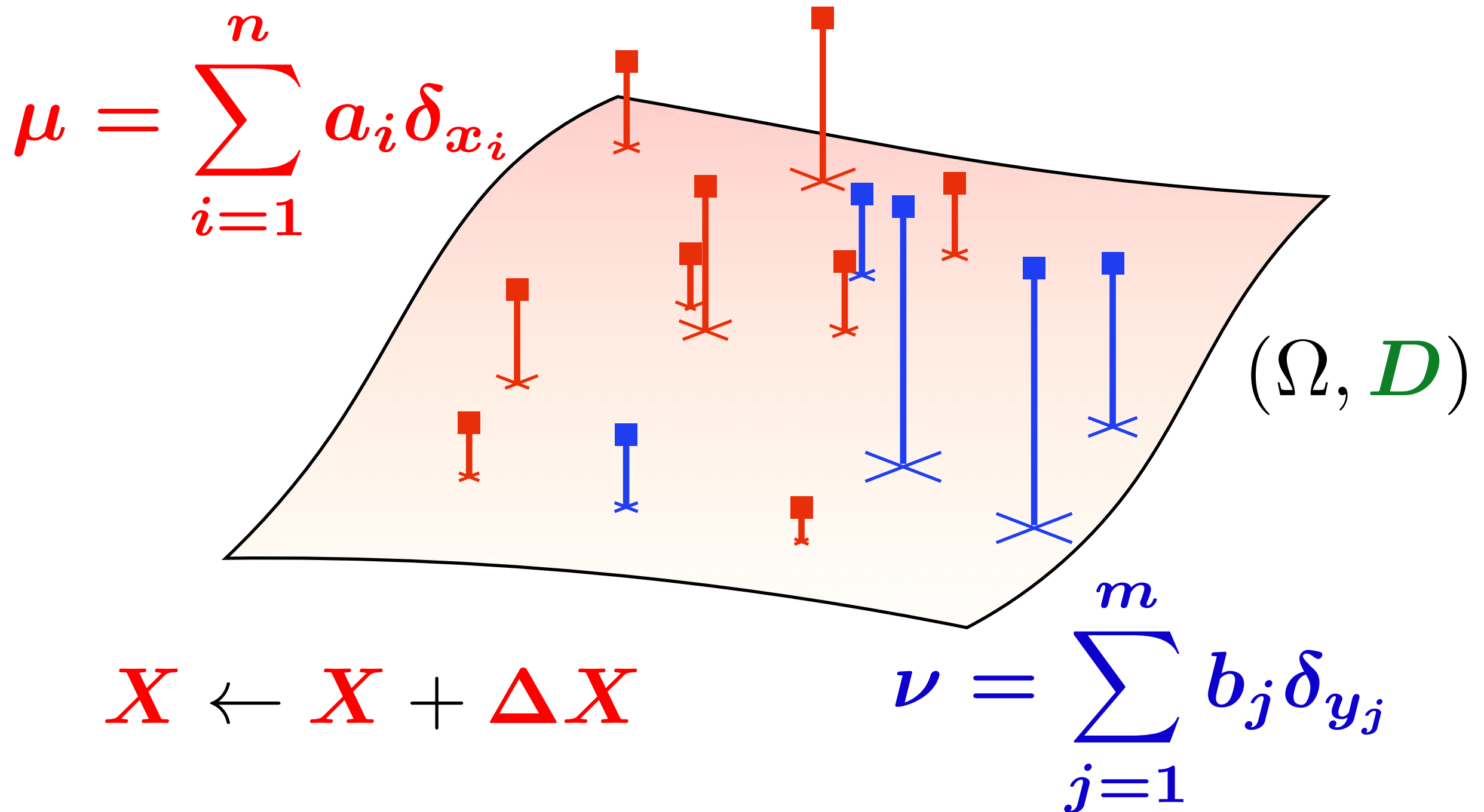
Sinkhorn $\rightsquigarrow$ *Differentiability*

$$W((a, X + \Delta X), (b, Y)) = W((a, X), (b, Y)) + ??$$



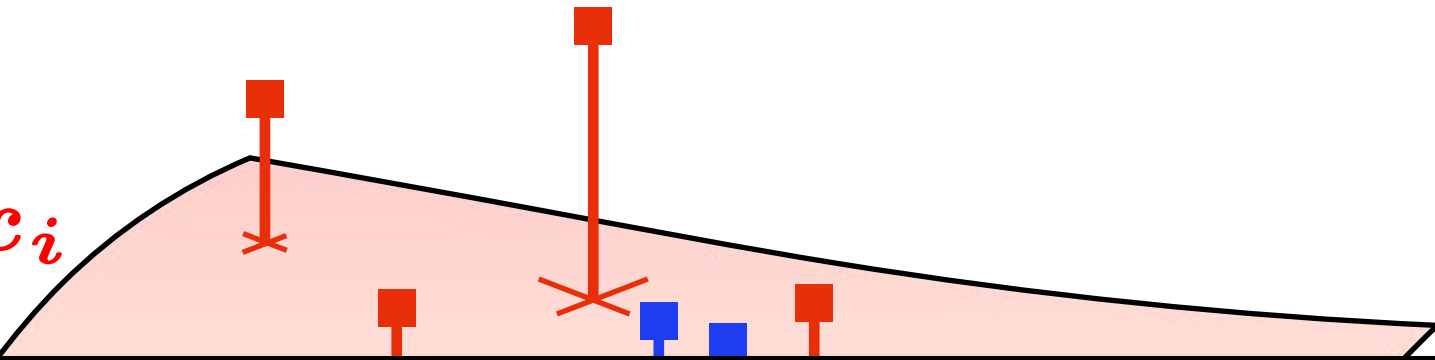
Sinkhorn $\rightsquigarrow$ *Differentiability*

$$W((a, X + \Delta X), (b, Y)) = W((a, X), (b, Y)) + ??$$



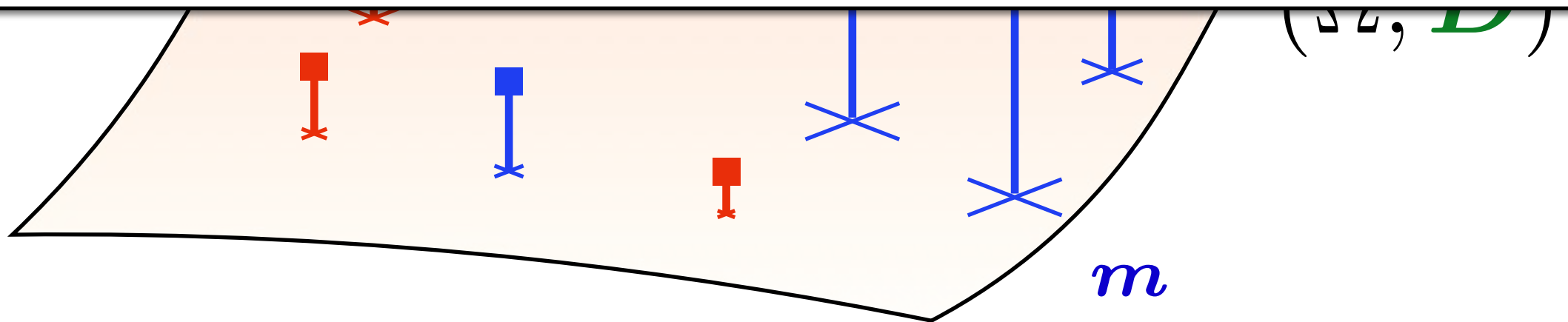
Sinkhorn $\rightsquigarrow$ *Differentiability*

$$W((a, X + \Delta X), (b, Y)) = W((a, X), (b, Y)) + ??$$

$$\mu = \sum_{i=1}^n a_i \delta_{x_i}$$


A diagram illustrating the measure μ as a distribution over a set of points x_i . It shows a red shaded area representing the support of the measure, with several red squares representing the points x_i . Vertical red lines with 'X' marks at the top indicate the weights a_i assigned to each point.

Changes in objective can be handled with Danskin's theorem.



$$X \leftarrow X + \Delta X$$

$$\nu = \sum_{j=1}^m b_j \delta_{y_j}$$

How to decrease W ? change weights

$$W_p^p(\mu, \nu) = \max_{\substack{\alpha \in \mathbb{R}^n, \beta \in \mathbb{R}^m \\ \alpha \oplus \beta \leq M_{\mathbf{x}\mathbf{y}}}} \alpha^T a + \beta^T b.$$

DUAL

Prop. $W(\mu, \nu)$ is convex w.r.t. a ,

$$\partial_a W = \arg_{\alpha} \max_{\alpha \oplus \beta \leq M_{\mathbf{x}\mathbf{y}}} \alpha^T a + \beta^T b.$$

Prop. $W_\gamma(\mu, \nu)$ is convex and differentiable w.r.t. a , $\nabla_a W_\gamma = \alpha_\gamma^\star = \gamma \log u$

How to decrease W ? change locations

$$W_2^2(\mu, \nu) = \min_{\substack{P \in \mathbb{R}_+^{n \times m} \\ P \mathbf{1}_m = \mathbf{a}, P^T \mathbf{1}_n = \mathbf{b}}} \langle P, \mathbf{1}_n \mathbf{1}_d^T X^2 + Y^{2T} \mathbf{1}_d \mathbf{1}_m - 2X^T Y \rangle$$

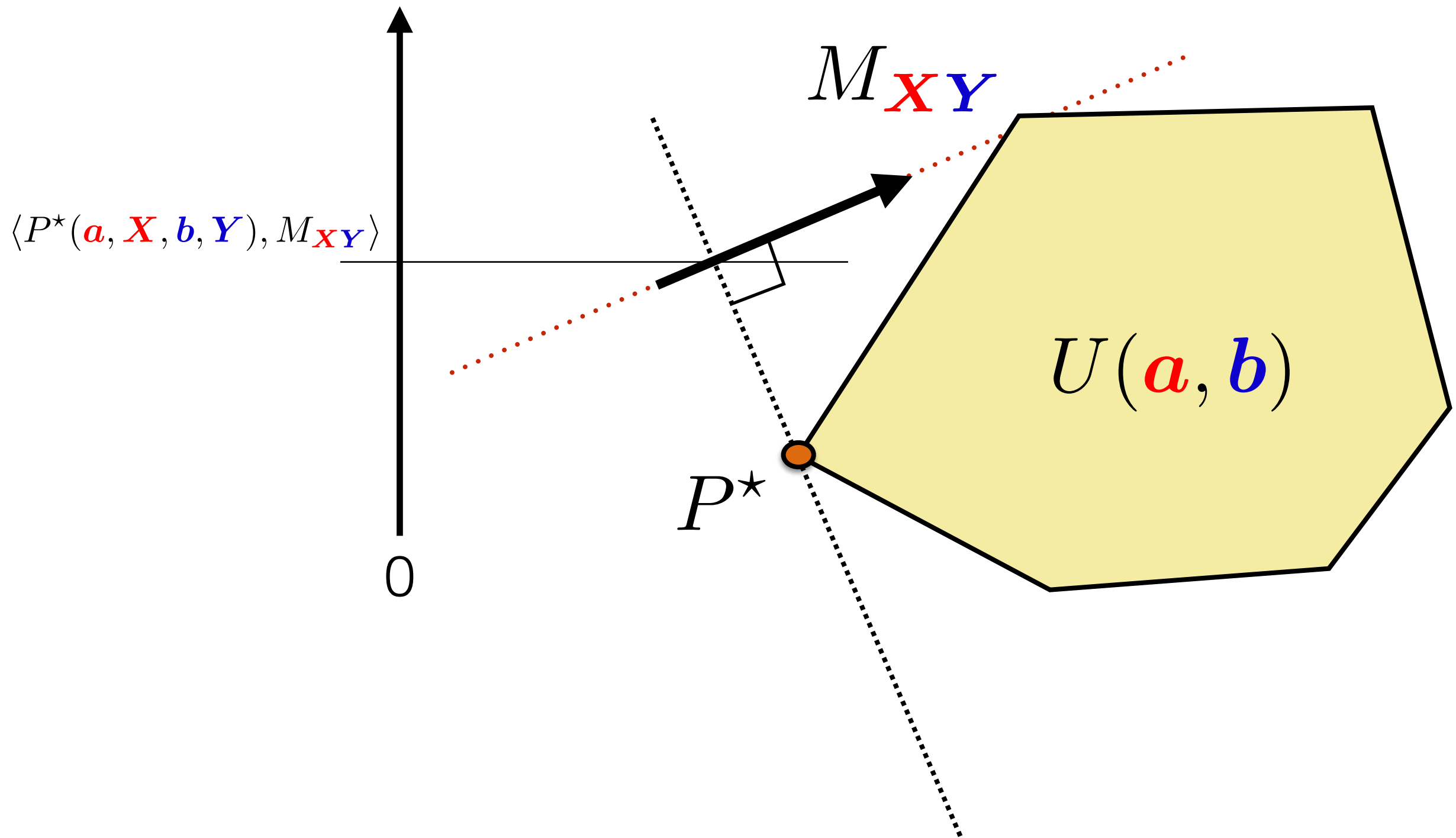
PRIMAL

Prop. $p = 2, \Omega = \mathbb{R}^d$. $W(\mu, \nu)$ decreases if
$$X \leftarrow Y P^{\star T} \mathbf{D}(\mathbf{a}^{-1}).$$

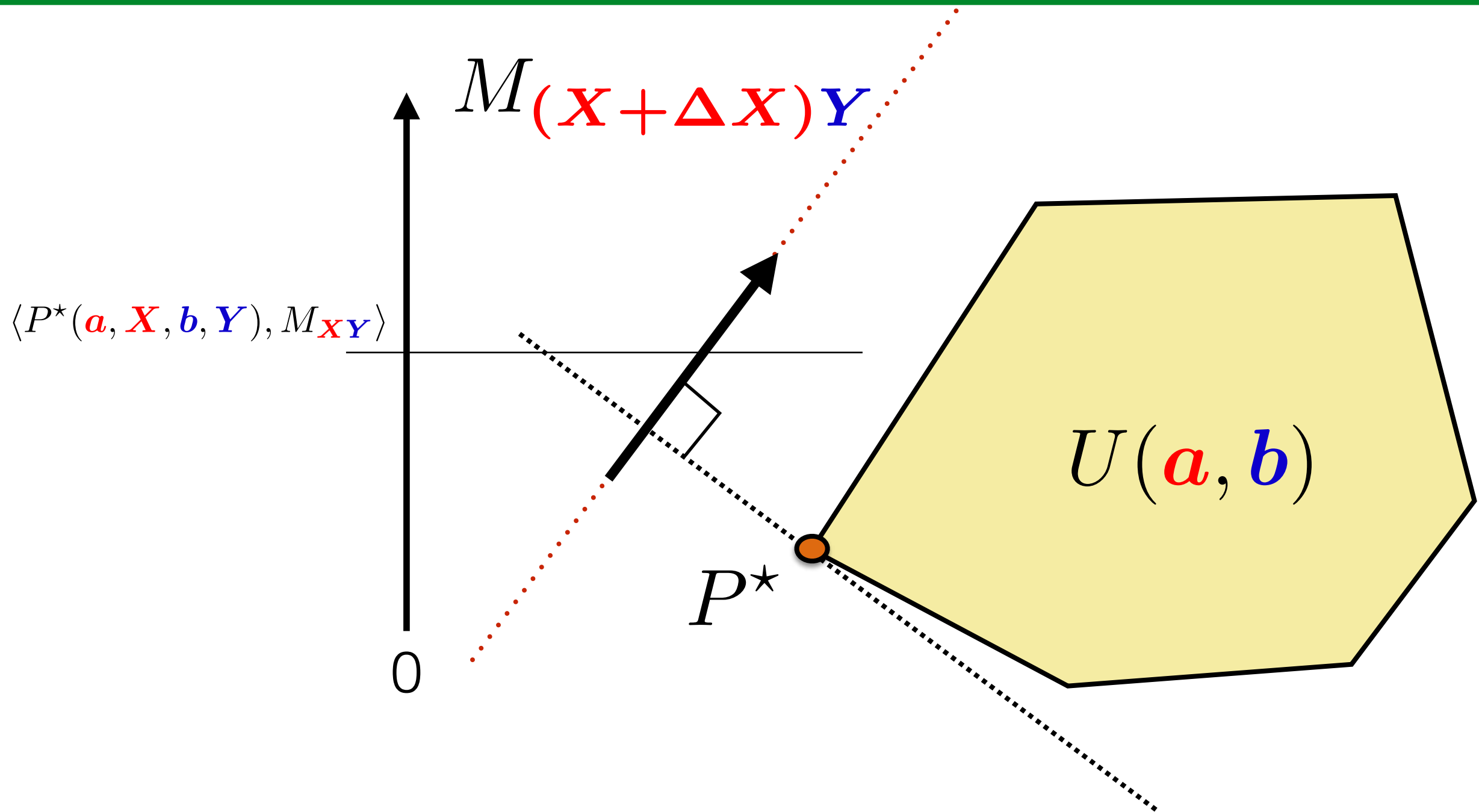
Prop. $p = 2, \Omega = \mathbb{R}^d$. $W_\gamma(\mu, \nu)$ is differentiable w.r.t. X , with

$$\nabla_X W_\gamma = X - Y P_\gamma^T \mathbf{D}(\mathbf{a}^{-1}).$$

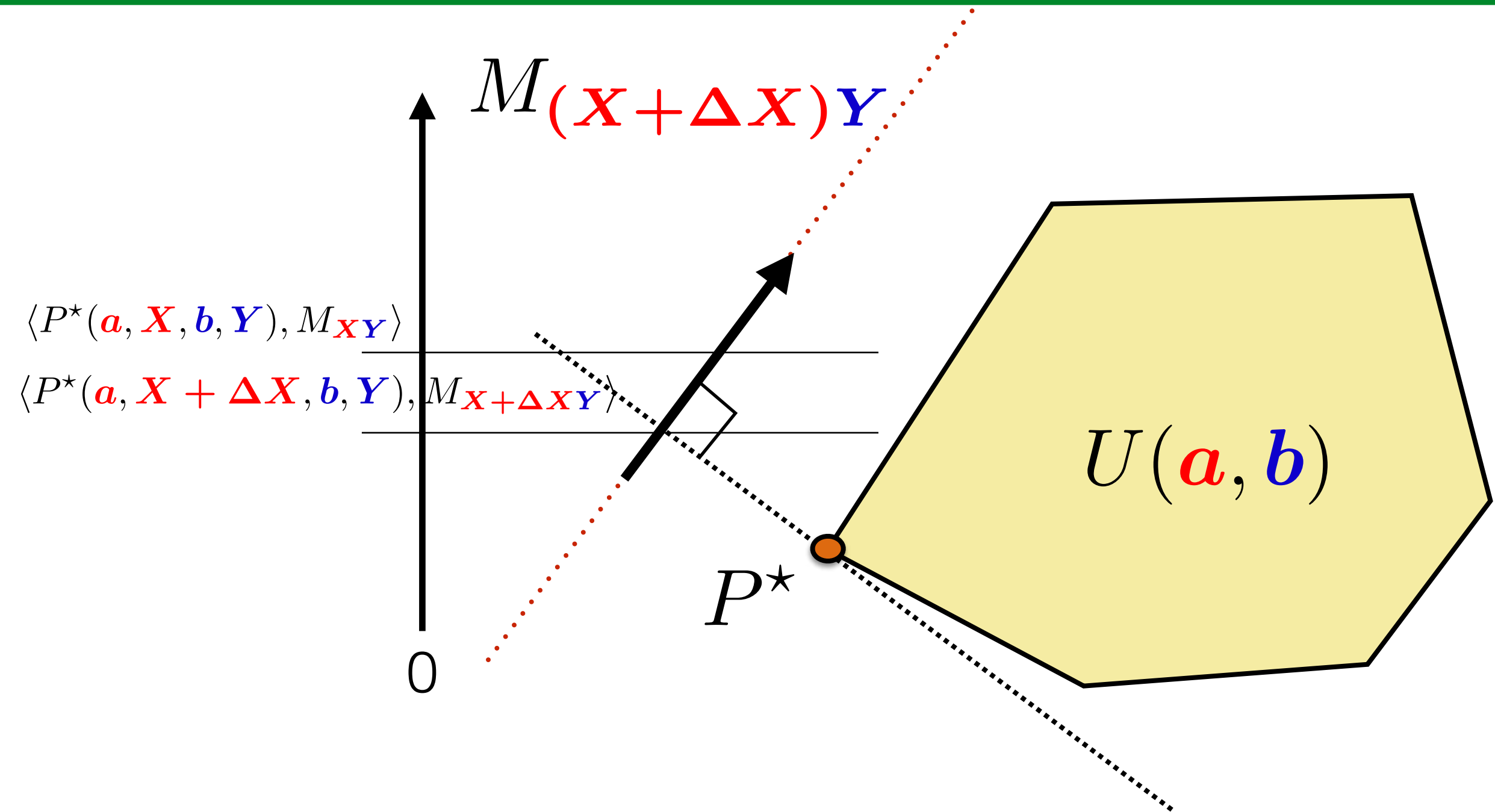
Solving the OT Problem



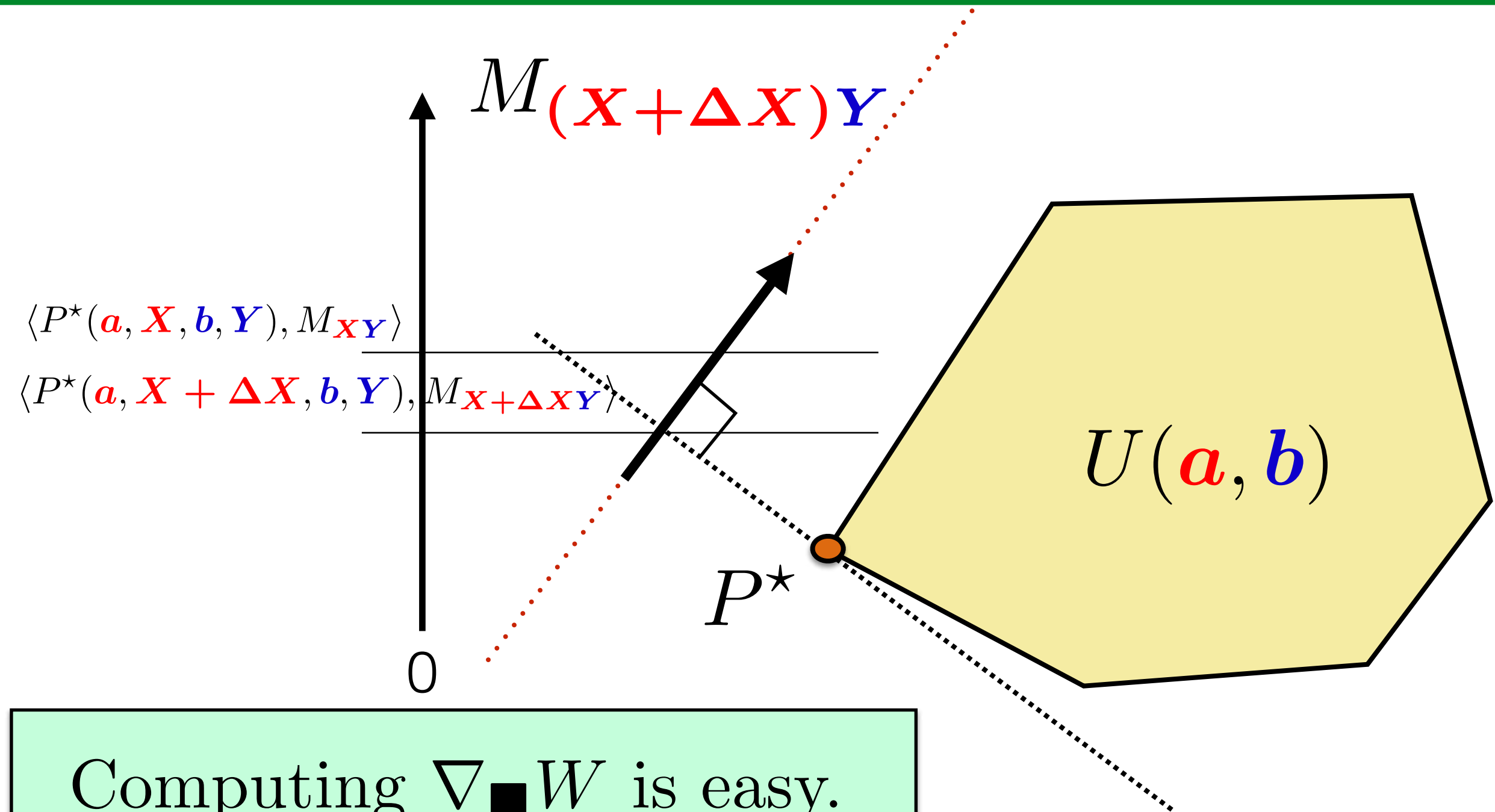
Solving the OT Problem



Solving the OT Problem



Solving the OT Problem

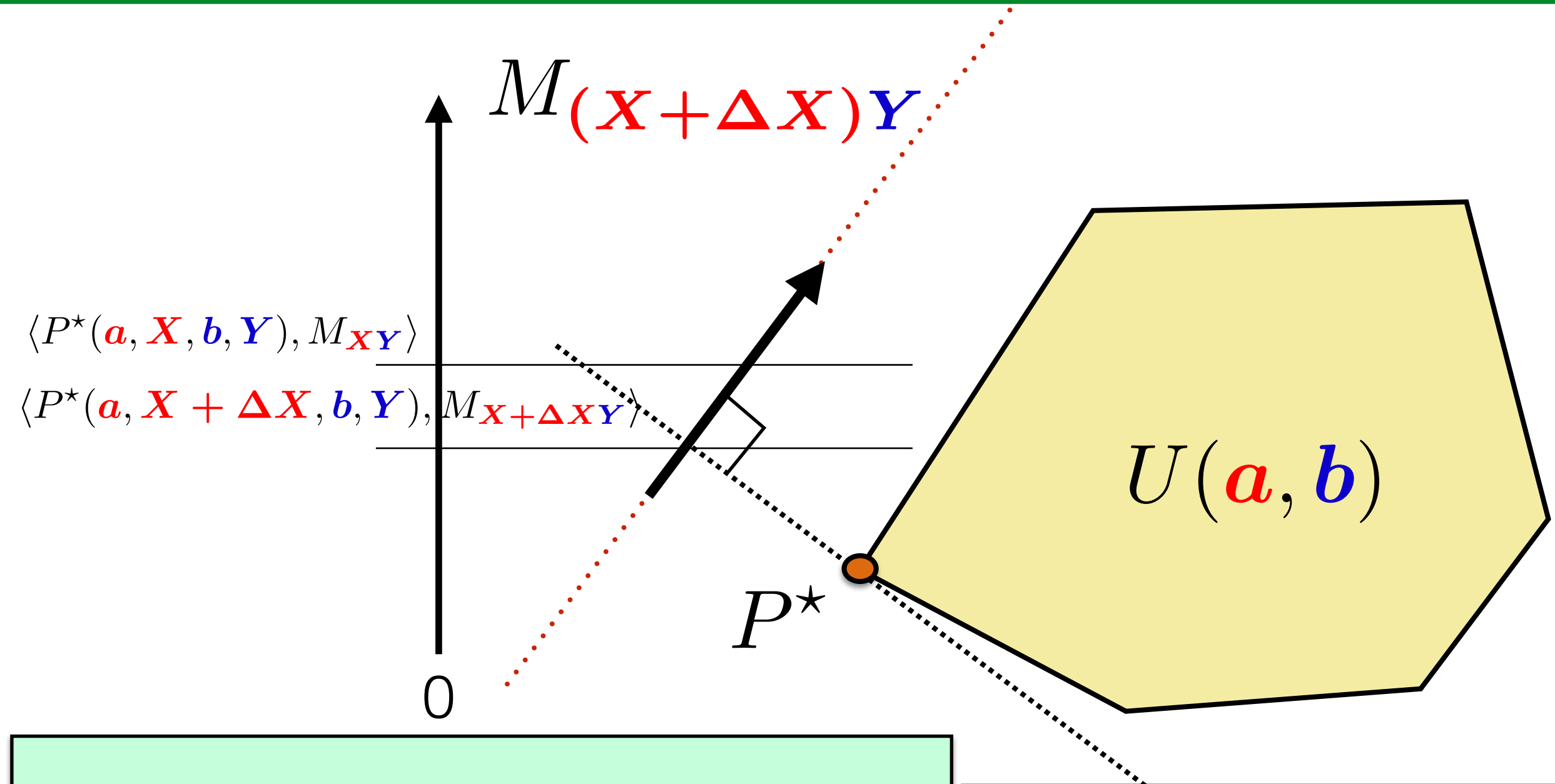


Computing $\nabla_{\blacksquare} W$ is easy.

here $\blacksquare$ can be anything,

e.g. a or X , parameter θ for cost c_θ

Solving the OT Problem



Computing $\nabla_{\blacksquare} W$ is easy.

here $\blacksquare$ can be anything,

e.g. $\mathbf{a}$ or $\mathbf{X}$, parameter θ for cost $\mathbf{c}_\theta$

Computing $\frac{\partial P^\star}{\partial \blacksquare}$ is harder. 



Sinkhorn: A Programmer View

Def. For $L \geq 1$, define

$$\mathbf{P}_L \stackrel{\text{def}}{=} \text{diag}(\mathbf{u}_L) \mathbf{K} \text{diag}(\mathbf{v}_L),$$

where

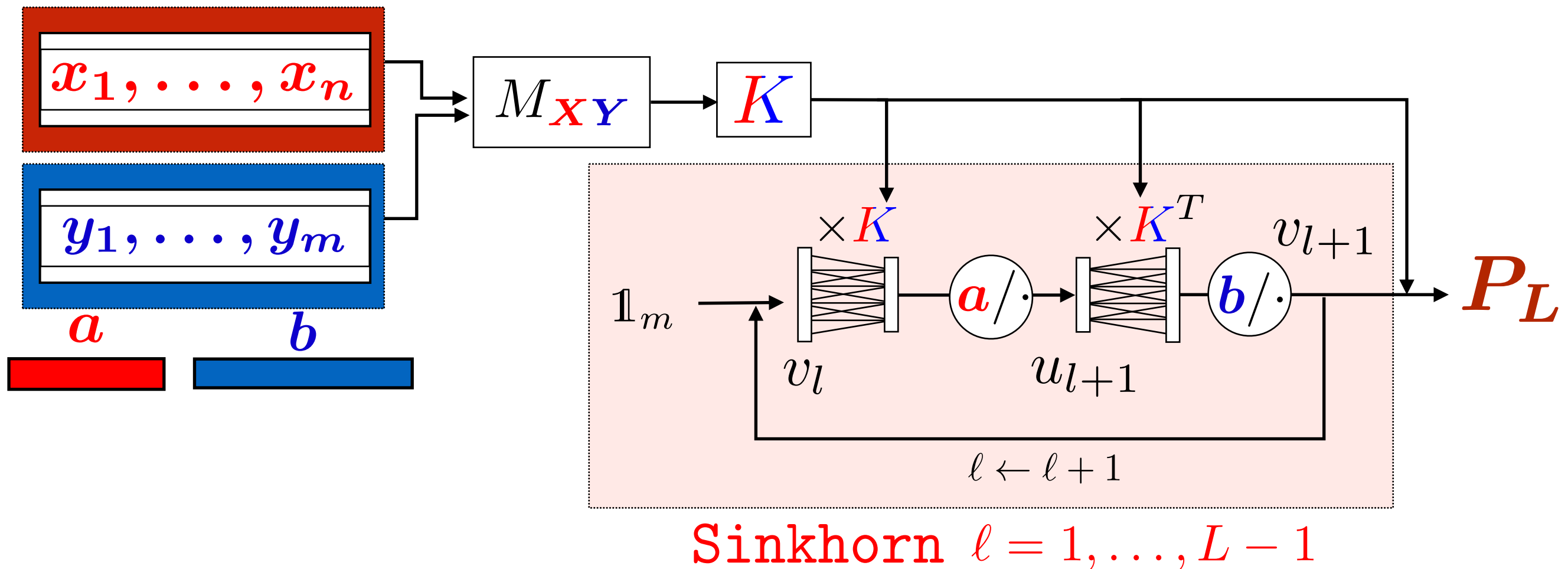
$$\mathbf{v}_0 = \mathbf{1}_m; l \geq 0, \mathbf{u}_l \stackrel{\text{def}}{=} \mathbf{a} / \mathbf{K} \mathbf{v}_l, \mathbf{v}_{l+1} \stackrel{\text{def}}{=} \mathbf{b} / \mathbf{K}^T \mathbf{u}_l.$$

Prop. $\frac{\partial \mathbf{P}_L}{\partial \mathbf{X}}, \frac{\partial \mathbf{P}_L}{\partial \mathbf{a}}$ can be computed recursively, in $O(L)$ kernel $\mathbf{K} \times$ vector products.

Sinkhorn: A Programmer View

Def. For $L \geq 1$, define

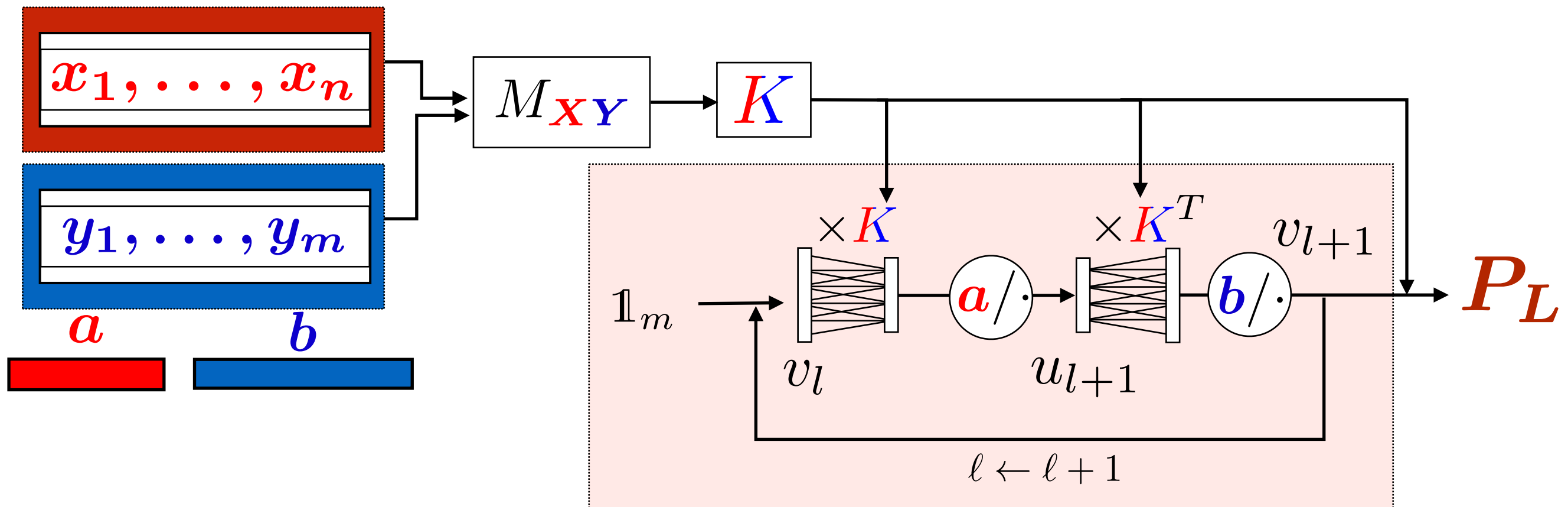
$$W_L(\mu, \nu) \stackrel{\text{def}}{=} \langle P_L, M_{\mathbf{x}\mathbf{y}} \rangle,$$



Sinkhorn: A Programmer View

Def. For $L \geq 1$, define

$$W_L(\mu, \nu) \stackrel{\text{def}}{=} \langle P_L, M_{\mathbf{x}\mathbf{y}} \rangle,$$



Sinkhorn $\ell = 1, \dots, L-1$

[Adams'11] [Hashimoto'16] [Bonneel'16][Shalit'16]

Sinkhorn: A Mathematician View

$$F : \mu, \nu, c, \varepsilon \rightarrow (\alpha^*, \beta^*)$$

$$H(\mu, \nu, c, \varepsilon; \alpha, \beta) := \begin{bmatrix} e^{\frac{\alpha \oplus \beta - c}{\varepsilon}} \mathbf{1}_m - a \\ e^{\frac{\beta \oplus \alpha - c^T}{\varepsilon}} \mathbf{1}_n - b \end{bmatrix} = 0$$

At the optimum,

$$H(\mu, \nu, c, \varepsilon; F(\mu, \nu, c, \varepsilon)) = 0$$

Sinkhorn: A Mathematician View

$$F : \mu, \nu, c, \varepsilon \rightarrow (\alpha^*, \beta^*)$$

$$H(\mu, \nu, c, \varepsilon; \alpha, \beta) := \begin{bmatrix} e^{\frac{\alpha \oplus \beta - c}{\varepsilon}} \mathbf{1}_m = a \\ e^{\frac{\beta \oplus \alpha - c^T}{\varepsilon}} \mathbf{1}_n = b \end{bmatrix} = 0$$

Using the implicit function theorem

$$J_{F, \blacksquare} = - \left(J_{H, (\alpha, \beta)}(\blacksquare, (\alpha^*, \beta^*)) \right)^{-1} J_{H, \blacksquare}(\blacksquare, (\alpha^*, \beta^*))$$

$$J_{F, \blacksquare}^T = - J_{H, \blacksquare}(\blacksquare, (\alpha^*, \beta^*))^T \left(J_{H, (\alpha, \beta)}(\blacksquare, (\alpha^*, \beta^*)) \right)^{-T}$$

The best of both worlds

When computing the gradient of a loss whose evaluation depends on the optimal transport solution, backward mode differentiation is more efficient. This involves evaluating:

$$J_{F, \blacksquare}^T \mathbf{z}$$

$$-J_{H, \blacksquare}(\blacksquare, (\alpha^*, \beta^*))^T \left(\left(J_{H, (\alpha, \beta)}(\blacksquare, (\alpha^*, \beta^*)) \right)^T \right)^{-1} \mathbf{z}$$

The best of both worlds

When computing the gradient of a loss whose evaluation depends on the optimal transport solution, backward mode differentiation is more efficient. This involves evaluating:

$$J_{F, \blacksquare}^T \mathbf{z}$$

$$-J_{H, \blacksquare}(\blacksquare, (\alpha^*, \beta^*))^T \left(\underbrace{\left(J_{H, (\alpha, \beta)}(\blacksquare, (\alpha^*, \beta^*)) \right)^T}_{\text{jax.vjp}} \right)^{-1} \mathbf{z}$$

The best of both worlds

When computing the gradient of a loss whose evaluation depends on the optimal transport solution, backward mode differentiation is more efficient. This involves evaluating:

$$J_{F, \blacksquare}^T \mathbf{z}$$

$$-J_{H, \blacksquare}(\blacksquare, (\alpha^*, \beta^*))^T \left(\underbrace{\left(J_{H, (\alpha, \beta)}(\blacksquare, (\alpha^*, \beta^*)) \right)^T}_{\text{jax.vjp}} \right)^{-1} \mathbf{z}$$

`jax.linalg.sparse.cg`

The best of both worlds

When computing the gradient of a loss whose evaluation depends on the optimal transport solution, backward mode differentiation is more efficient. This involves evaluating:

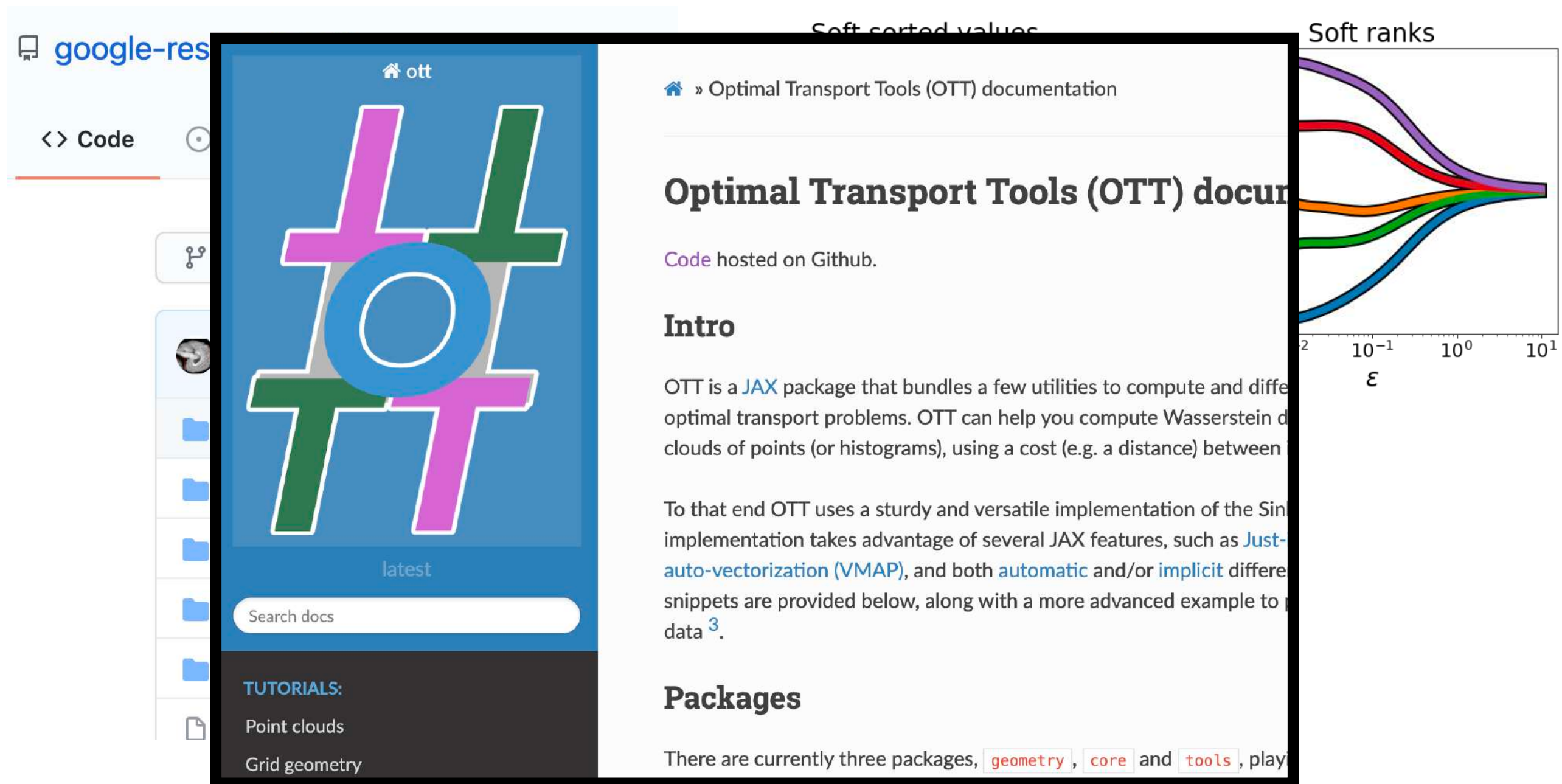
$$J_{F, \blacksquare}^T \mathbf{z}$$

$$- \underset{\text{jax.vjp}}{J_{H, \blacksquare}(\blacksquare, (\alpha^*, \beta^*))^T} \left(\underset{\text{jax.vjp}}{\left(J_{H, (\alpha, \beta)}(\blacksquare, (\alpha^*, \beta^*)) \right)^T} \right)^{-1} \mathbf{z}$$

`jax.linalg.sparse.cg`

For more details

<https://github.com/google-research/ott>



google-res

<> Code

ott

Optimal Transport Tools (OTT) documentation

Code hosted on Github.

Intro

OTT is a [JAX](#) package that bundles a few utilities to compute and differentiate optimal transport problems. OTT can help you compute Wasserstein distances between clouds of points (or histograms), using a cost (e.g. a distance) between

To that end OTT uses a sturdy and versatile implementation of the Sinkhorn algorithm. This implementation takes advantage of several JAX features, such as [Just-in-time auto-vectorization \(VMAP\)](#), and both [automatic](#) and/or [implicit differentiation](#). Code snippets are provided below, along with a more advanced example to compute data ³.

Packages

There are currently three packages, `geometry`, `core` and `tools`, play

Soft sorted values

Soft ranks

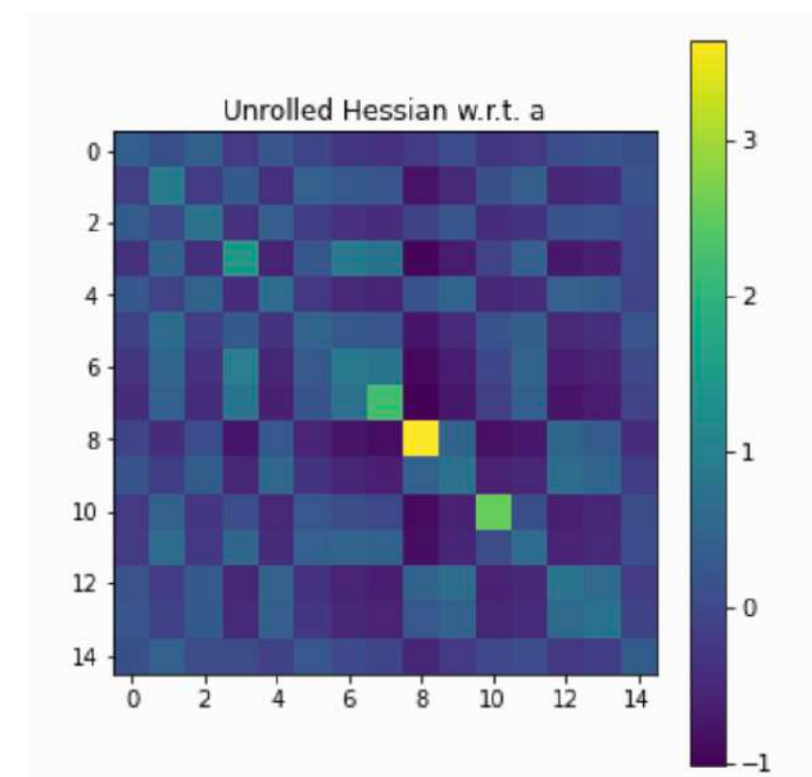
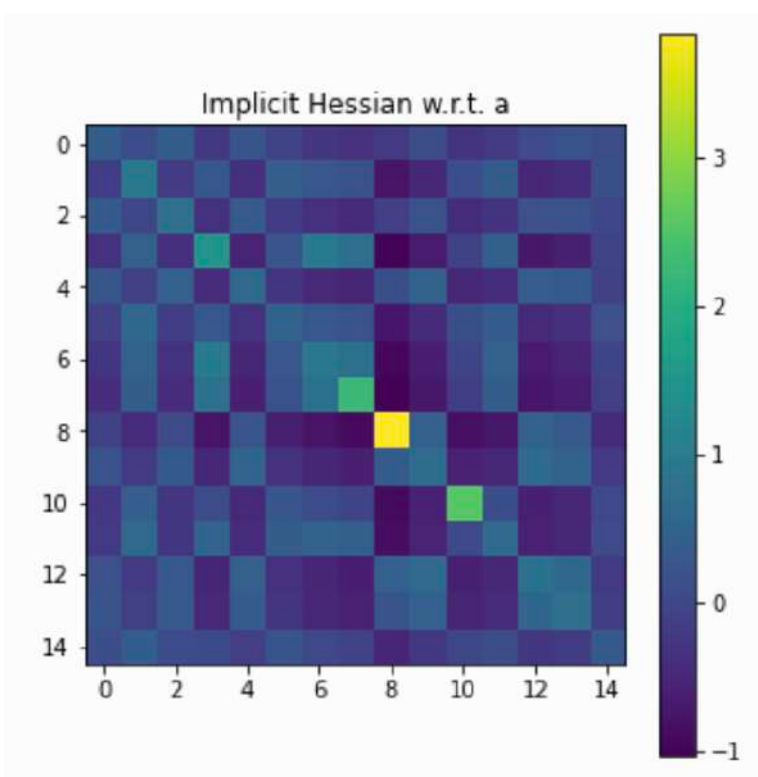
ϵ

For instance...computing Hessians.

<https://ott-jax.readthedocs.io/en/stable/notebooks/Hessians.html>

```
[9]: def loss(a, x, implicit):  
      return sinkhorn_divergence.sinkhorn_divergence(  
          pointcloud.PointCloud, x, y, # this part defines geometry  
          a=a, b=b, # this sets weights  
          sinkhorn_kwargs={'implicit_differentiation': implicit, 'use_danskin': False} # to be  
          ).divergence
```

```
[17]: for arg in [0,1]:  
      # Compute Hessians using either unrolling or implicit differentiation.  
      hess_loss_imp = jax.jit(jax.hessian(lambda a, x: loss(a, x, True),  
                                      argnums=arg))
```



To conclude

To sum up

- **Optimal matchings** used everywhere in ML, to
 - Compute a distance (W) between measures.
 - register/map/match points, to disambiguate.
- **Regularization** is needed to:
 - Scale up/robustify these tools;
 - Ensure gradients/jacobians exist and are not 0.
- **To differentiate** through regularised OT, either
 - Automatic differentiation (unrolling)
 - Implicit differentiation

A few research topics around OT

- **Sinkhorn fixed-point iterations** speed
 - Using convolutions **[Solomon+15]**
 - Faster kernel multiplications (*many refs!*)
 - Accelerations (Anderson **[Chizat+20]**)
 - Improved convergence proofs
- Generalize Sinkhorn to **continuous spaces?**
(Schrödinger bridges) **[Chen+14]** **[Bortoli+21]**
- **Other** regularizers (L_2) / constraints (low-rank)

A few research topics around OT

- **Unbalanced OT formulations**
 - Penalized **[FZMAP15][Chizat+15]**, approaches can handle (with Sinkhorn) different marginals
- **Brenier / Monge topics**
 - Estimate Monge maps using ICNN **[VM+19]**
 - Links with Normalizing Flows **[Huang+20]**
- **Proximal optimization** in spaces of measures.
 - **[JKO'98]** starting to make an impact in practice, see for instance **[Bunne+21]**

A few research topics around OT

- **OT for Heterogeneous spaces**
 - Use quadratic assignment problem to compute machines (GW, [Memoli'11])
 - Computational challenges, stats unknown.
- **Exploit matching differentiability**
 - soft-sorting [CTV19] and -quantiles [CTNV'20]
 - to impute missing values [Muzellec+20]