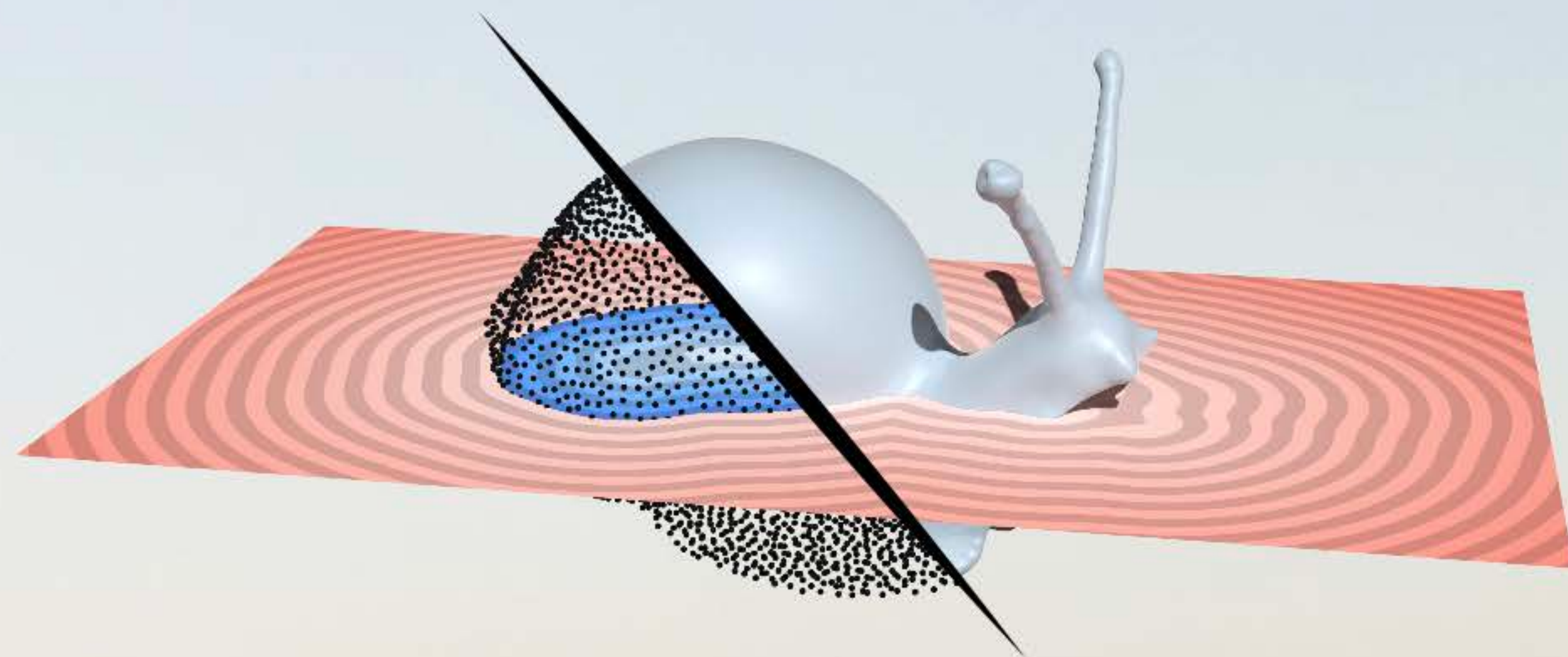
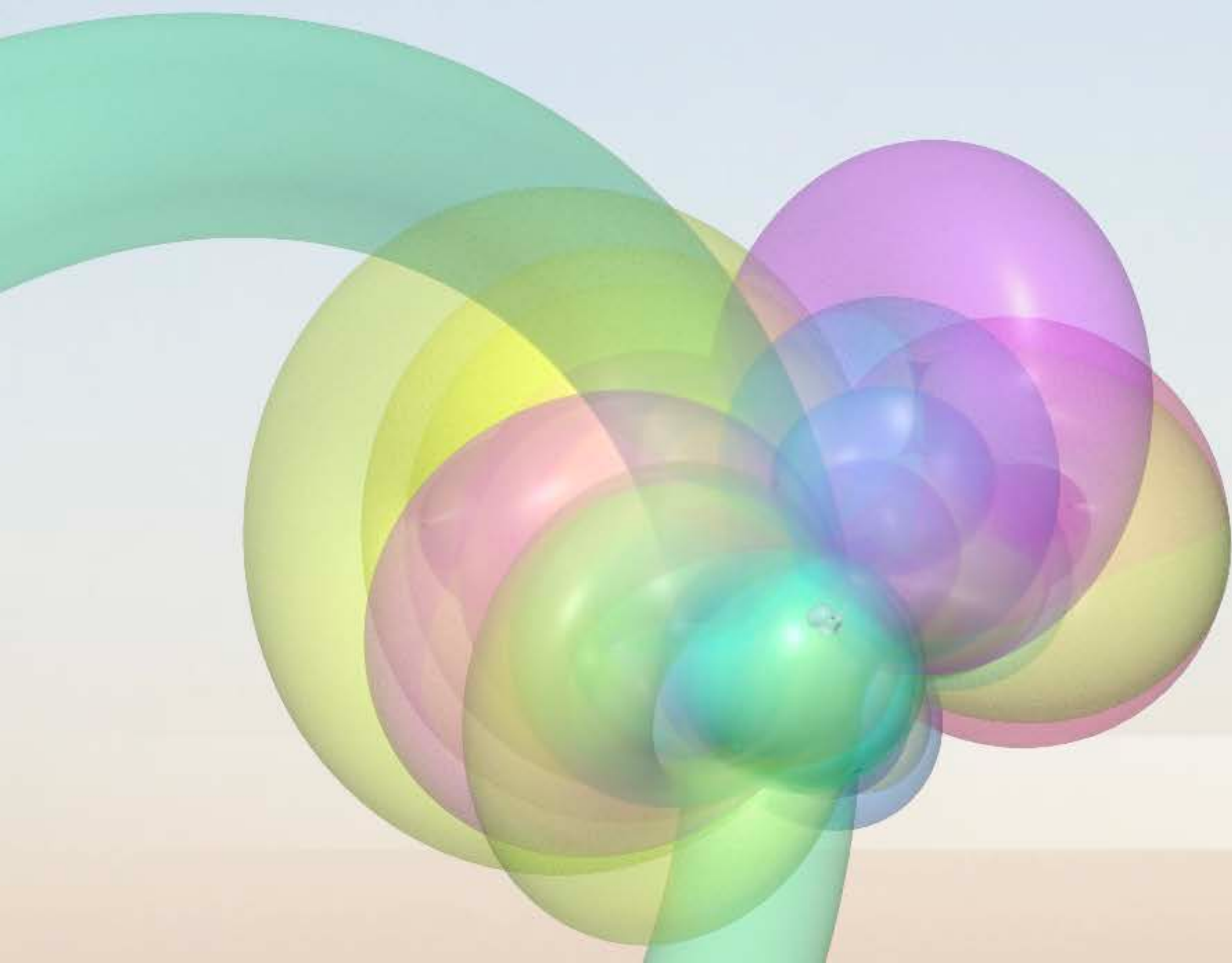


POINTS AS TORI: FAST POINTWISE SIGNED DISTANCE FOR POINT CLOUDS

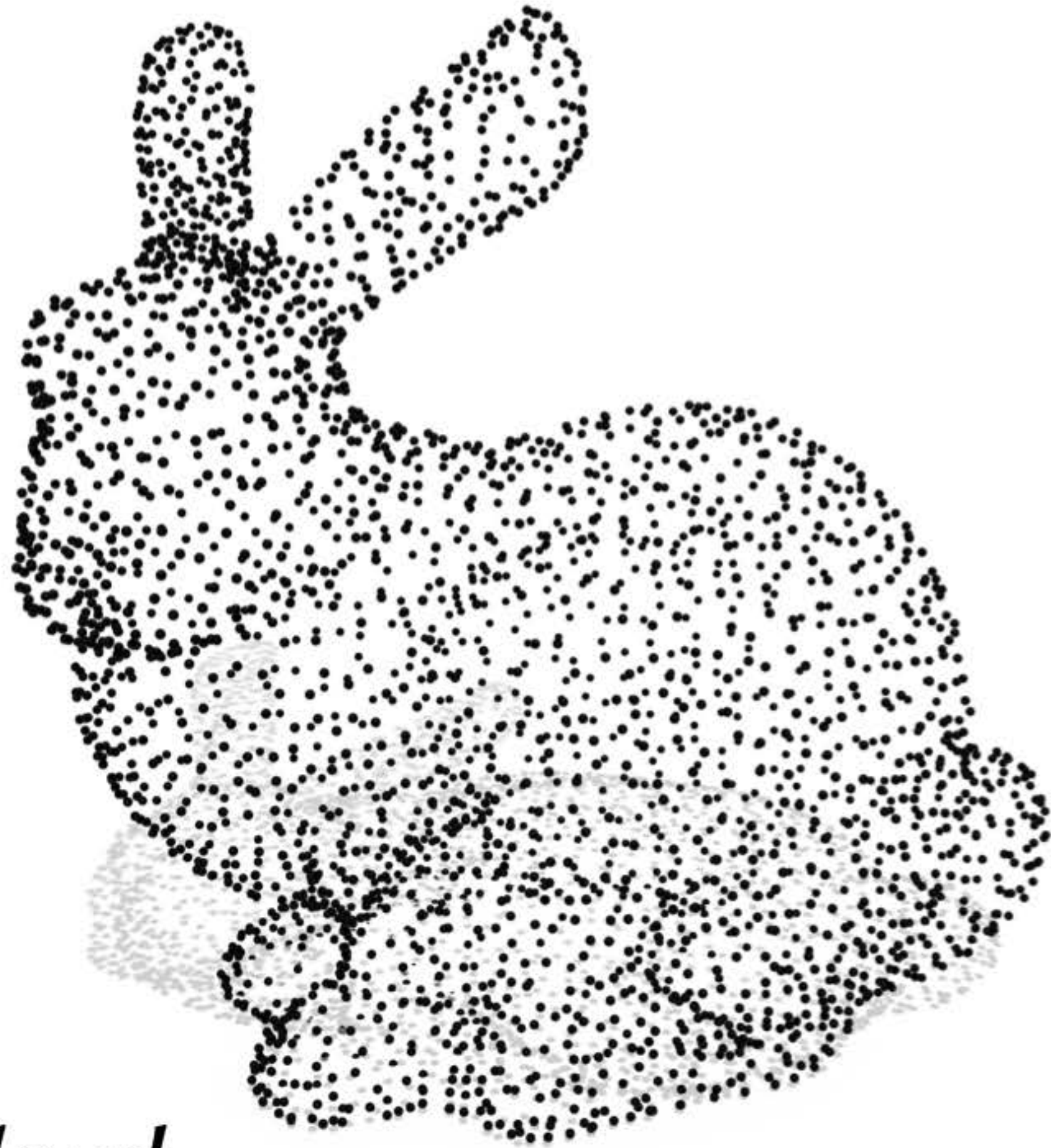
Nicole Feng, Ioannis Gkioulekas*, Keenan Crane*

Carnegie Mellon University



The problem

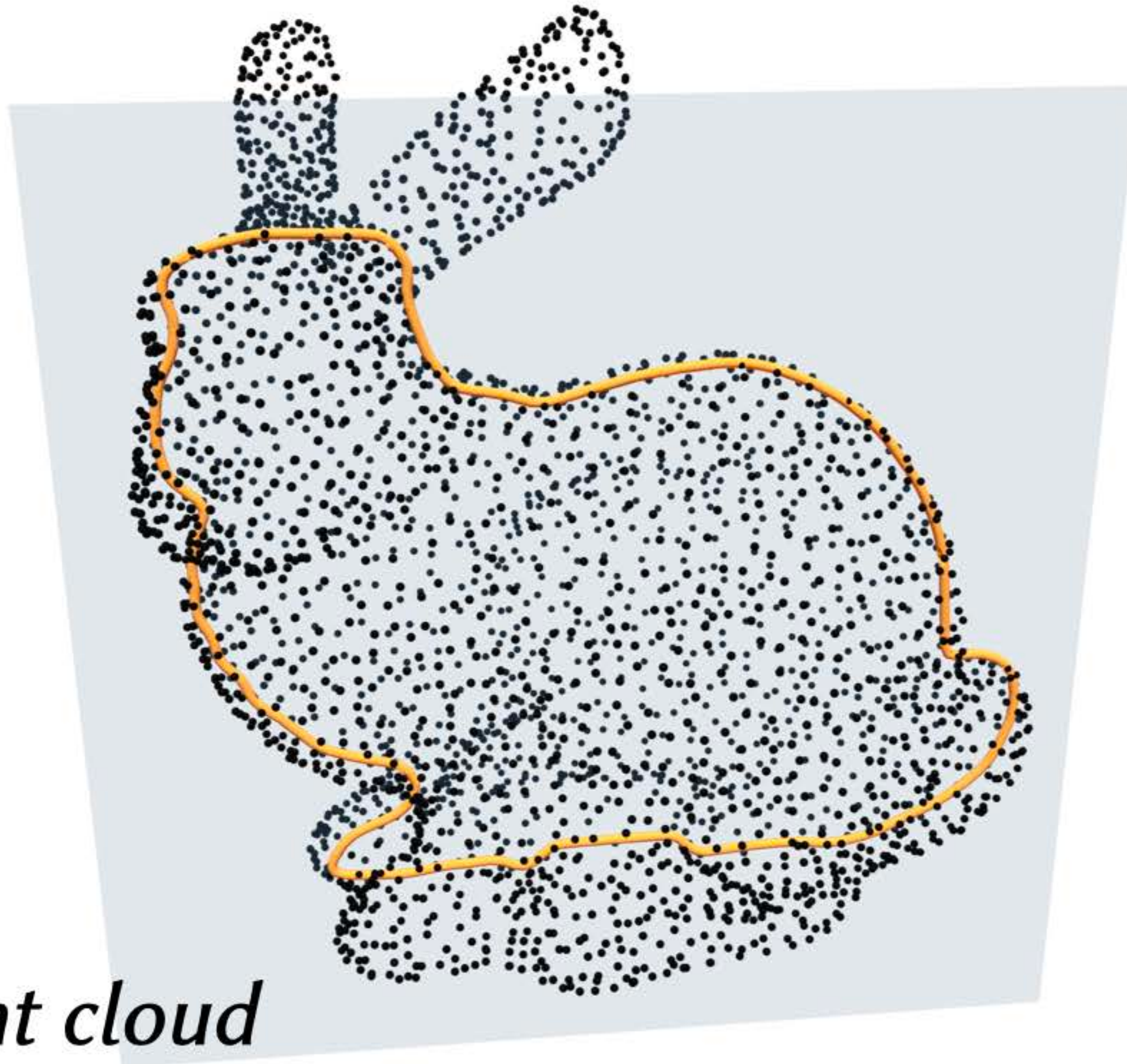
The problem



point cloud

scanning, vision, modeling...

The problem



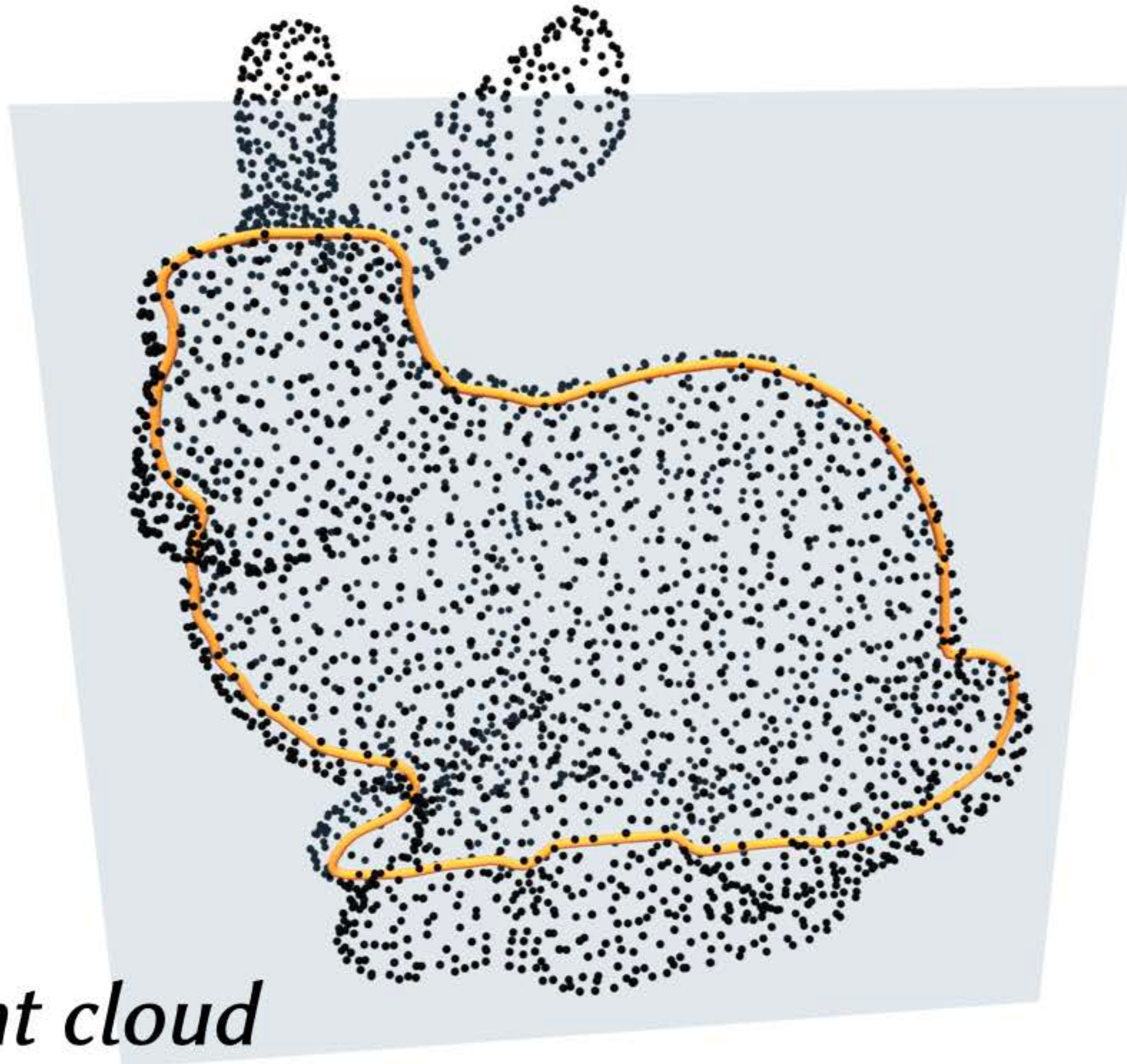
point cloud

scanning, vision, modeling...

signed distance function (SDF)



The problem



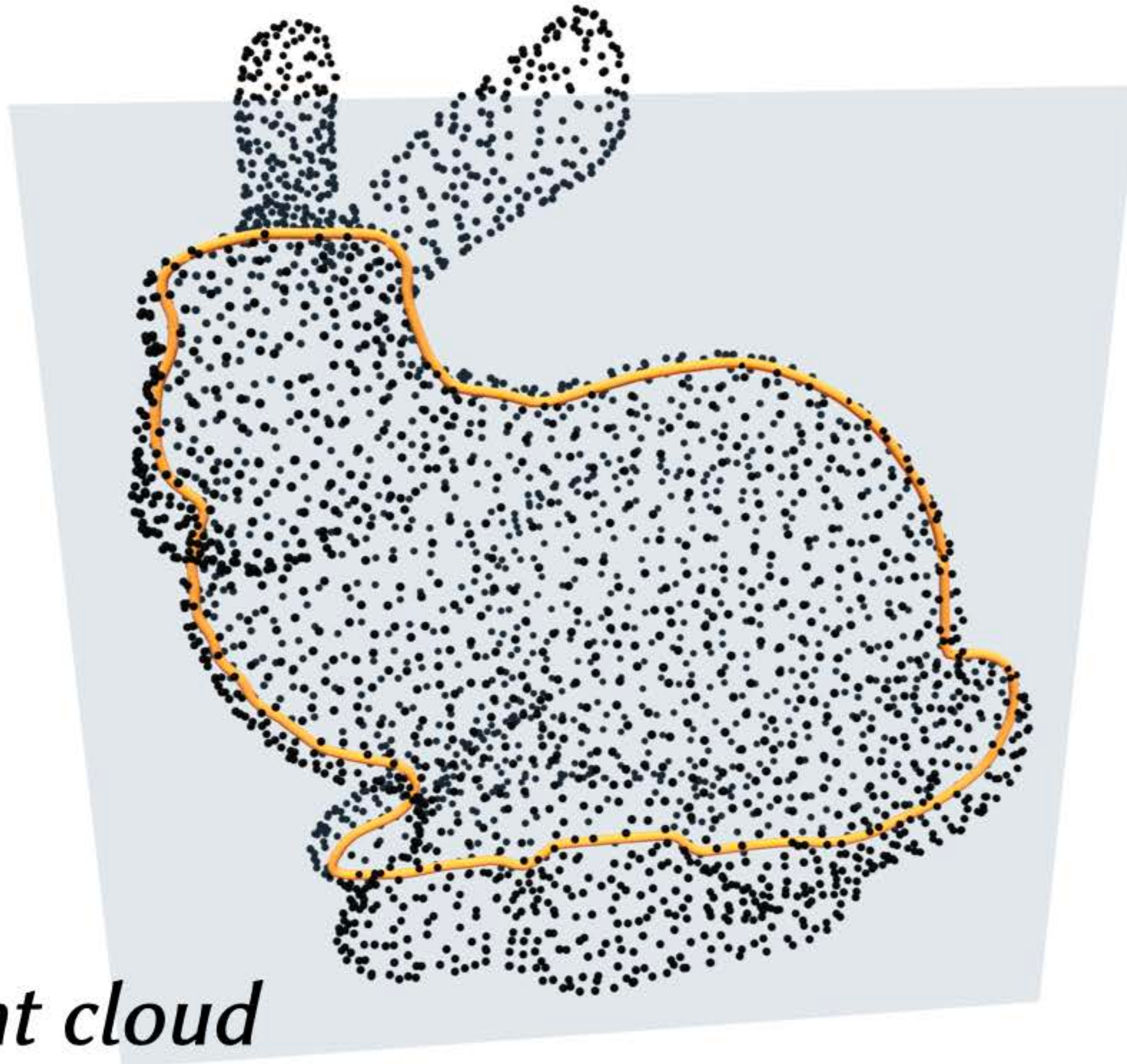
point cloud

scanning, vision, modeling...

signed distance function (SDF)



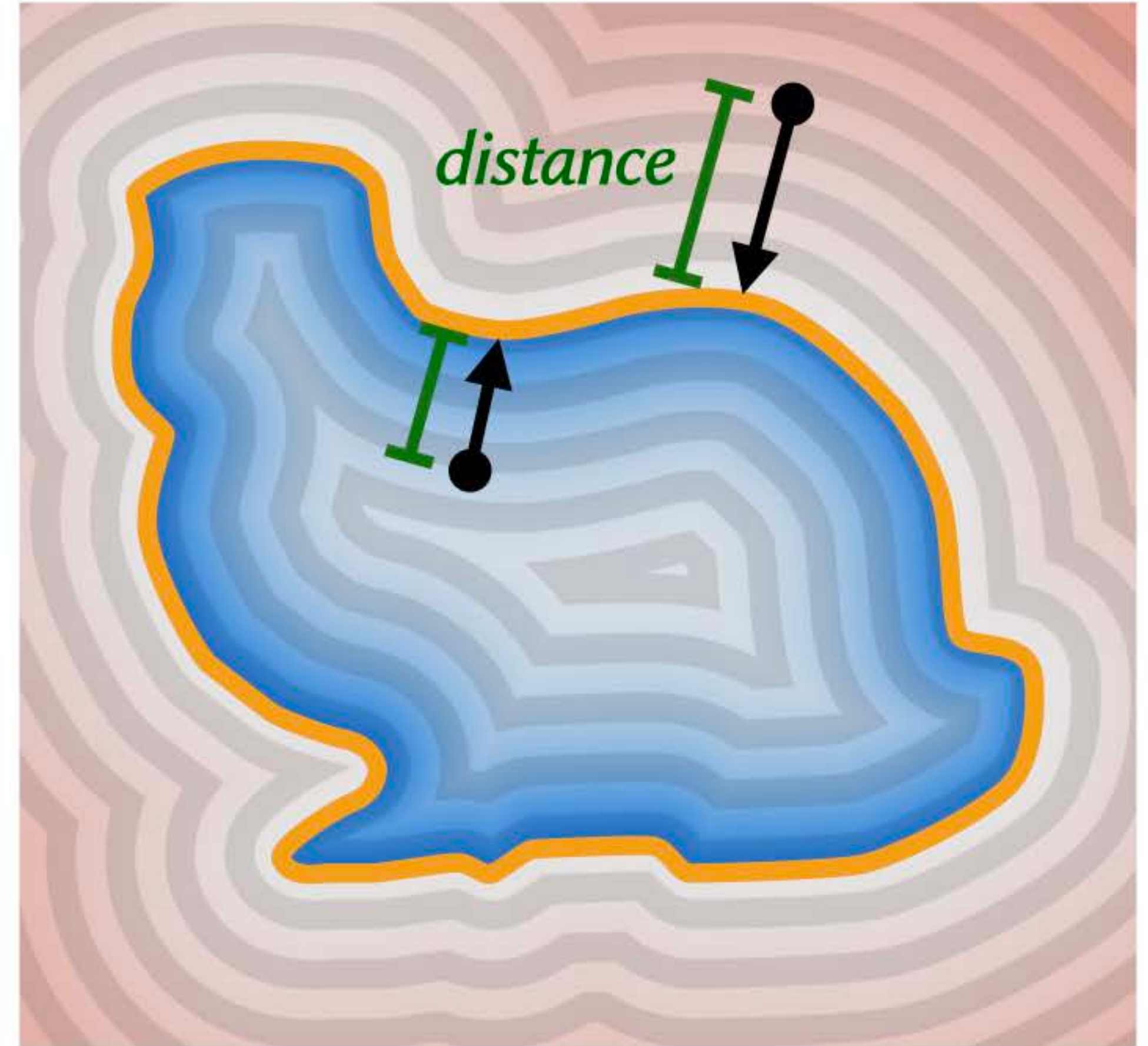
The problem



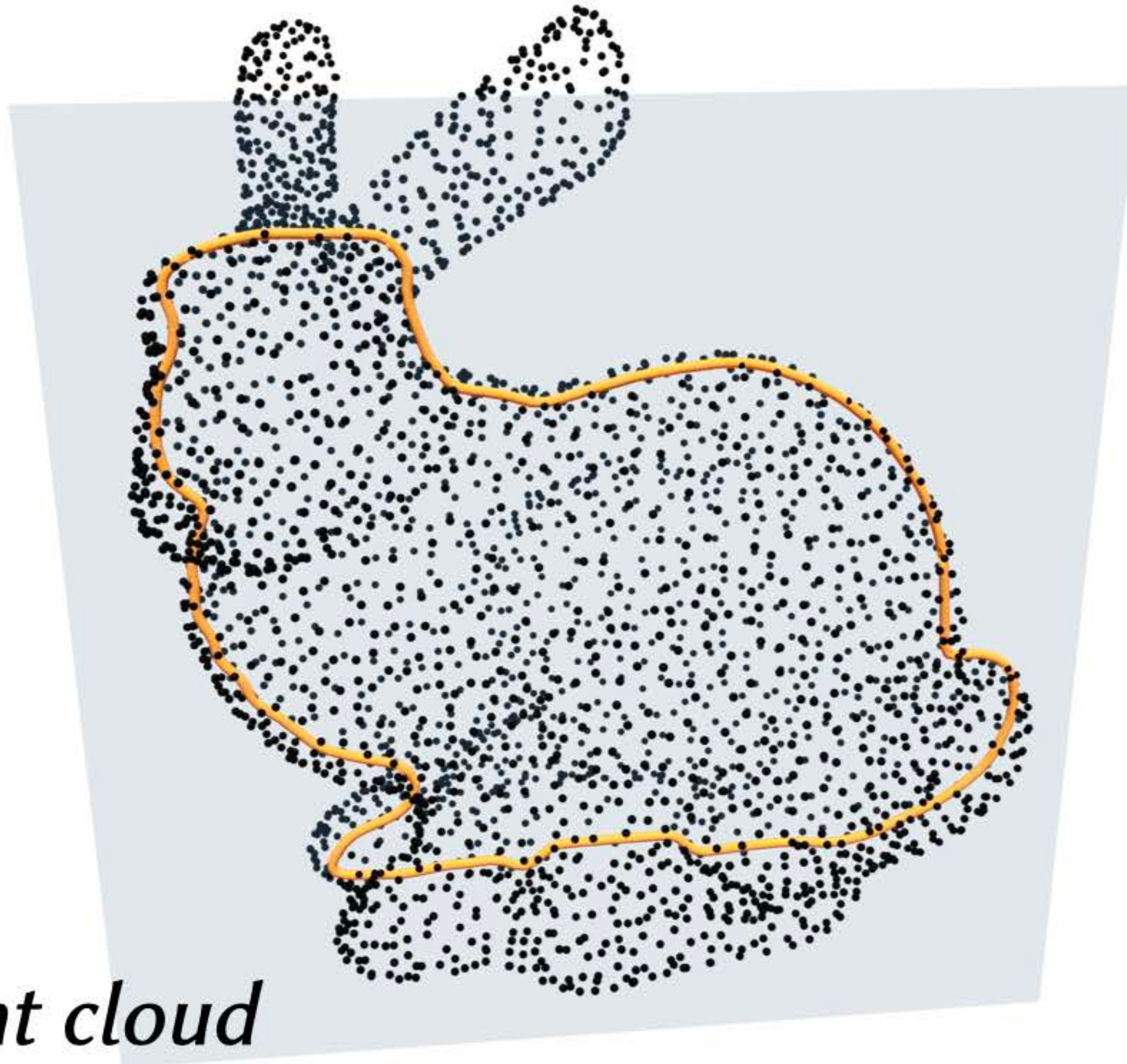
point cloud

scanning, vision, modeling...

signed distance function (SDF)



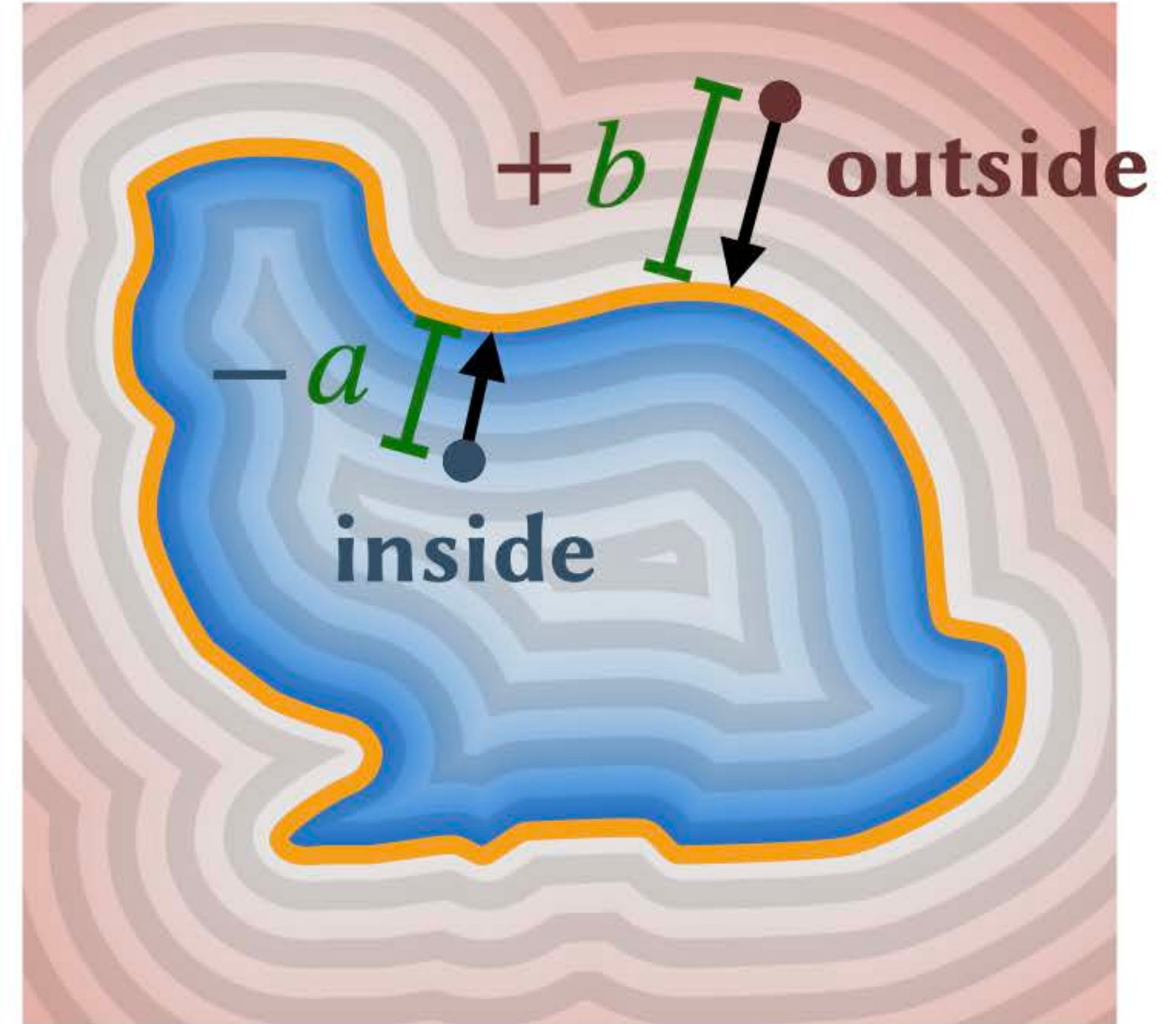
The problem



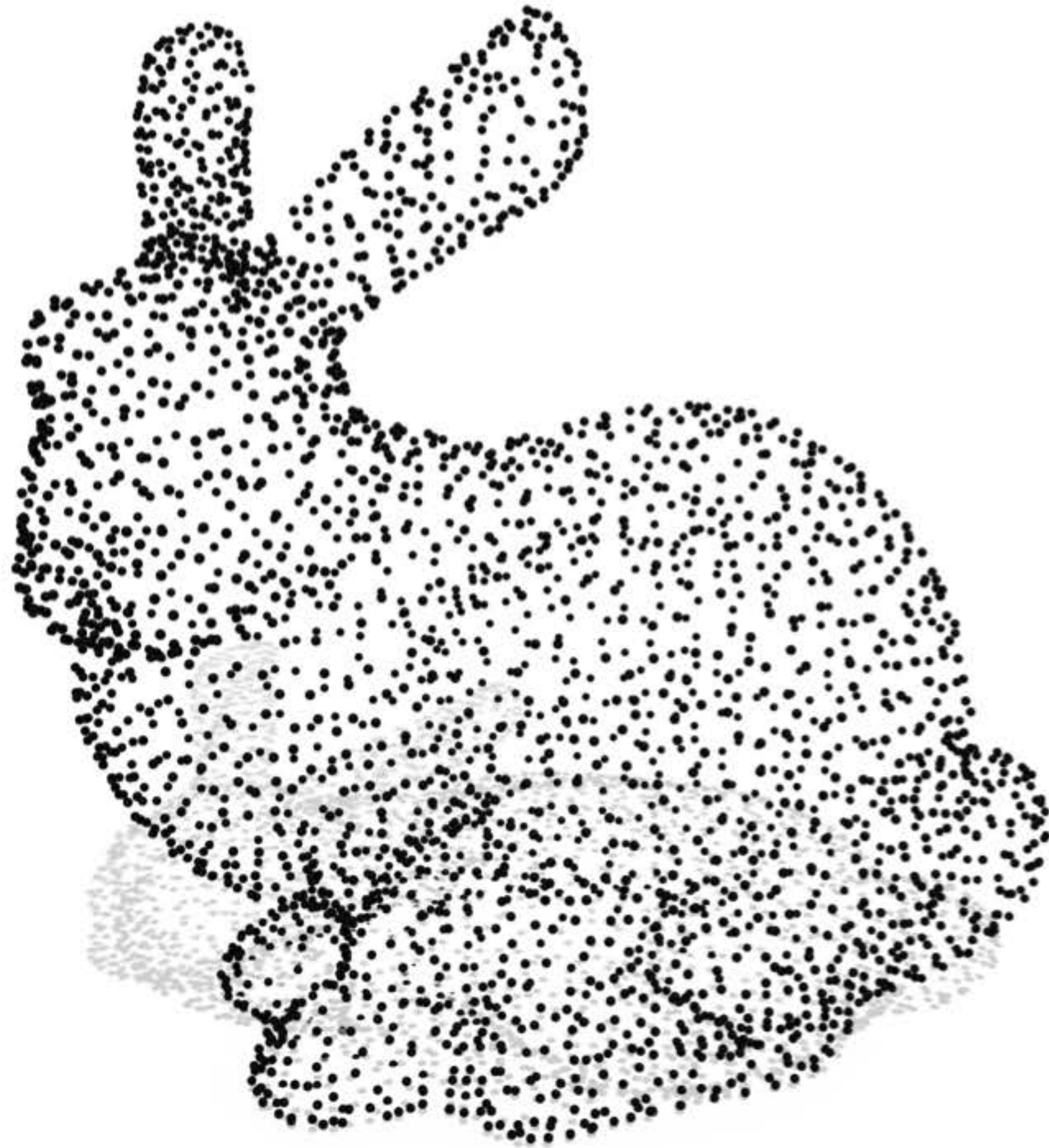
point cloud

scanning, vision, modeling...

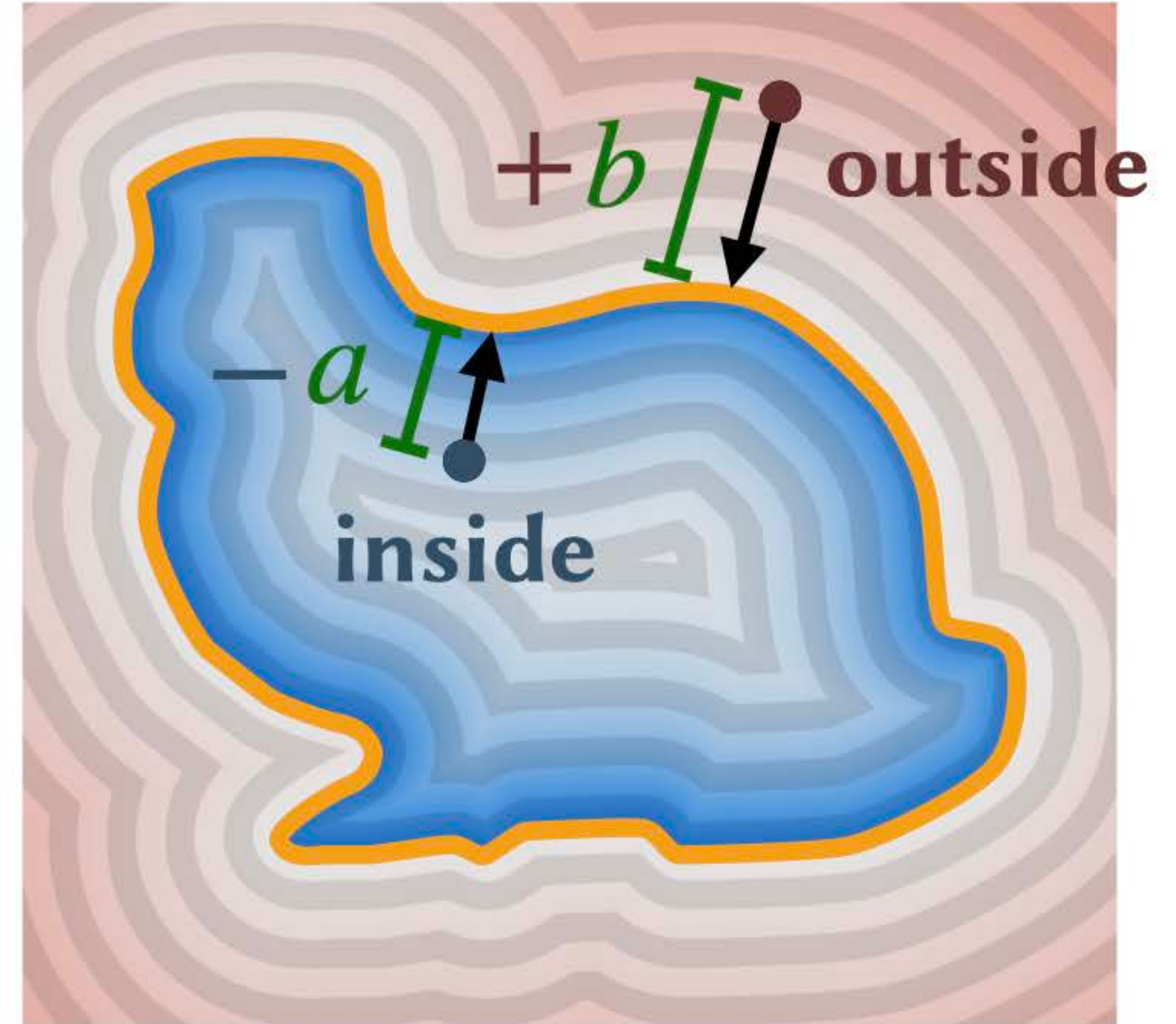
signed distance function (SDF)



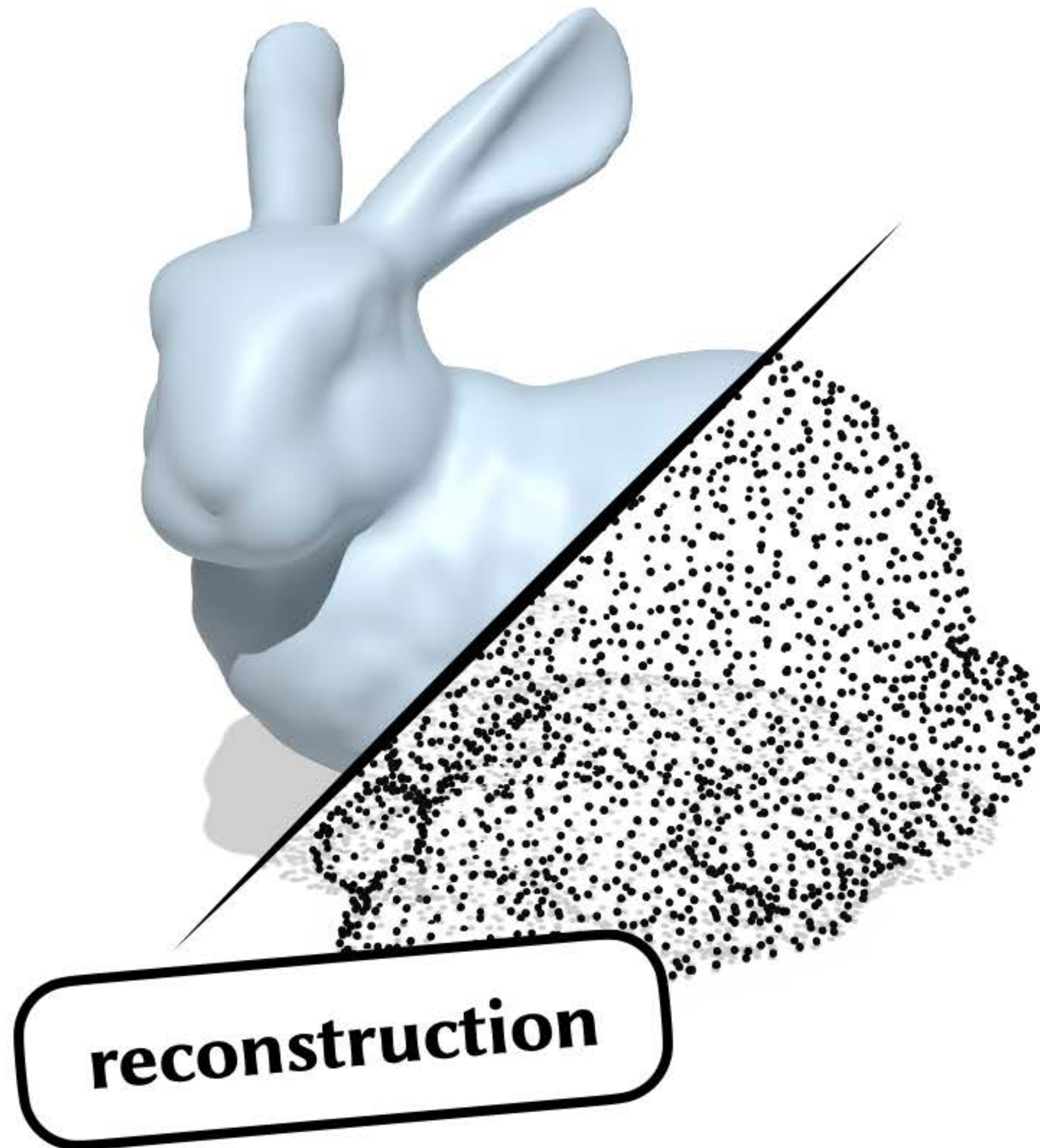
Signed distance requires both reconstruction and distance



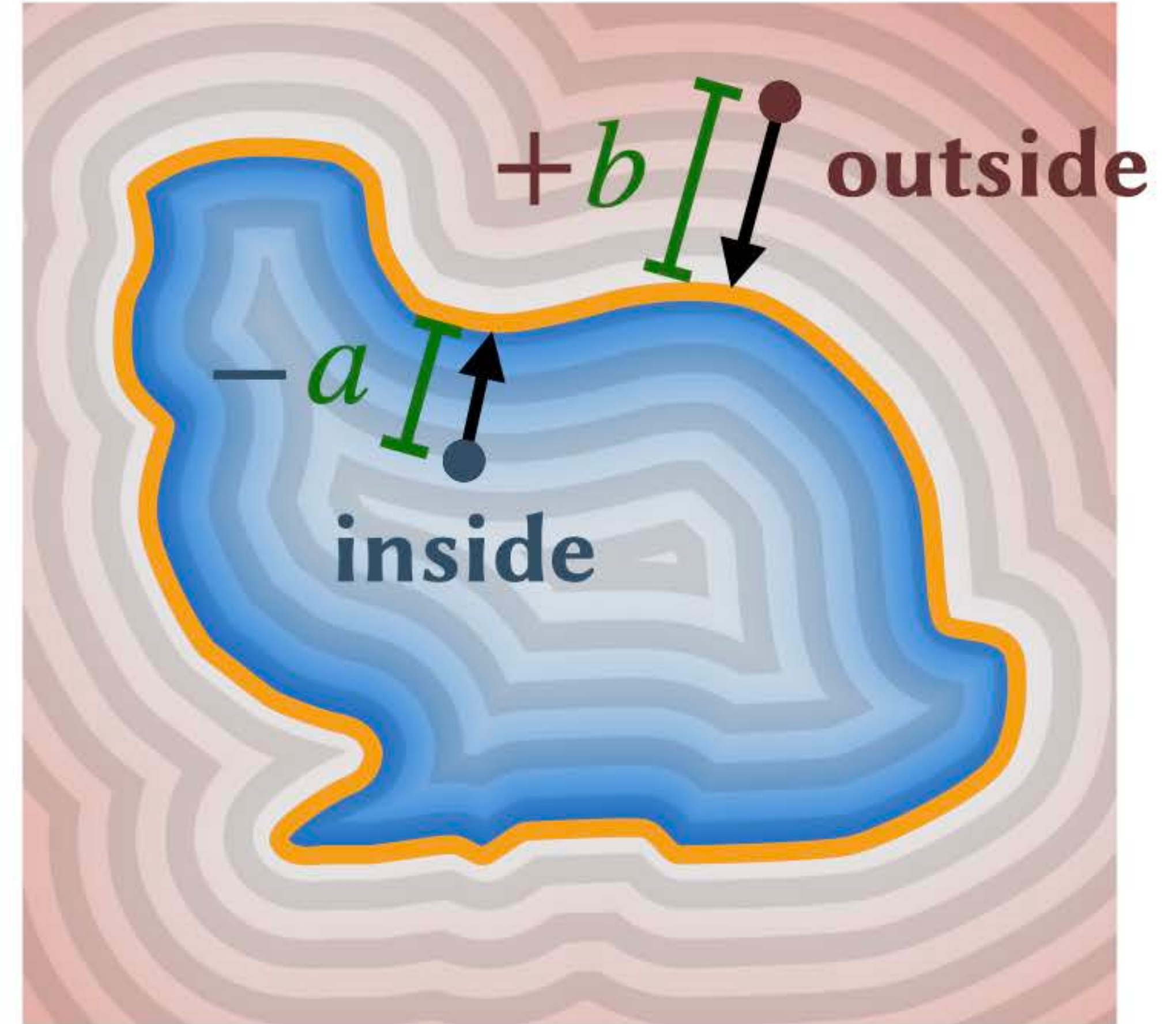
signed distance function (SDF)



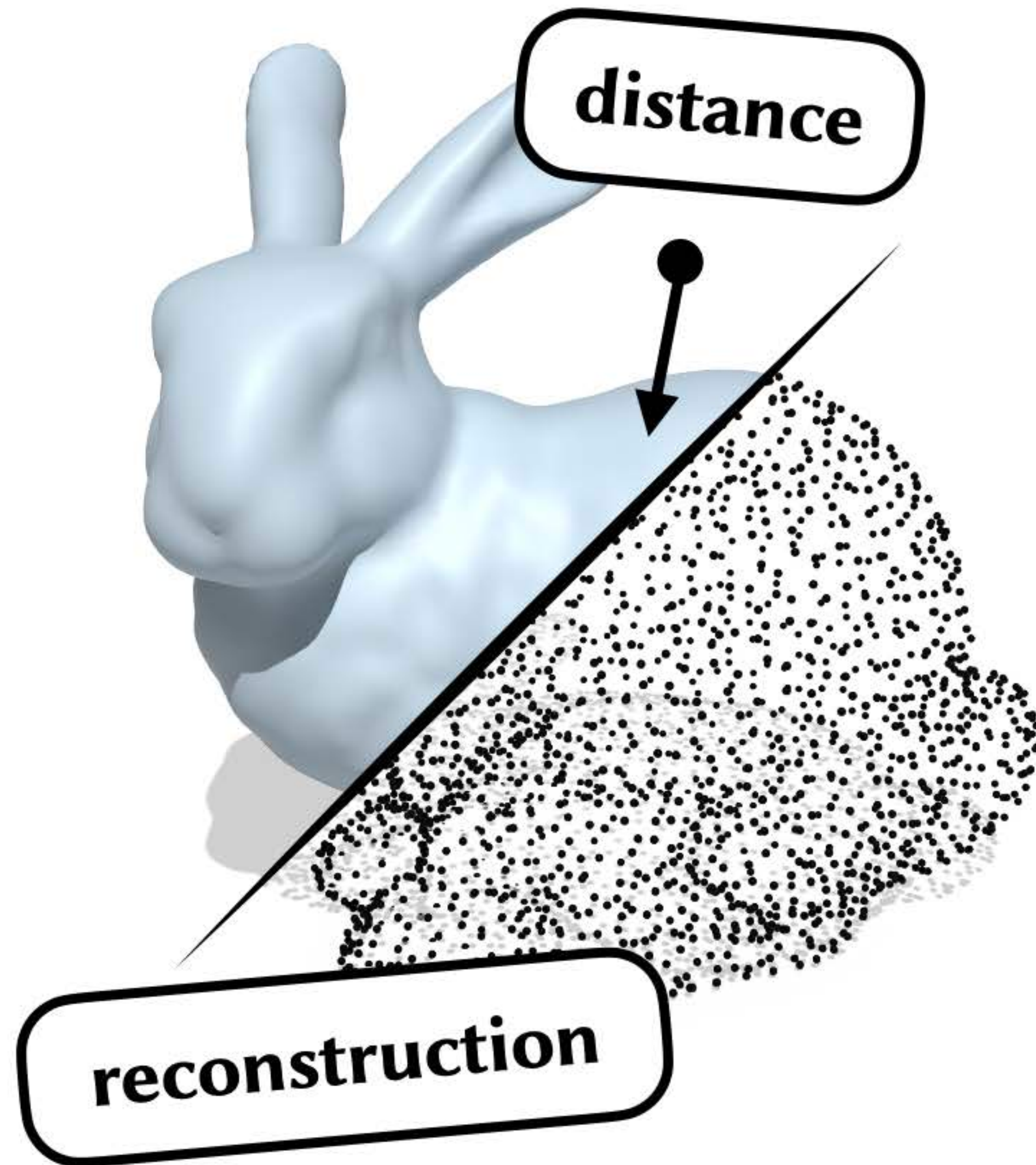
Signed distance requires both reconstruction and distance



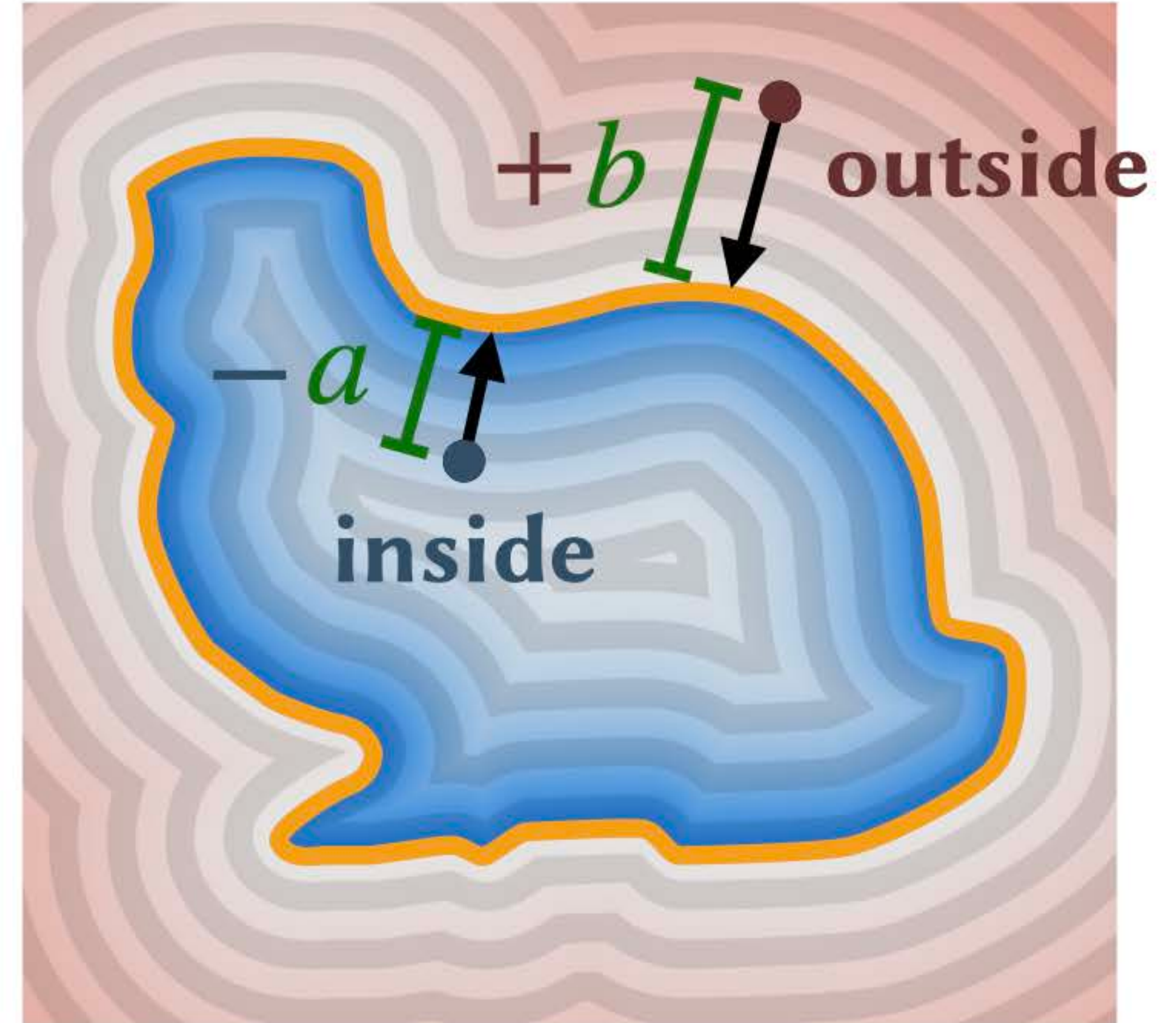
signed distance function (SDF)



Signed distance requires both reconstruction and distance



signed distance function (SDF)



Generalized distance algorithms are expensive

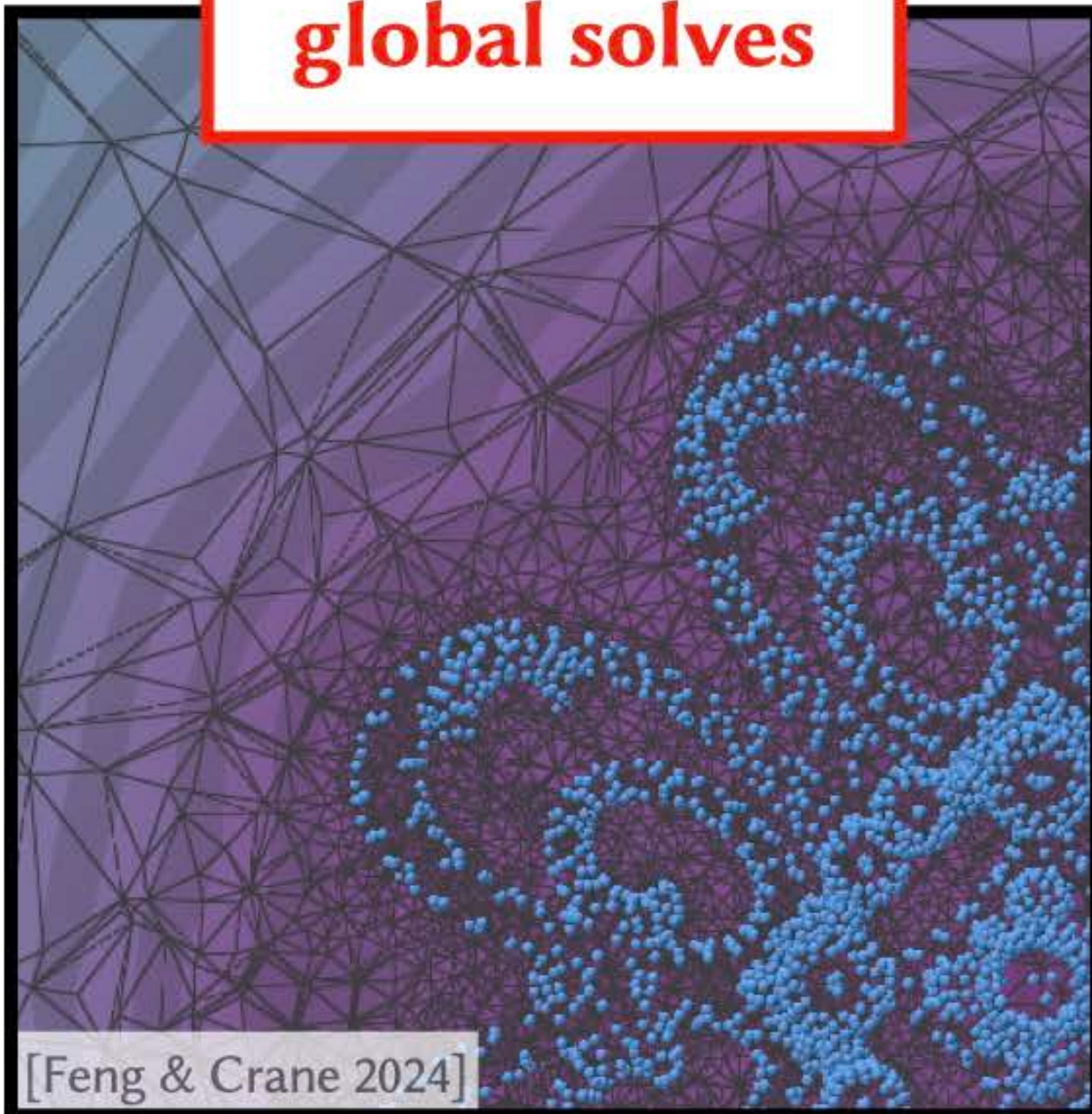
Generalized distance algorithms are expensive

*SSD: Smoothed Signed Distance
Surface Reconstruction*
[Calakli & Taubin 2011]

*A Heat Method for
Generalized Signed Distance*
[Feng & Crane 2024]

etc....

global solves



Generalized distance algorithms are expensive

SSD: Smoothed Signed Distance Surface Reconstruction
[Calakli & Taubin 2011]

A Heat Method for Generalized Signed Distance
[Feng & Crane 2024]

etc....

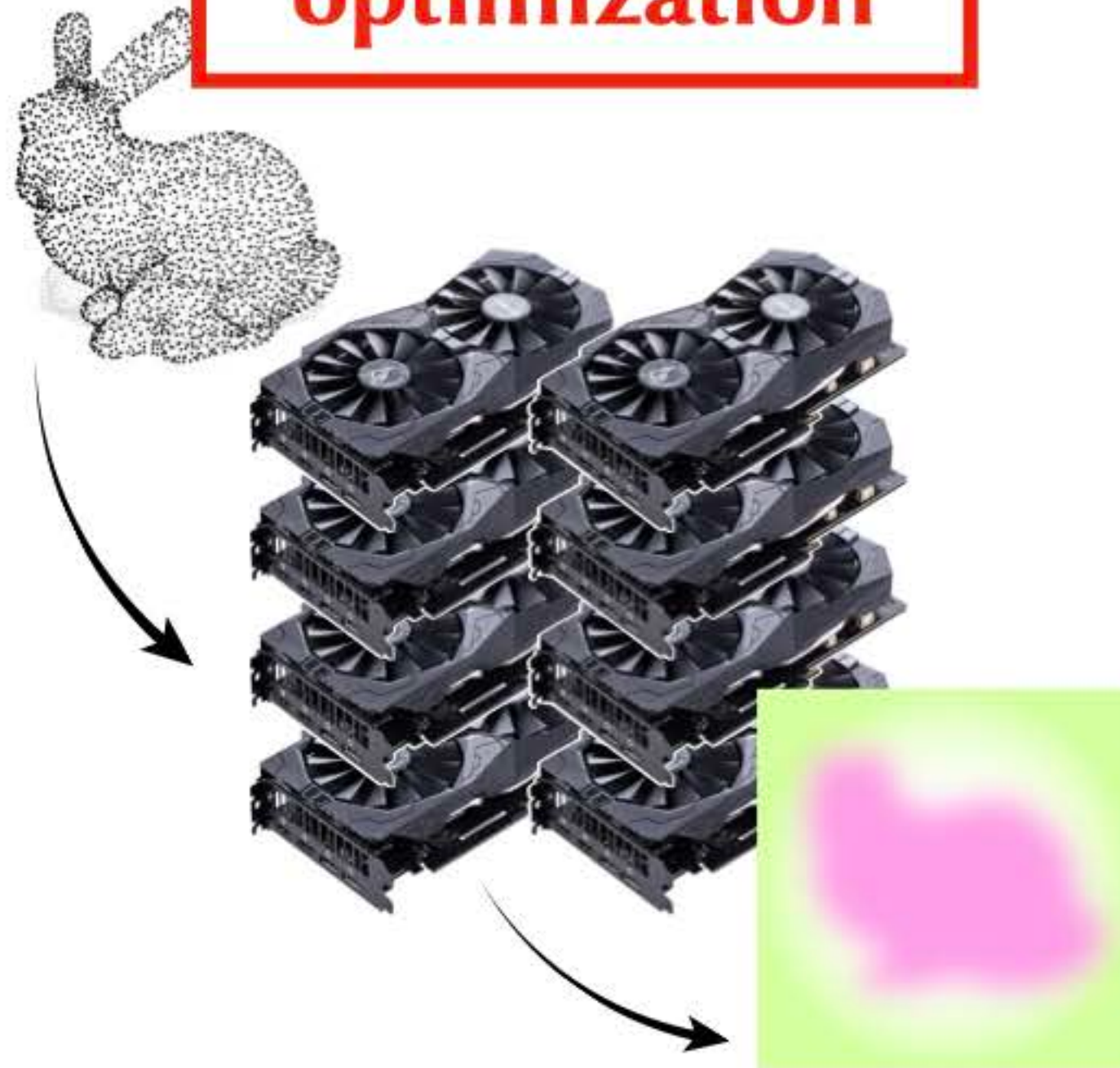
global solves

Phase Transitions, Distance Functions, and Implicit Neural Representations
[Lipman 2019]

Neural Unsigned Distance Fields for Implicit Function Learning
[Chibane et al. 2020]

etc....

expensive optimization



Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces
[Ma et al. 2021]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian
[Wang et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation
[Yang et al. 2023]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds
[Park et al. 2023]

1-Lipschitz Neural Distance Fields
[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition
[Wang et al. 2025]

SDFs from Unoriented Point Clouds using Neural Variational Heat Distances
[Weidemaier et al. 2025]

efunc: An Efficient Function Representation without Neural Networks
[Zhang & Wonka 2025]

[Feng & Crane 2024]

Generalized distance algorithms are expensive

SSD: Smoothed Signed Distance Surface Reconstruction
[Calakli & Taubin 2011]

A Heat Method for Generalized Signed Distance
[Feng & Crane 2024]

etc....

global solves

Phase Transitions, Distance Functions, and Implicit Neural Representations
[Lipman 2019]

Neural Unsigned Distance Fields for Implicit Function Learning
[Chibane et al. 2020]

etc....

expensive optimization

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces
[Ma et al. 2021]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian
[Wang et al. 2023]

StEik: Stochastic Estimation of Neural Signed Distance Functions

not true distance

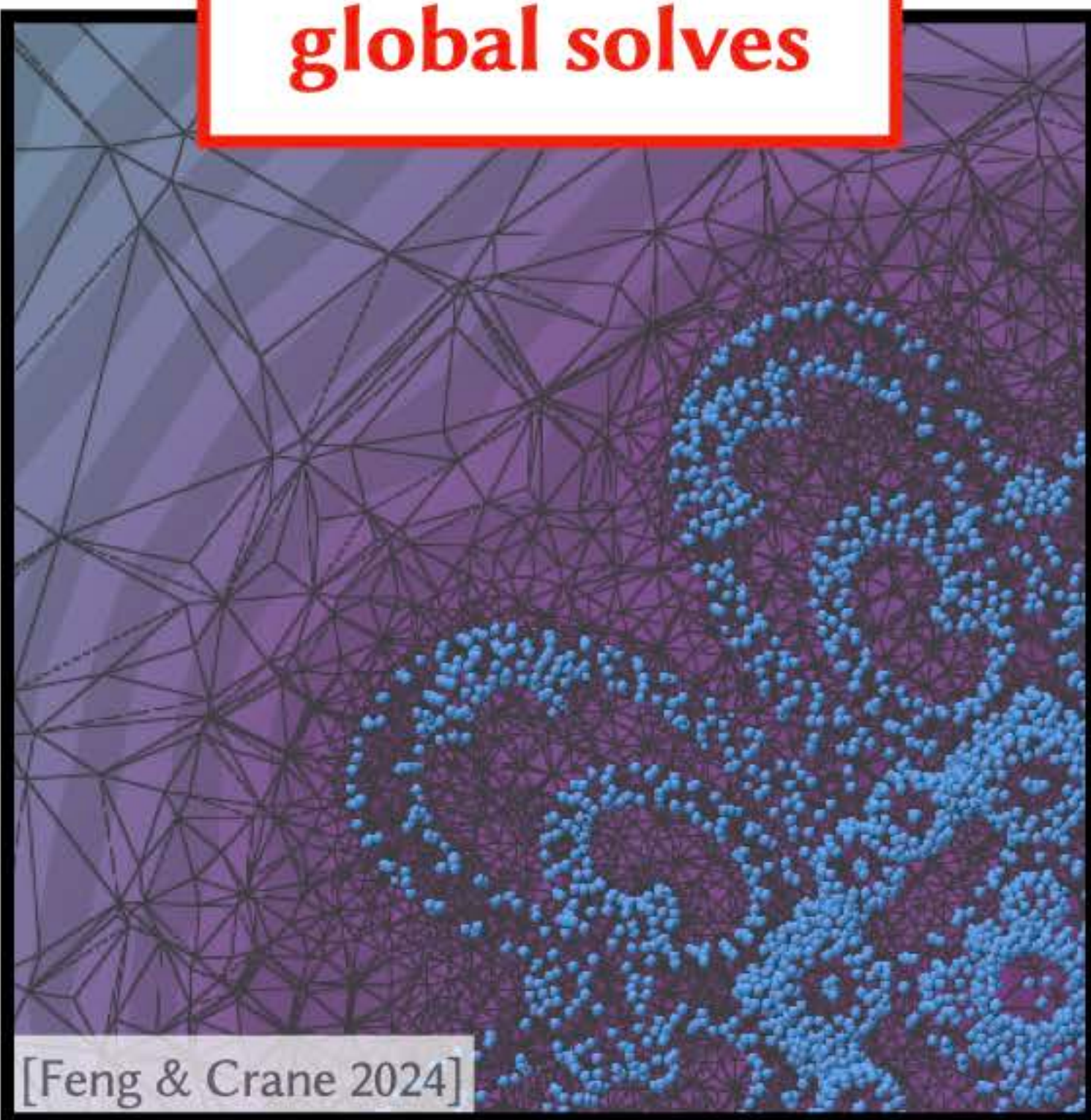
Surface Reconstruction in Point Clouds

[Lipman 2019]

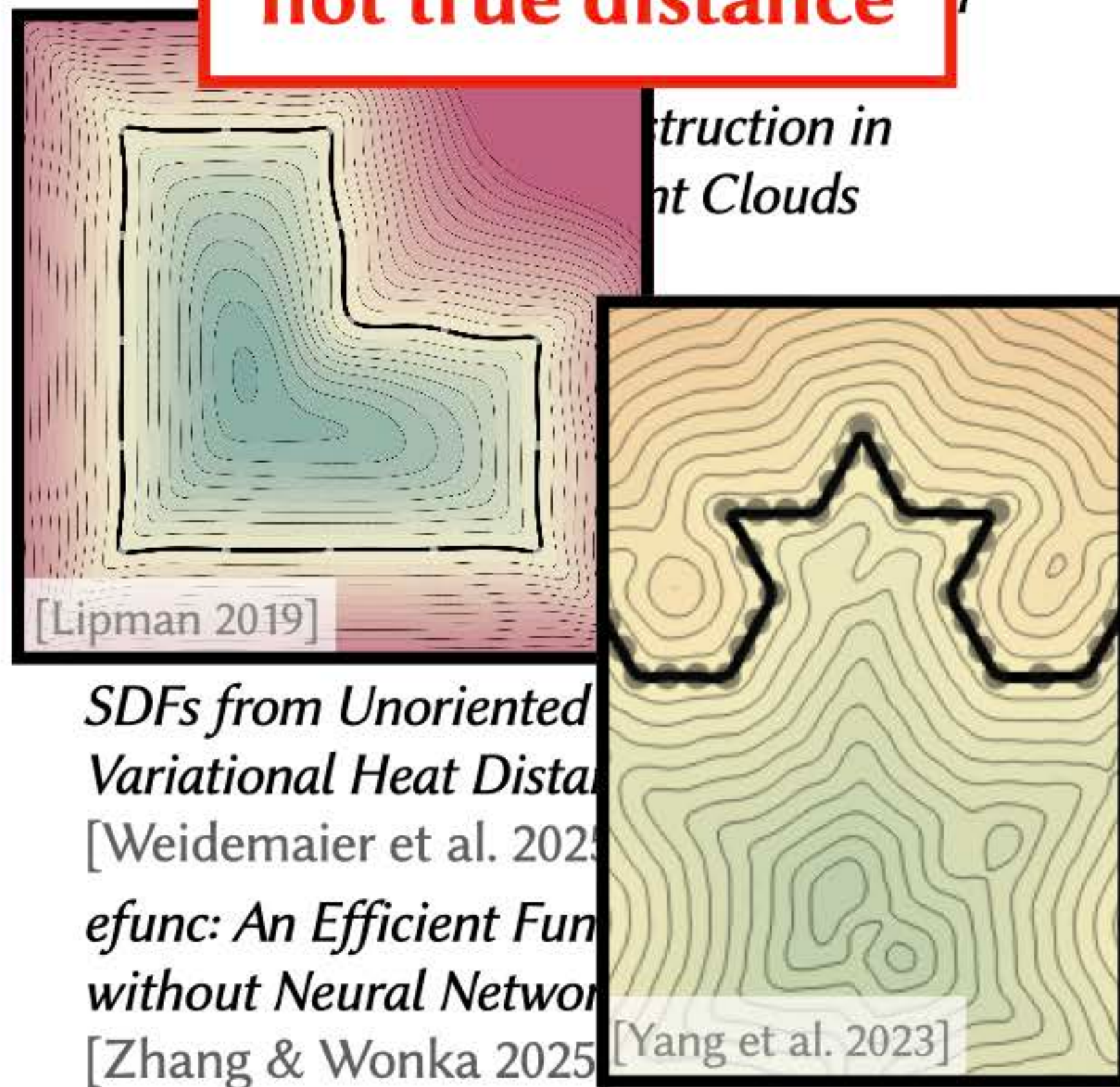
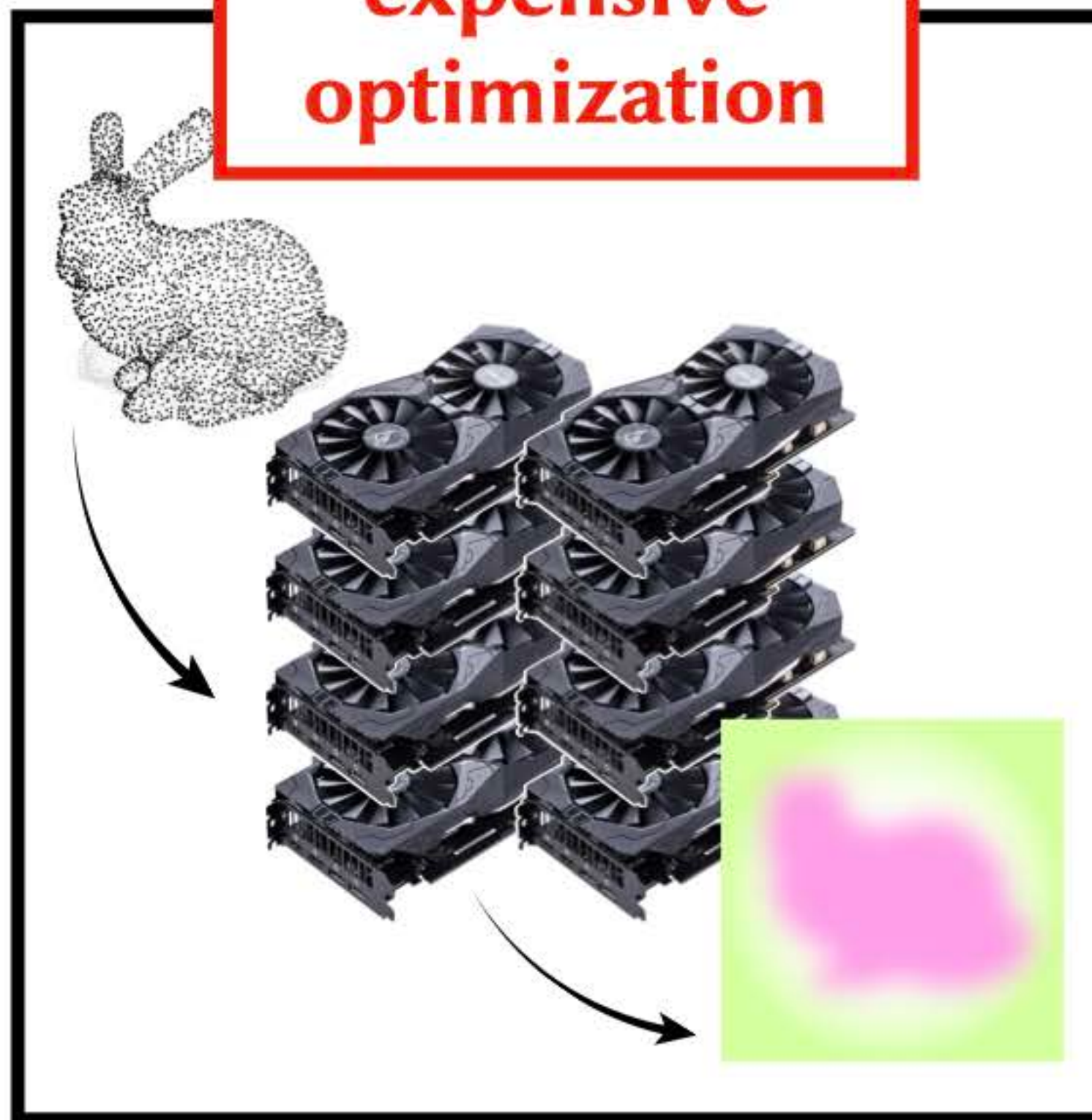
SDFs from Unoriented Variational Heat Distance
[Weidemaier et al. 2025]

efunc: An Efficient Function without Neural Networks
[Zhang & Wonka 2025]

Neural Signed Distance Functions
[Yang et al. 2023]



[Feng & Crane 2024]



Fast reconstruction algorithms don't compute distance

Fast reconstruction algorithms don't compute distance

meta-balls

Generalization of Algebraic Surface Drawing
[Blinn 1982]

Convolution Surfaces
[Bloomenthal & Shoemake 1991]

Alternate Distance Metrics for Implicit Surface Modeling
[Tigges et al. 1999]

Interpolating implicit surfaces from scattered surface data using compactly supported radial basis functions
[Morse et al. 2005]

etc. ...

winding numbers

Fast Winding Numbers for Soups and Clouds
[Barill et al. 2018]

3D Reconstruction with Fast Dipole Sums
[Chen et al. 2024]

moving least squares / partition of unity

Surfaces generated by moving least squares methods
[Lancaster & Salkauskas 1981]

The Approximation Power of Moving Least-Squares
[Levin 1998]

Variational Implicit Surfaces
[Turk & O'Brien 1999]

Reconstruction and representation of 3D objects with radial basis functions
[Carr et al. 2001]

Defining point-set surfaces
[Amenta & Kil 2004]

Interpolating and Approximating Implicit Surfaces from Polygon Soup
[Shen et al. 2004]

An adaptive MLS surface for reconstruction with guarantees
[Dey & Sun 2005]

Multi-level partition of unity implicits
[Ohtake et al. 2005]

Anisotropic point set surfaces
[Adamson & Alexa 2006]

Feature Preserving Point Set Surfaces based on Non-Linear Kernel Regression
[Öztireli et al. 2009, etc.]

A Curvature-Adaptive Implicit Surface Reconstruction for Irregularly Spaced Points
[Zagorchev & Goshtasby 2012]

Floating Scale Surface Reconstruction
[Fuhrmann & Goesele 2014]
etc. ...

convolutional “distance”

Provably good moving least squares
[Kolluri 2008]

A Schrödinger Equation for the Fast Computation of Approximate Euclidean Distance Functions
[Gurumoorthy & Rangarajan 2009]

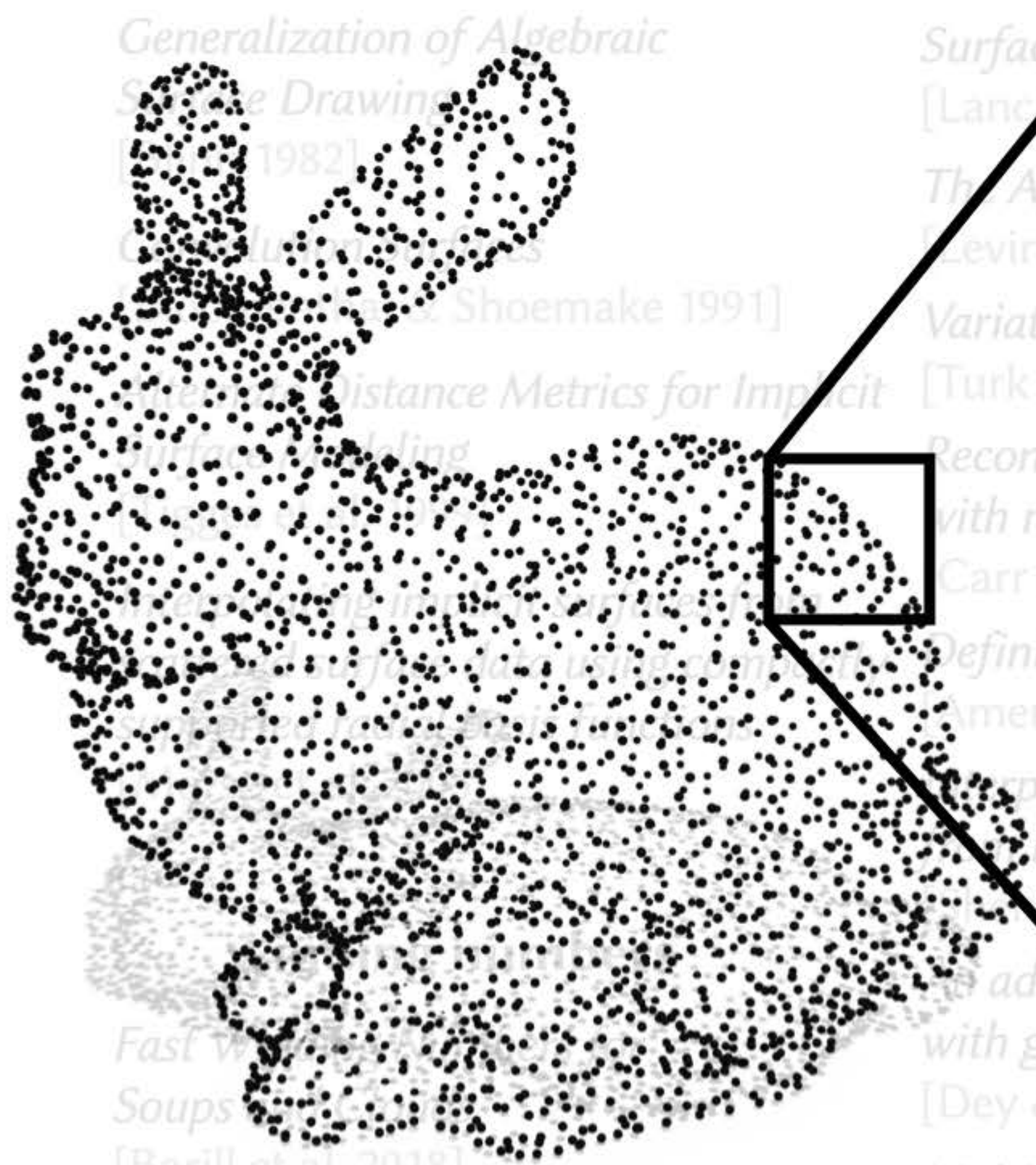
Fast Convolutional Distance Transform
[Karam et al. 2019]

Evaluating a distance function
[Abgrall 2022]

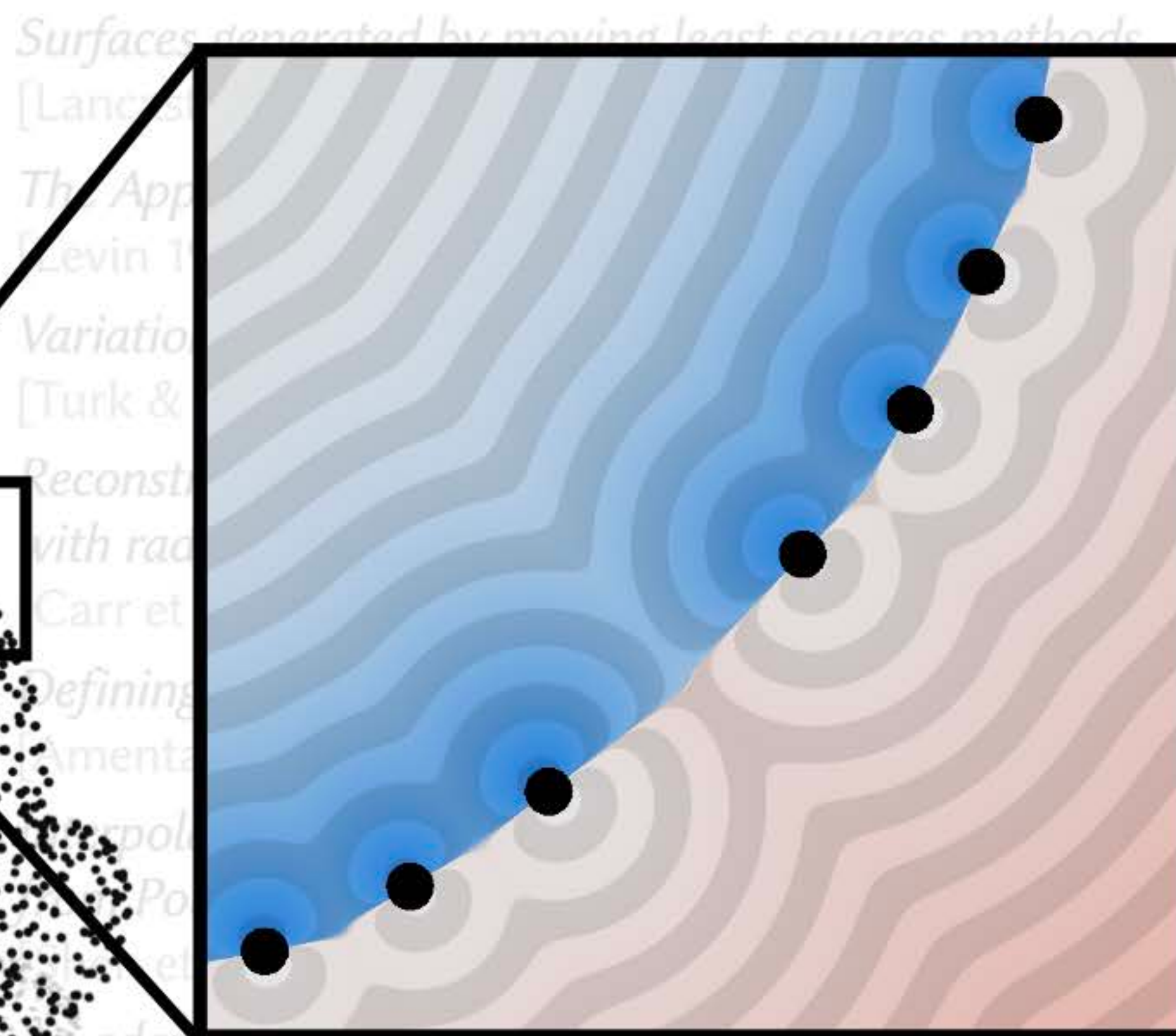
Fast evaluation of smooth distance constraints on co-dimensional geometry
[Madan & Levin 2022]

Fast reconstruction algorithms don't compute distance

meta-balls

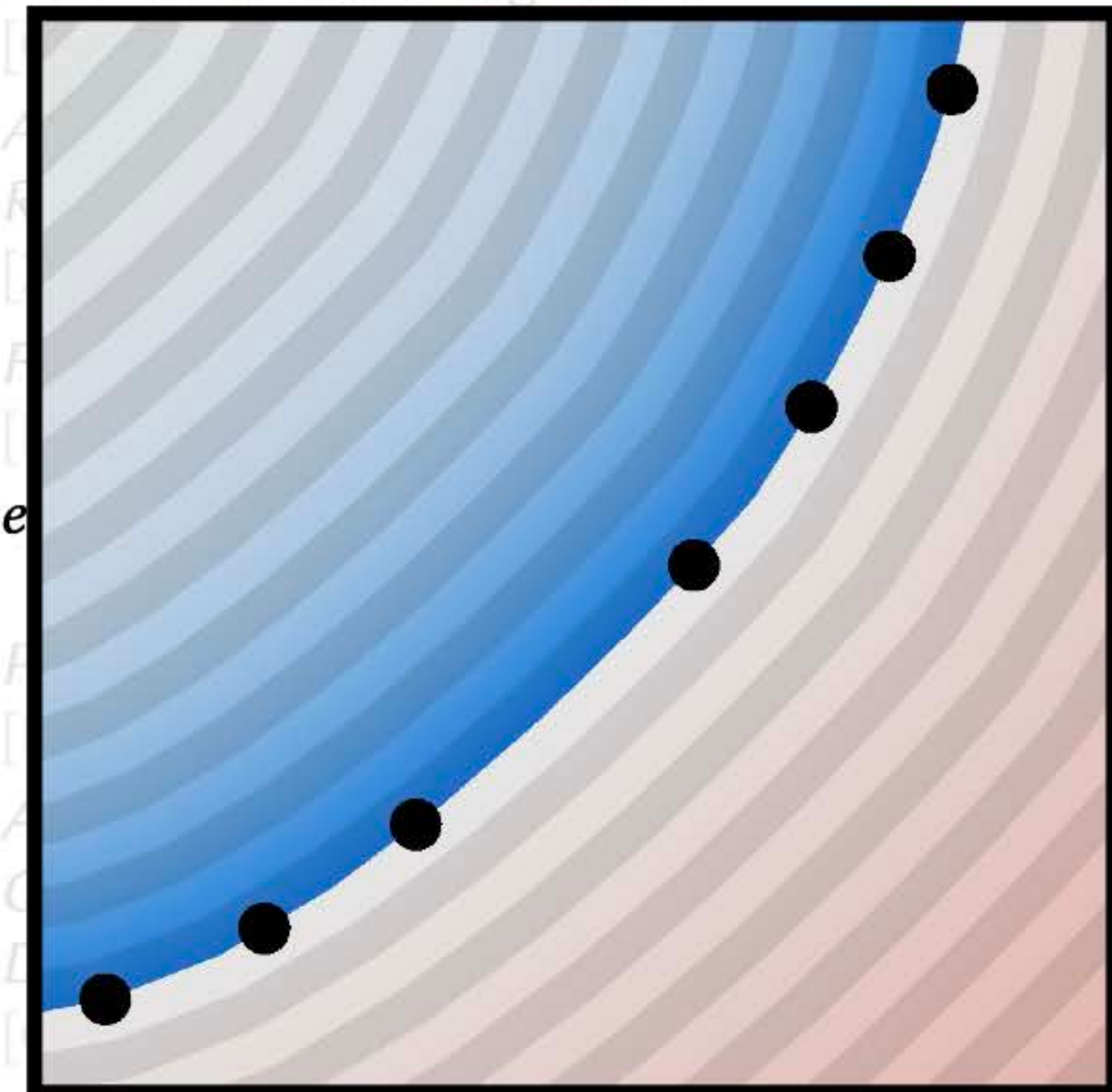


moving least squares / partition of unity



past "distance" methods

Feature Preserving Point Set Surfaces based on Non-Linear Kernel Regression



ground truth

Points as Tori (PAT)

Points as Tori (PAT)

Input: point cloud with normals, and a query point $\mathbf{x} \in \mathbb{R}^3$

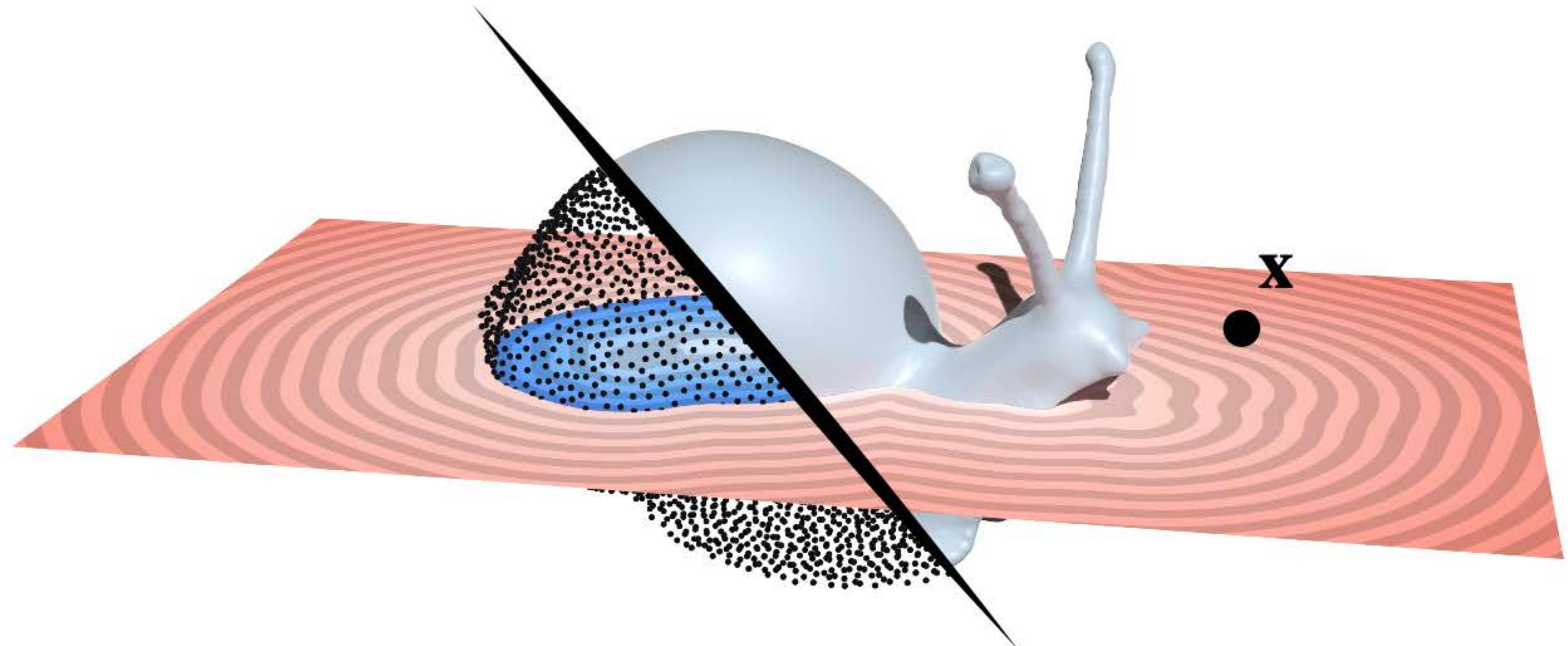


Points as Tori (PAT)

Input: point cloud with normals, and a query point $\mathbf{x} \in \mathbb{R}^3$



Output: signed distance at $\mathbf{x}$ to an approximate underlying surface



Points as Tori (PAT)

Points as Tori (PAT)

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

kernel density estimator

Points as Tori (PAT)

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

weights

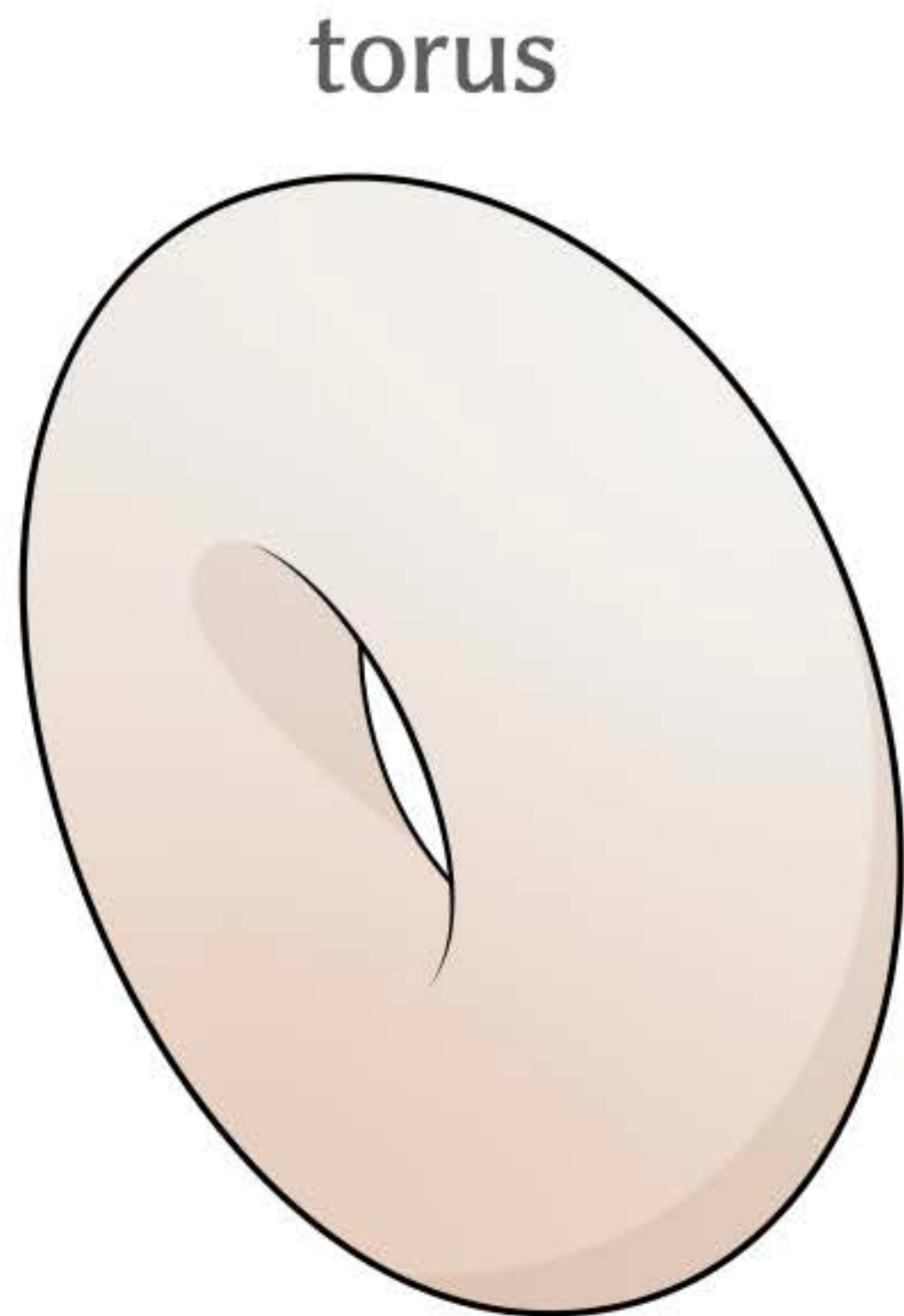
kernel density estimator

Points as Tori (PAT)

$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{i\text{-th function}}{g_i(x)} \overset{\text{weights}}{\exp(-\lambda \|x - p_i\|)}}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

kernel density estimator

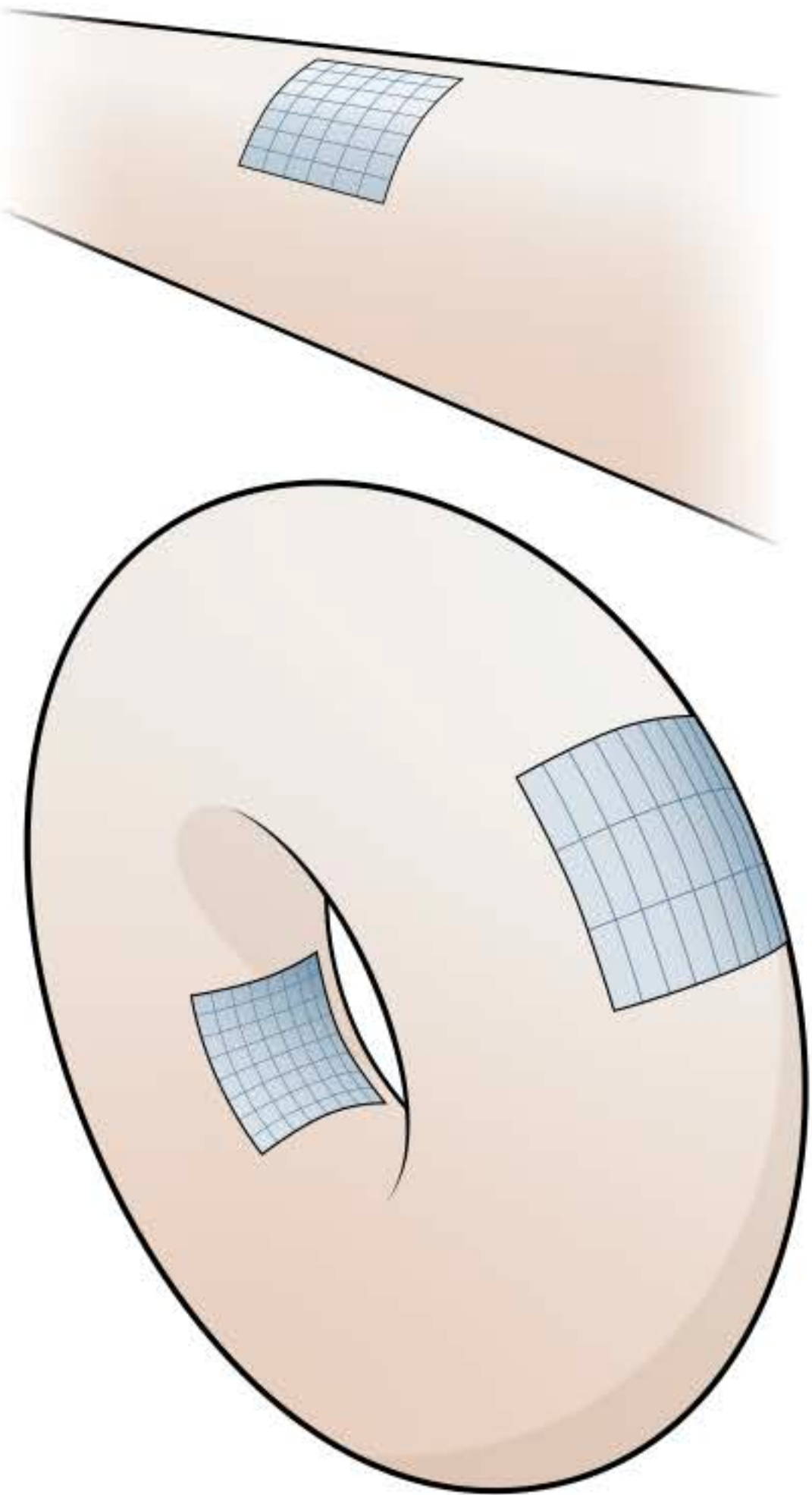
Points as Tori (PAT)



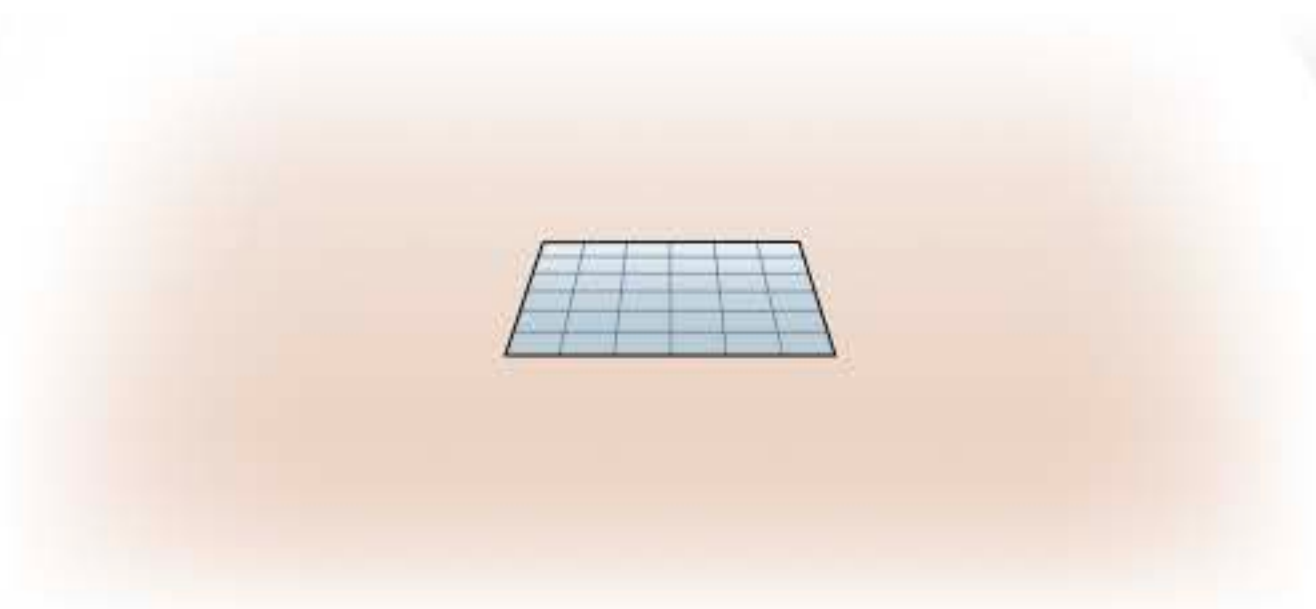
$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{i\text{-th function}}{g_i(x)} \overset{\text{weights}}{\exp(-\lambda \|x - p_i\|)}}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

kernel density estimator

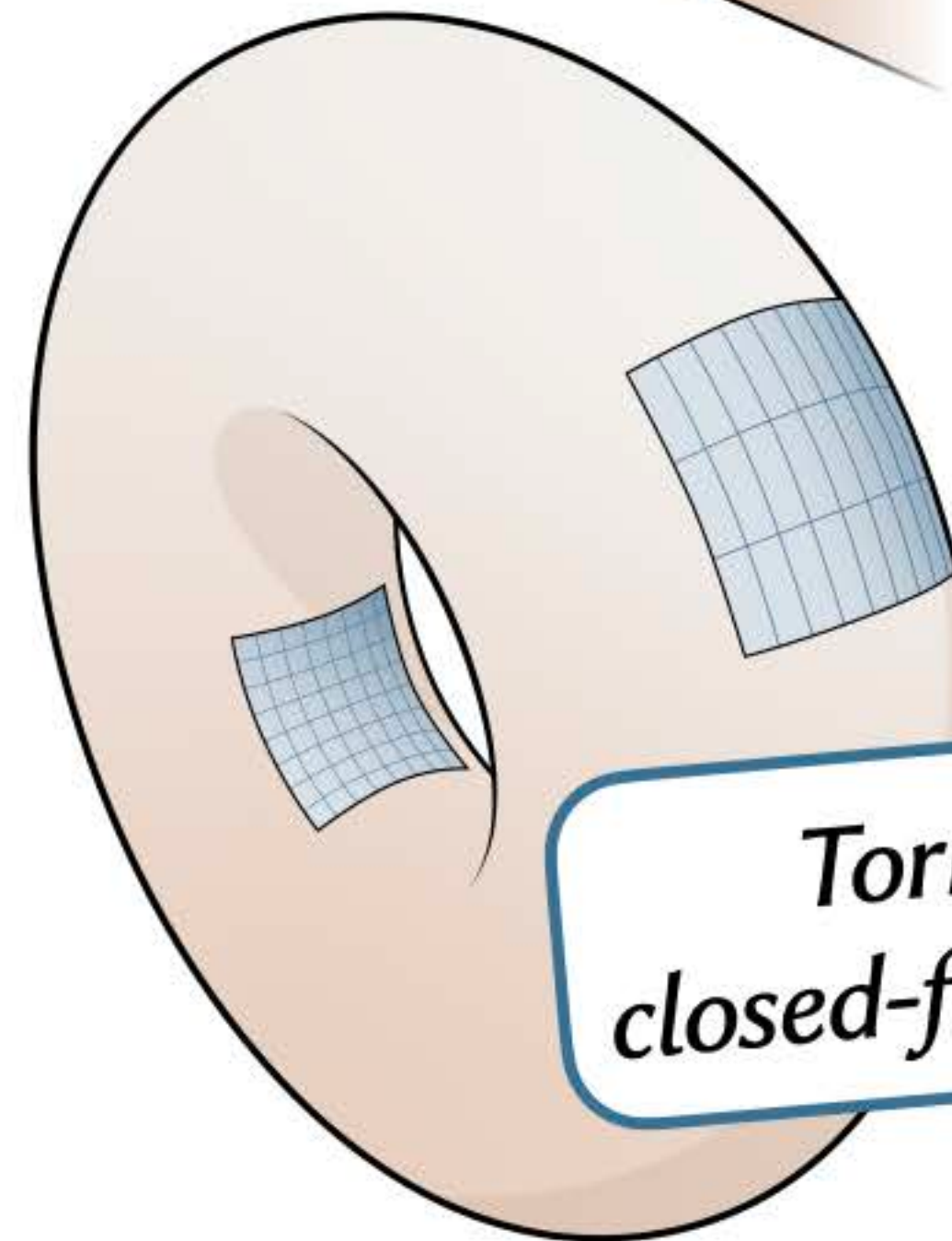
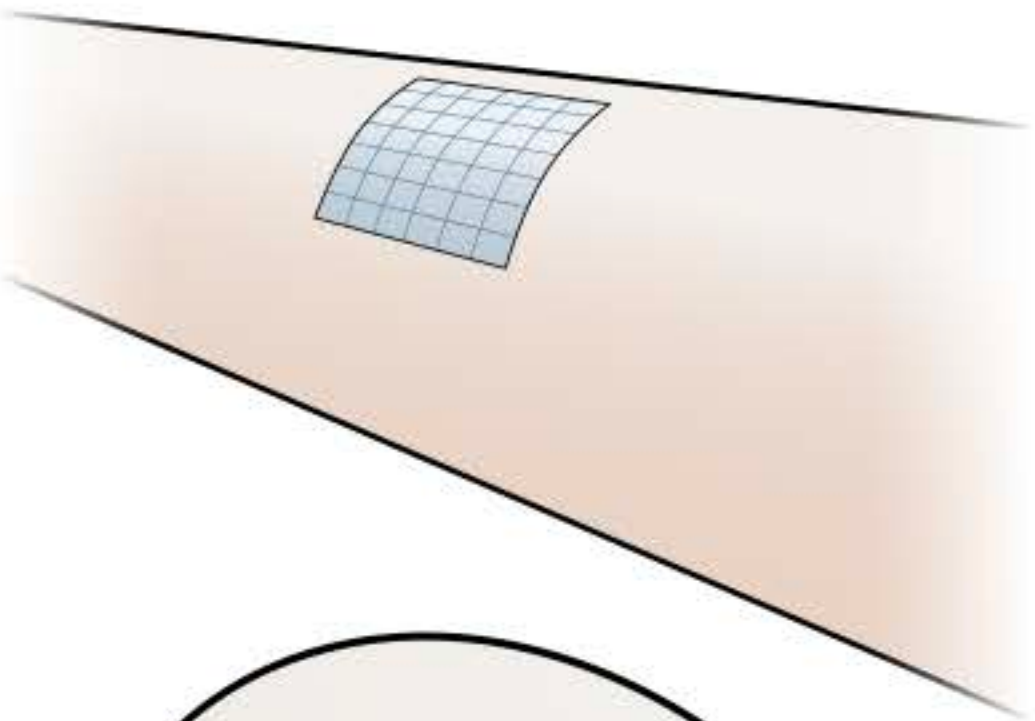
Points as Tori (PAT)



$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{\textit{i-th function}}}{g_i(x)} \overset{\text{\textit{weights}}}{\exp(-\lambda \|x - p_i\|)}}{\sum_{i=1}^{|P|} \underset{\text{\textit{kernel density estimator}}}{\exp(-\lambda \|x - p_i\|)}}$$



Points as Tori (PAT)



Tori have
closed-form SDFs!



$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{\textit{i-th function}}}{g_i(x)} \overset{\text{\textit{weights}}}{\exp(-\lambda \|x - p_i\|)}}{\sum_{i=1}^{|P|} \underset{\text{\textit{kernel density estimator}}}{\exp(-\lambda \|x - p_i\|)}}$$

Points as Tori (PAT)

global SDF

SDF of i -th torus

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

Tori have
closed-form SDFs!

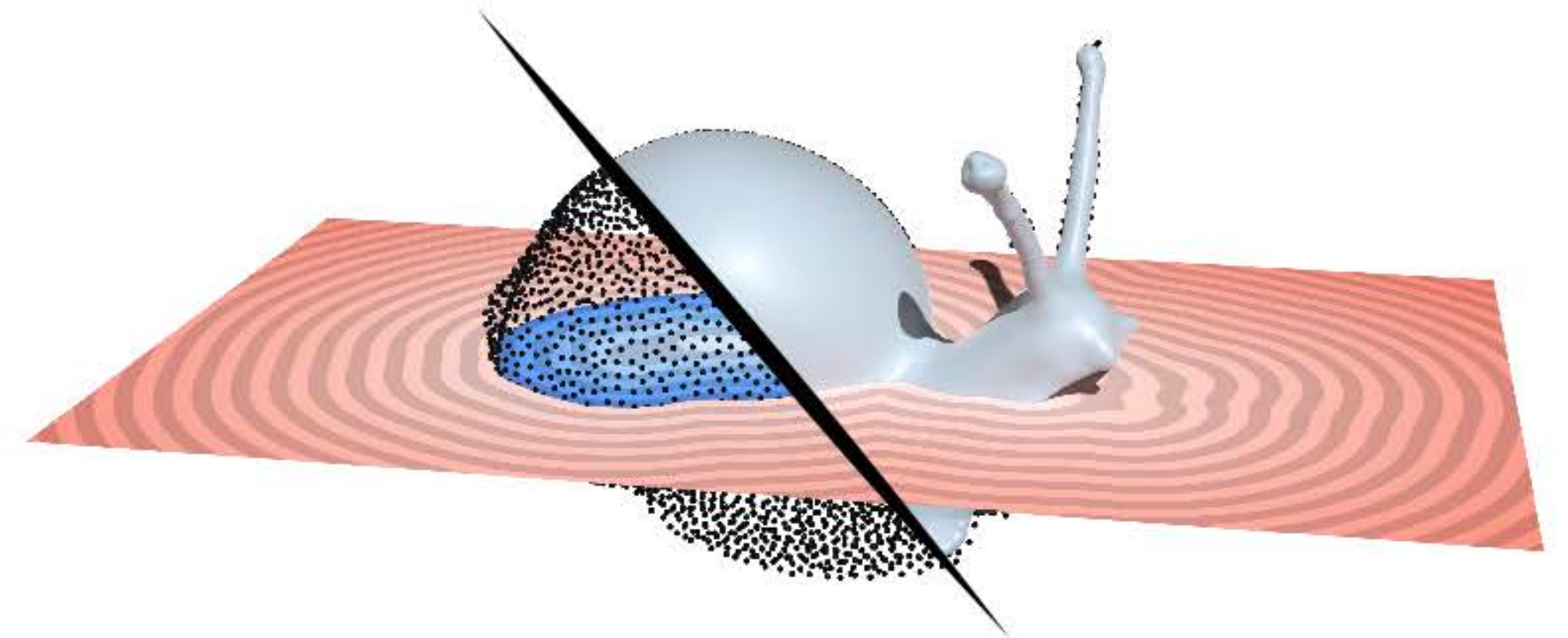
Points as Tori (PAT)

global SDF

SDF of i -th torus

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

Tori have
closed-form SDFs!



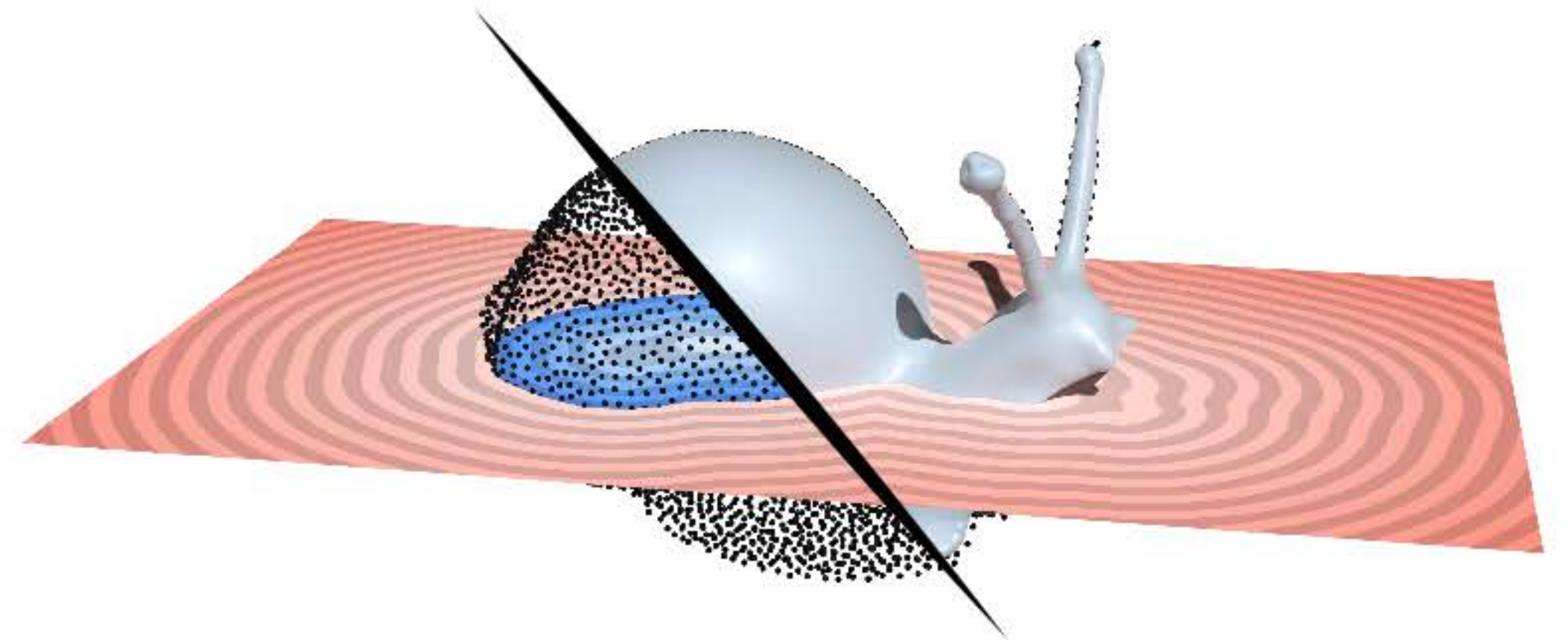
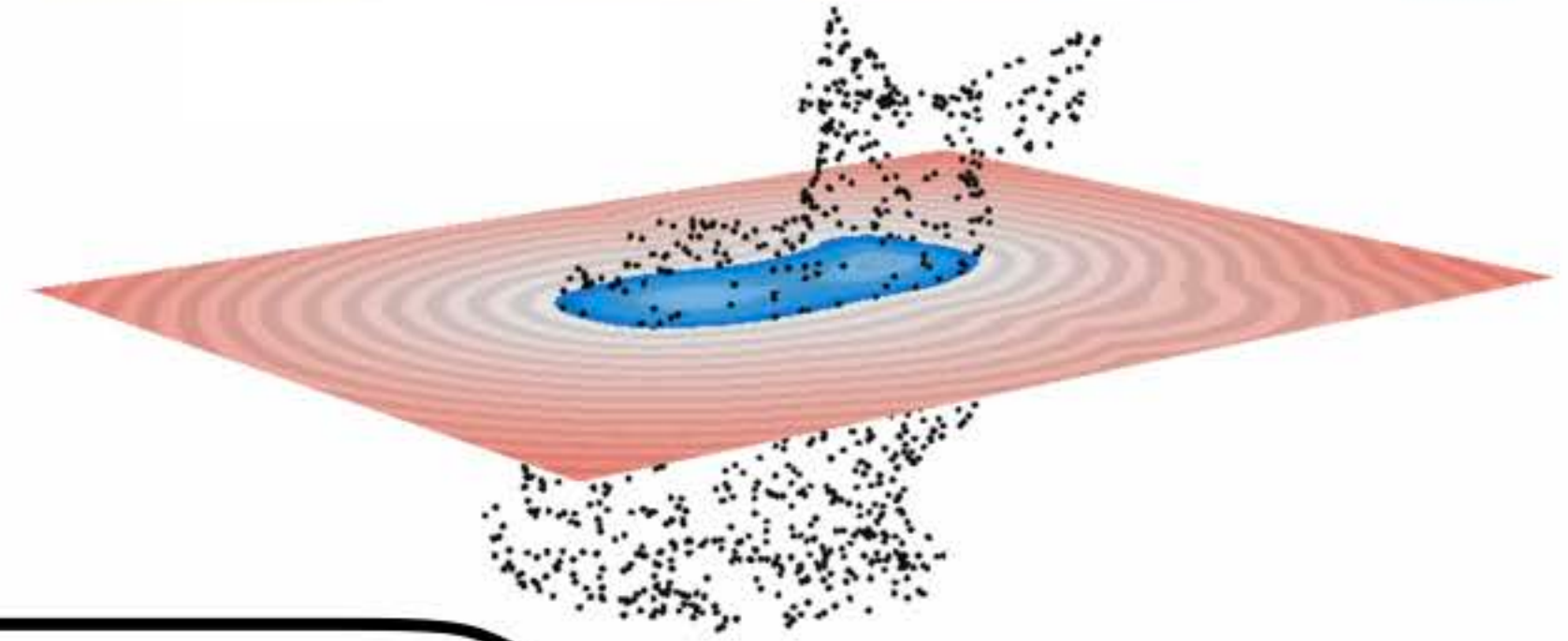
Points as Tori (PAT)

global SDF

SDF of i -th torus

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

Tori have
closed-form SDFs!



Points as Tori (PAT)

no global solves

global SDF

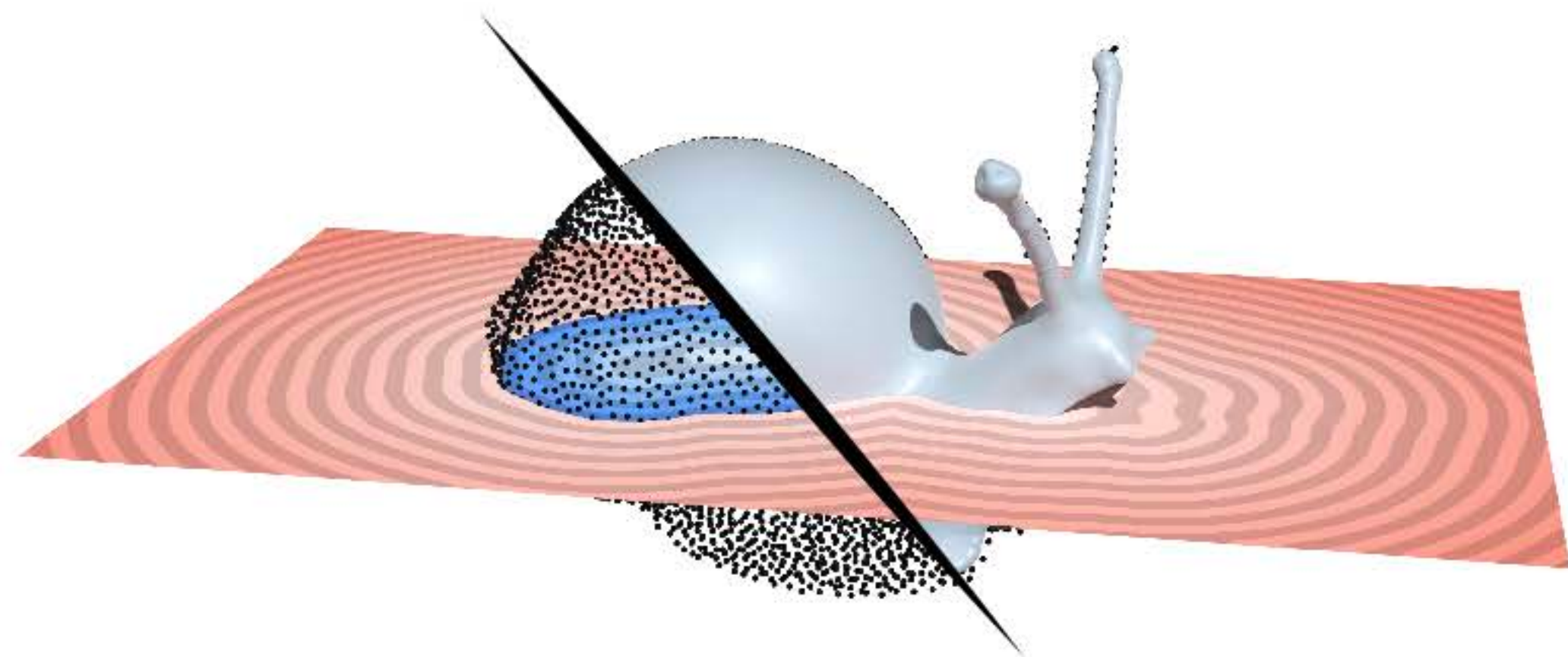
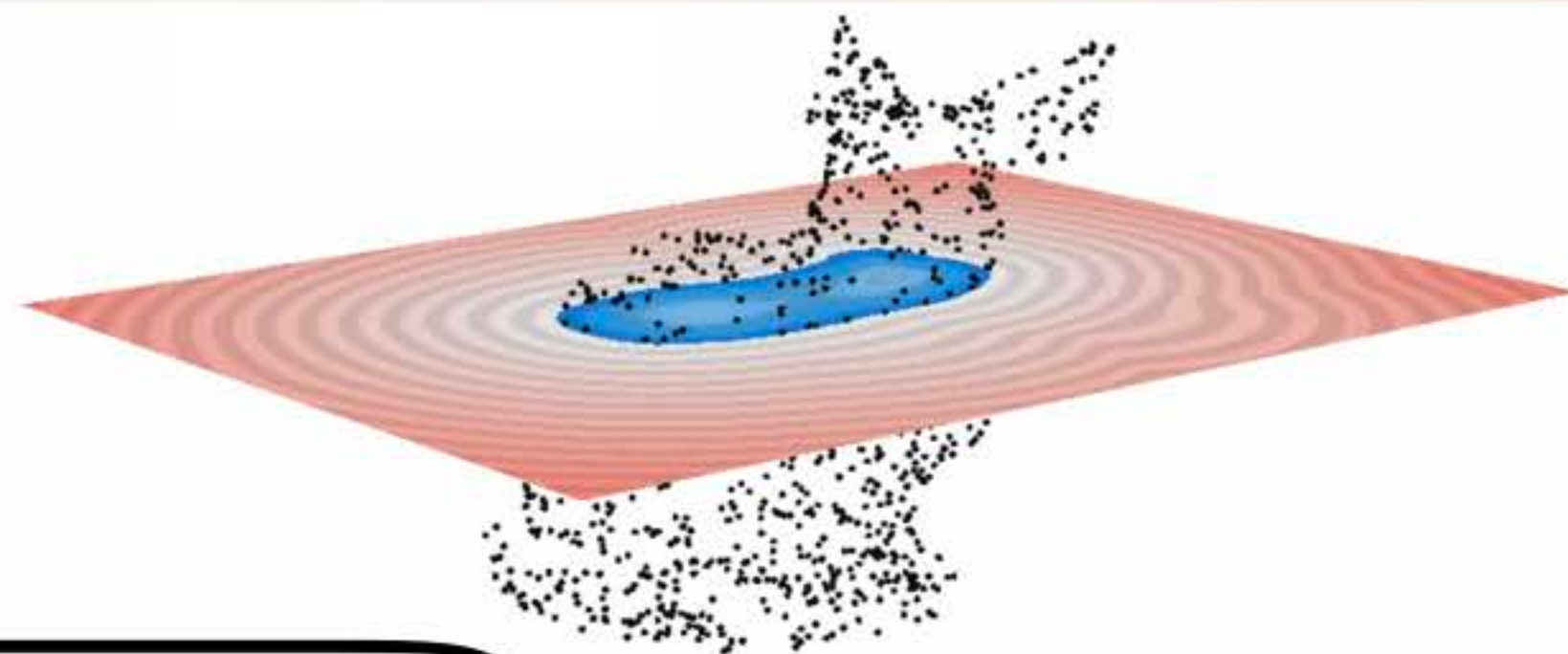
SDF of i -th torus

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

no discretization

pointwise
evaluation

Tori have
closed-form SDFs!



Two *convolutional formulas* for distance

Two *convolutional* formulas for distance

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

Formula #2:

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

Two *convolutional* formulas for distance

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

Formula #2:

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

*“convolutional
distance formulas”*

Two *convolutional* formulas for distance

Formula #1:

*summation over the
shape boundary Ω*

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

kernel

Formula #2:

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

summation over Ω

*“convolutional
distance formulas”*

Many special cases

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

Many special cases

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

Examples:

- *Varadhan's formula* [Varadhan 1967]
- *Hopf-Cole transformations* for distance [Belyaev & Fayolle 2015; Lipman 2021]
- *LogSumExp* distance [Madan & Levin 2022]
- *Kreisselmeier-Steinhauser* functions [Kreisselmeier & Steinhauser 1980]
- *Schrödinger distance transform* [Gurumoorthy & Rangarajan 2009]
- other instances [Karam et al. 2019, Abgrall 2022]
- discrete LogSumExp

Many special cases

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \overset{\text{exponential}}{\exp(-\lambda \|x - z\|)} dz \right]$$

Examples:

- *Varadhan's formula* [Varadhan 1967]
- *Hopf-Cole transformations for distance* [Belyaev & Fayolle 2015; Lipman 2021]
- *LogSumExp distance* [Madan & Levin 2022]
- *Kreisselmeier-Steinhauser functions* [Kreisselmeier & Steinhauser 1980]
- *Schrödinger distance transform* [Gurumoorthy & Rangarajan 2009]
- other instances [Karam et al. 2019, Abgrall 2022]
- discrete LogSumExp

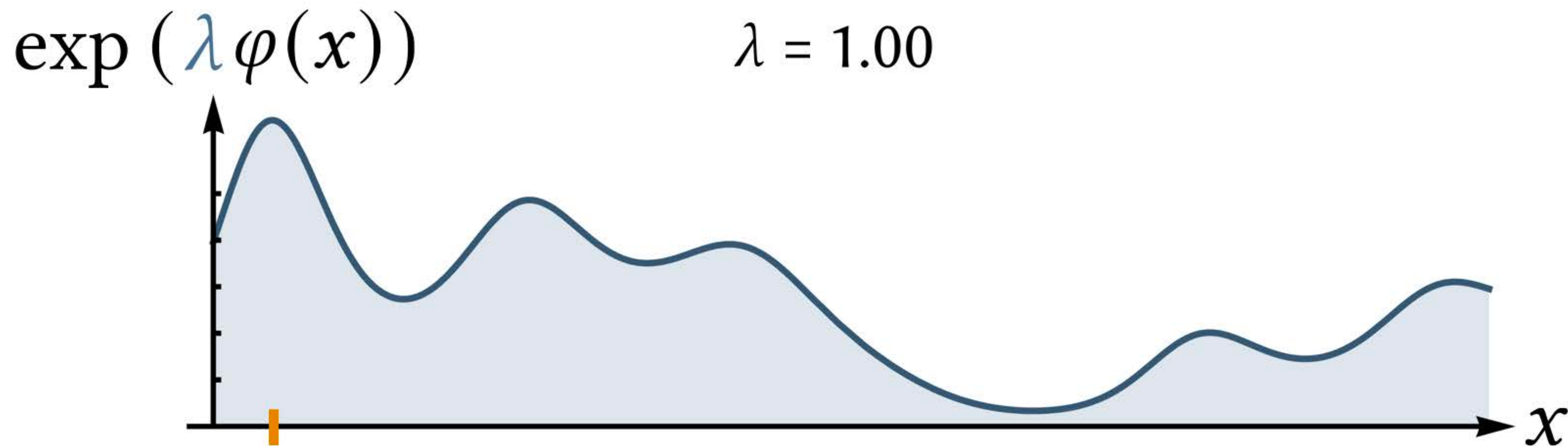
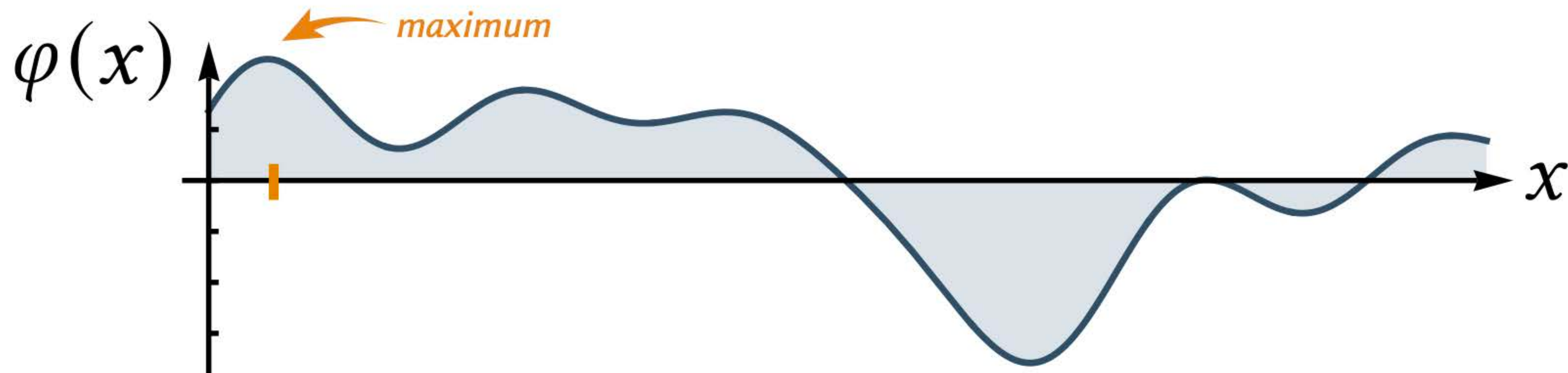
Consider the exponential of any function...

$$\varphi(x)$$

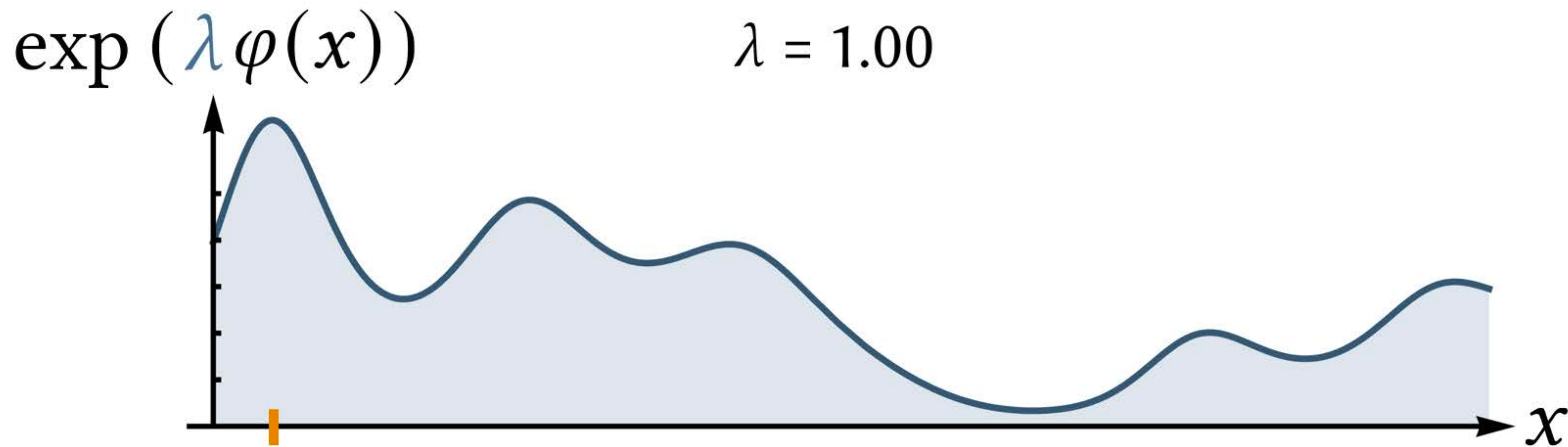
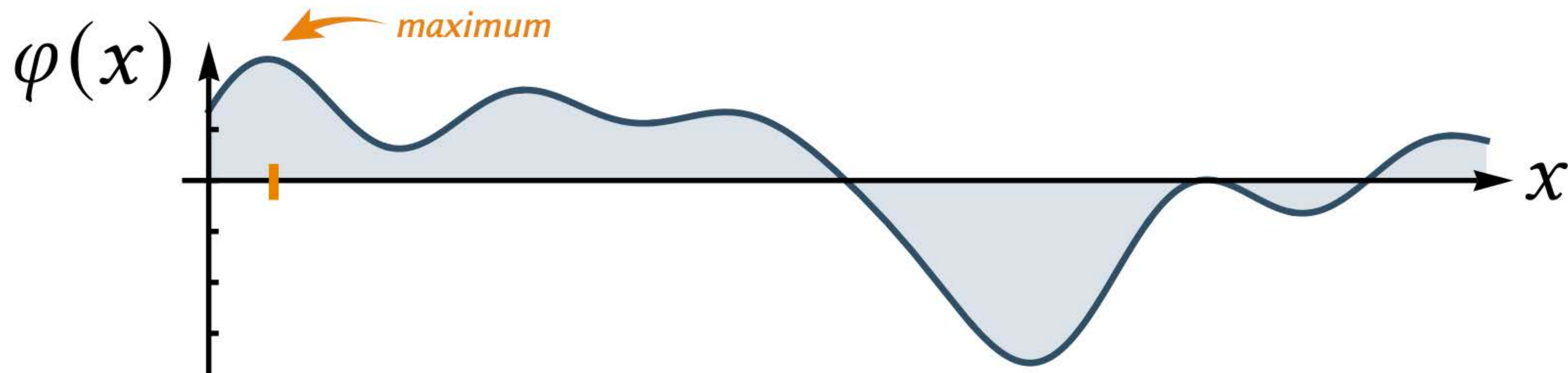
Consider the exponential of any function...

$$\exp (\varphi(x))$$

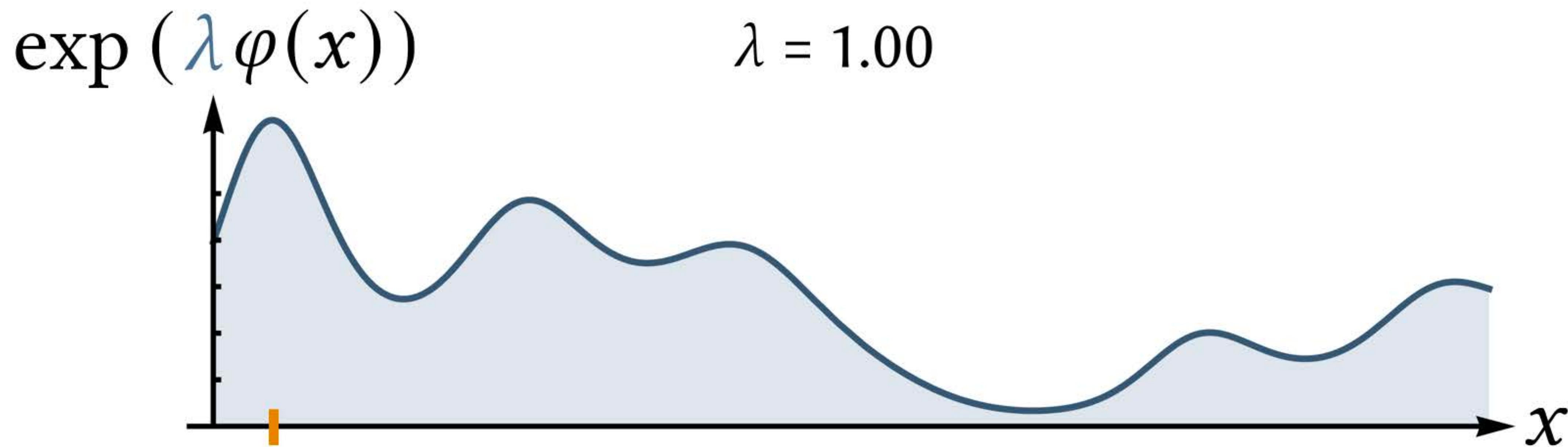
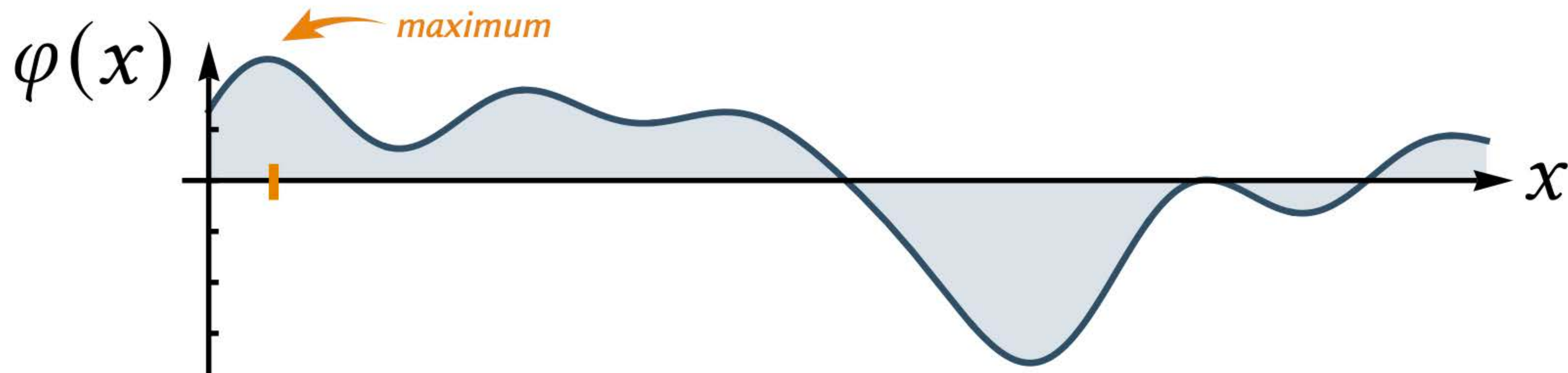
... the extremum value becomes exponentially more extreme



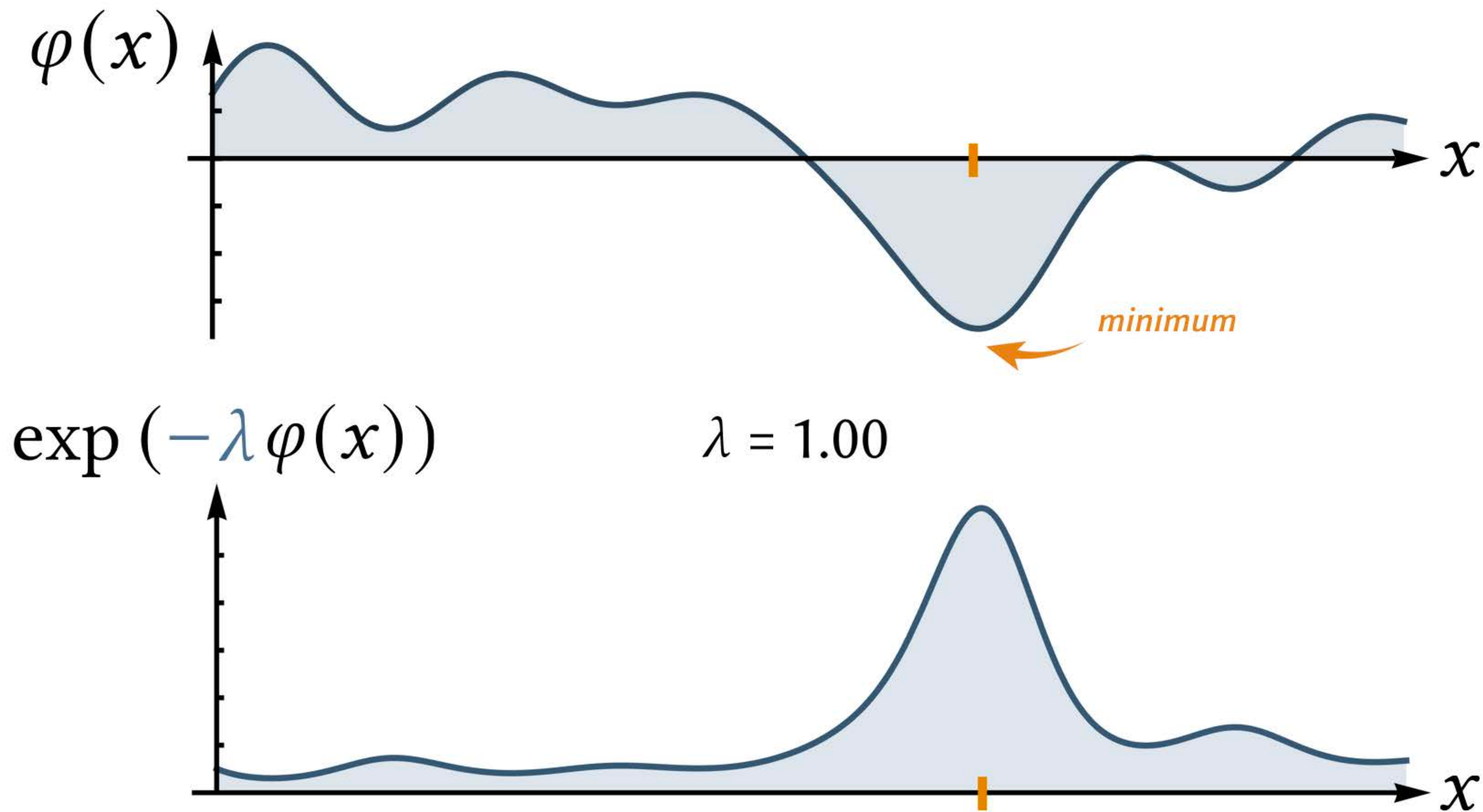
... the extremum value becomes exponentially more extreme



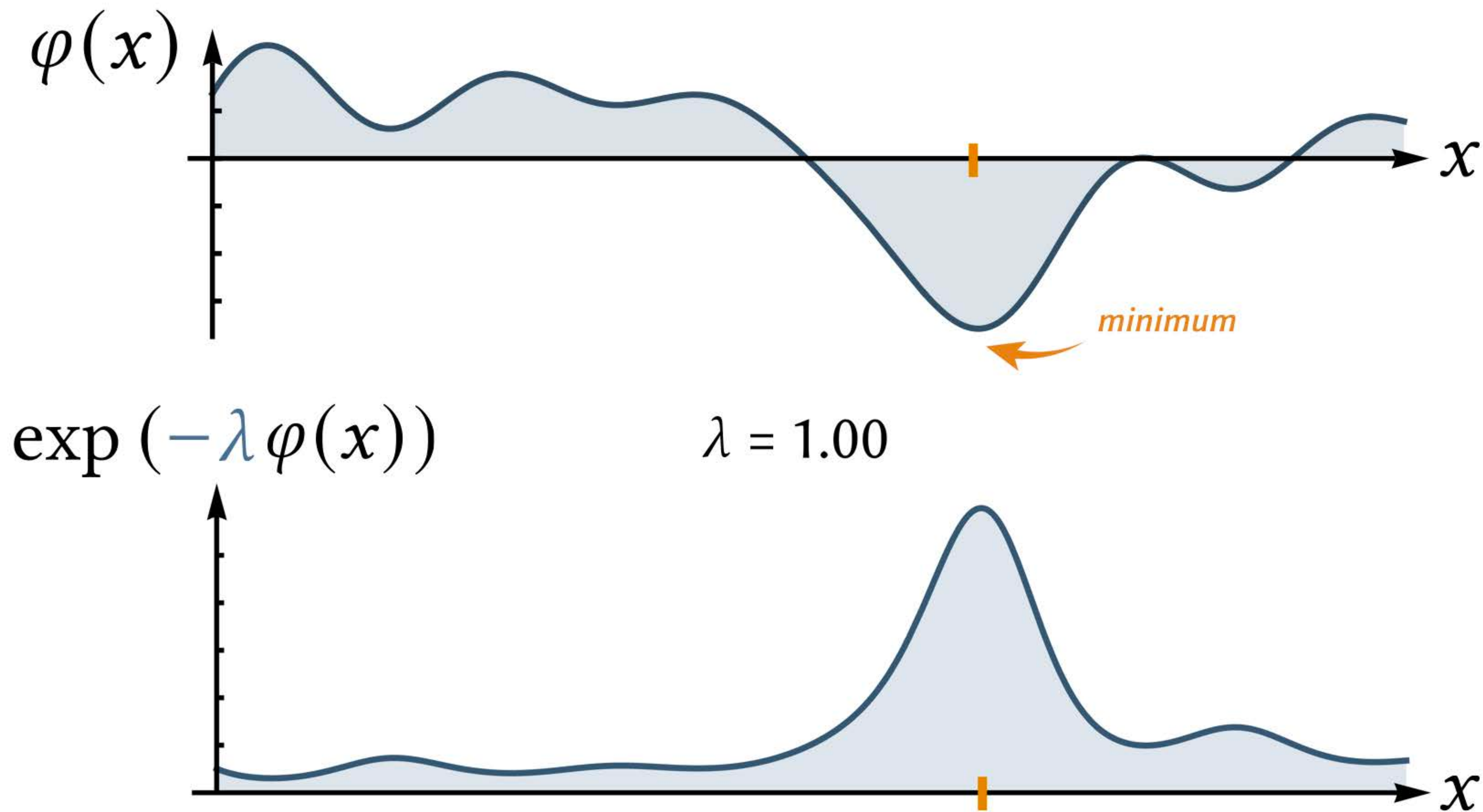
... the extremum value becomes exponentially more extreme



... the extremum value becomes exponentially more extreme



... the extremum value becomes exponentially more extreme



Asymptotic analysis

$$\int_{\Omega} h(z) \exp (-\lambda \varphi(z)) \mathrm{d} z$$

Asymptotic analysis

$$\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) \, dz \quad \stackrel{\lambda \rightarrow \infty}{\sim} \quad (2\pi/\lambda)^{d/2} \det(\nabla^2 \varphi(x^*))^{-1/2} h(x^*) \exp(-\lambda \varphi(x^*))$$

Laplace's method says...

... only the contribution at $x^ = \operatorname{argmin}_{z \in \Omega} \varphi(z)$ matters.*

Asymptotic analysis

$$\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) \, dz \quad \stackrel{\lambda \rightarrow \infty}{\sim} \quad (2\pi/\lambda)^{d/2} \det(\nabla^2 \varphi(x^*))^{-1/2} h(x^*) \exp(-\lambda \varphi(x^*))$$

Laplace's method says...

... only the contribution at $x^ = \operatorname{argmin}_{z \in \Omega} \varphi(z)$ matters.*

Applying $-\frac{1}{\lambda} \log(\cdot)$ to both sides,

$$-\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) \, dz \right] \sim \varphi(x^*) + O(\lambda^{-1}) \quad \text{as } \lambda \rightarrow \infty$$

Asymptotic analysis

$$\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) \, dz \quad \stackrel{\lambda \rightarrow \infty}{\sim} \quad (2\pi/\lambda)^{d/2} \det(\nabla^2 \varphi(x^*))^{-1/2} h(x^*) \exp(-\lambda \varphi(x^*))$$

Laplace's method says...

... only the contribution at $x^ = \operatorname{argmin}_{z \in \Omega} \varphi(z)$ matters.*

Applying $-\frac{1}{\lambda} \log(\cdot)$ to both sides,

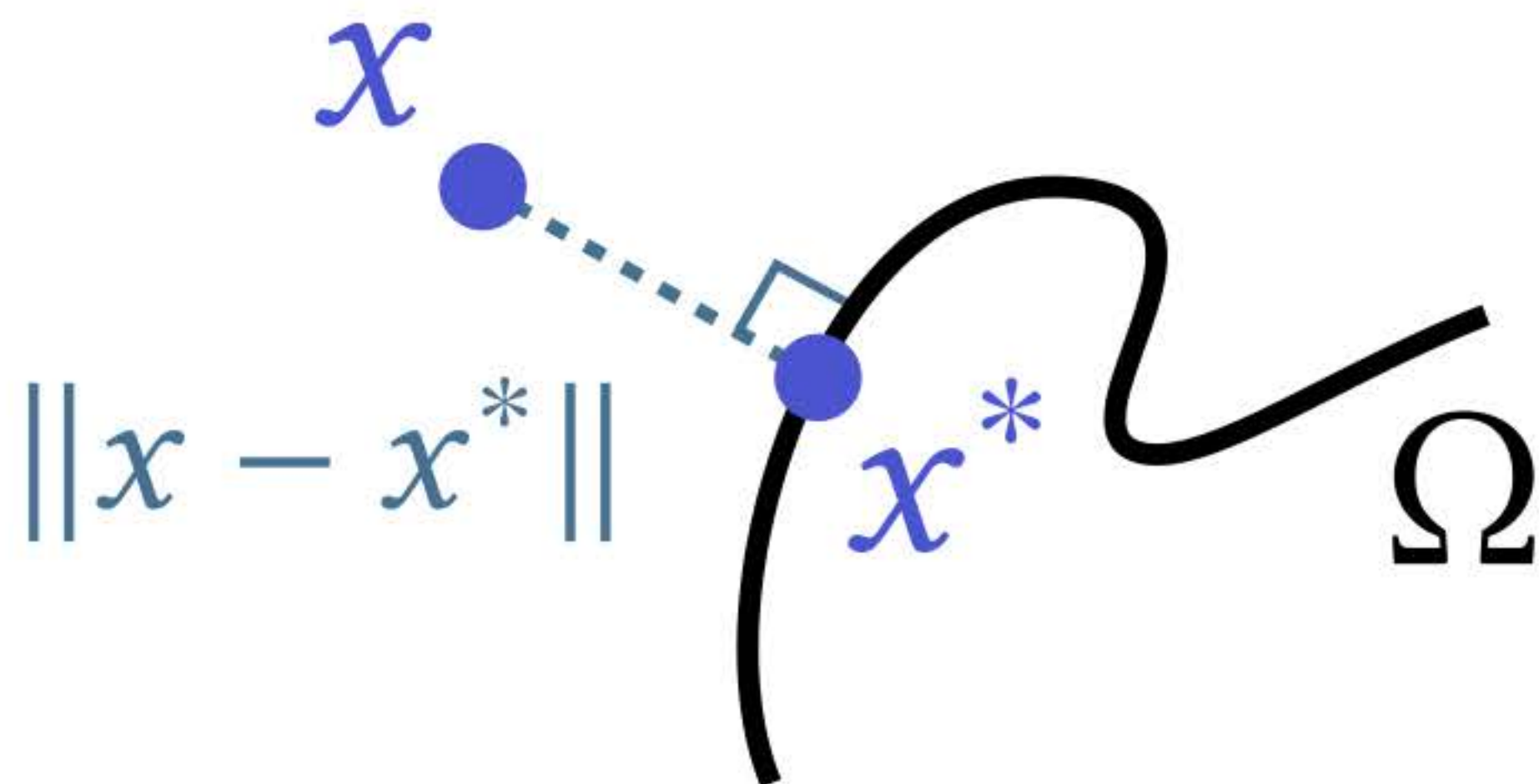
$$-\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) \, dz \right] \sim \underbrace{\varphi(x^*)}_{\text{minimum of } \varphi} + O(\lambda^{-1}) \quad \text{as } \lambda \rightarrow \infty$$

Deriving convolutional distance

Use $\varphi_x(z) = \|x - z\|$ to get $\|x - x^*\| = \min_{z \in \Omega} \|x - z\|$ as $\lambda \rightarrow \infty$.

distance

distance from x to closest point in Ω



$$-\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) dz \right] \sim \varphi(x^*) + O(\lambda^{-1}) \quad \text{as } \lambda \rightarrow \infty$$

A general convolutional formula for distance

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) \, dz \right]$$

where $\varphi_x(z) = \|x - z\|$.

A general convolutional formula for distance

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \varphi(z)) dz \right]$$

where $\varphi_x(z) = \|x - z\|$.

Examples:

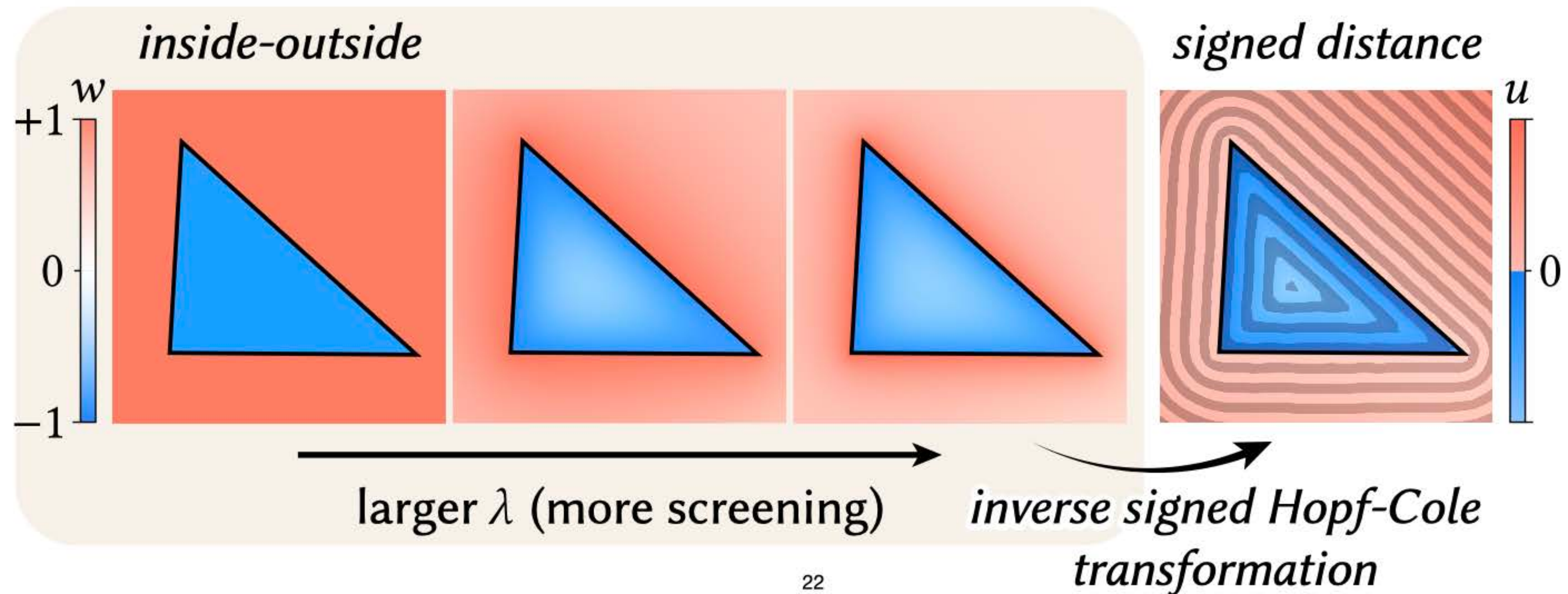
- *Varadhan's formula* [Varadhan 1967]
- *Hopf-Cole transformations for distance* [Belyaev & Fayolle 2015; Lipman 2021]
- *LogSumExp distance* [Madan & Levin 2022]
- *Kreisselmeier-Steinhauser functions* [Kreisselmeier & Steinhauser 1980]
- *Schrödinger distance transform* [Gurumoorthy & Rangarajan 2009]
- other instances [Karam et al. 2019, Abgrall 2022]
- discrete LogSumExp

We unify signed distance with reconstruction

We unify signed distance with reconstruction

eikonal equation $\rightarrow$ *jump screened Laplace equation* $\rightarrow$ *screened Poisson equation* $\rightarrow$ *boundary integral*

$$\begin{array}{ll}
 \|\nabla u(x)\|^2 &= 1, \quad x \notin \Omega \\
 u(x) &= 0 \quad x \in \Omega \\
 \frac{\partial u}{\partial n}(x) &= 1 \quad x \in \Omega
 \end{array}
 \quad
 \begin{array}{ll}
 \Delta w(x) - \lambda^2 w(x) &= 0, \quad x \notin \Omega \\
 w^\pm(x) &= \pm 1 \quad x \in \Omega \\
 \frac{\partial w^+}{\partial n}(x) - \frac{\partial w^-}{\partial n}(x) &= 0 \quad x \in \Omega
 \end{array}
 \quad
 \begin{array}{ll}
 \Delta w(x) - \lambda^2 w(x) &= -2(\nabla \cdot n(x))\mu_\Omega(x) \\
 w(x) &= 2 \int_\Omega \frac{\partial G^\lambda(x, z)}{\partial n(z)} dz
 \end{array}$$

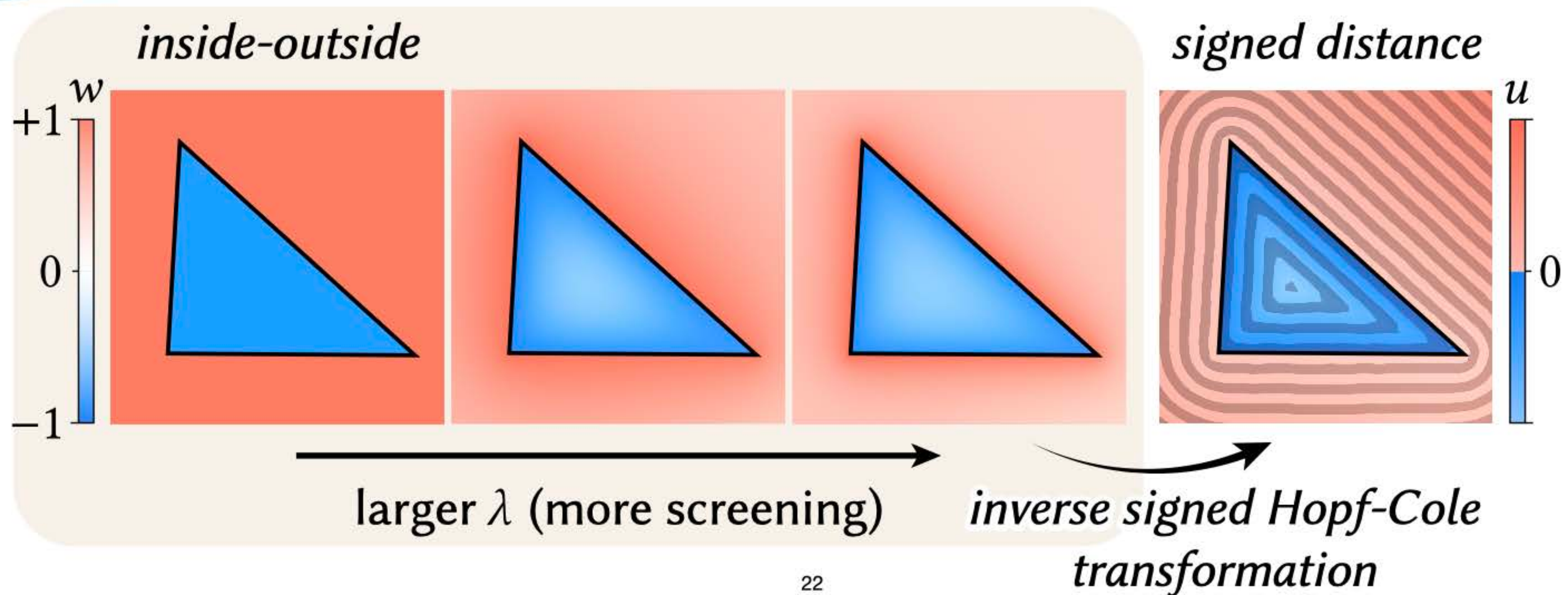


We unify signed distance with reconstruction

eikonal equation $\rightarrow$ jump screened Laplace equation $\rightarrow$ screened Poisson equation $\rightarrow$ boundary integral

$$\begin{aligned} \|\nabla u(x)\|^2 &= 1, & x \notin \Omega & & \Delta w(x) - \lambda^2 w(x) &= 0, & x \notin \Omega & & \Delta w(x) - \lambda^2 w(x) &= -2(\nabla \cdot n(x))\mu_\Omega(x) & & w(x) &= 2 \int_\Omega \frac{\partial G^\lambda(x, z)}{\partial n(z)} dz \\ u(x) &= 0 & x \in \Omega & & w^\pm(x) &= \pm 1 & x \in \Omega & & & & & & & \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega & & & & & & & & & & & \end{aligned}$$

$\lambda \rightarrow 0$: winding numbers
(jump harmonic functions) [Feng et al. 2023]



We unify signed distance with reconstruction

eikonal equation $\rightarrow$ jump screened Laplace equation $\rightarrow$ Poisson equation $\rightarrow$ boundary integral

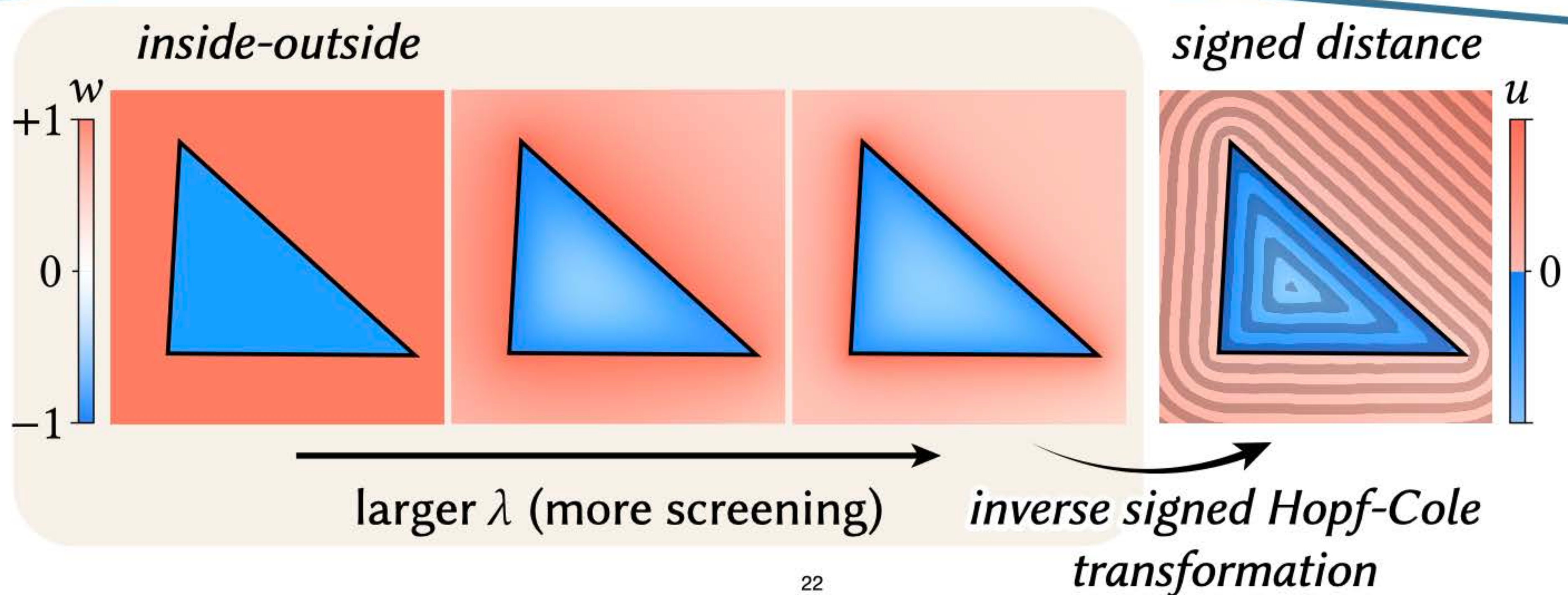
$$\begin{aligned} \|\nabla u(x)\|^2 &= 1, & x \notin \Omega & \\ u(x) &= 0 & x \in \Omega & \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega & \end{aligned} \quad \begin{aligned} \Delta w(x) - \lambda^2 w(x) &= 0, & x \notin \Omega & \\ w^\pm(x) &= \pm 1 & x \in \Omega & \end{aligned}$$

$\lambda \rightarrow 0$: winding numbers
(jump harmonic functions) [Feng et al. 2023]

Poisson Surface Reconstruction
[Kazhdan et al. 2006]

$$\frac{\partial G^\lambda(x, z)}{\partial n(z)} dz$$

regularized winding number
[Chen et al. 2024]



We unify signed distance with reconstruction

eikonal equation $\rightarrow$ jump screened Laplace equation $\rightarrow$ Poisson equation $\rightarrow$ boundary integral

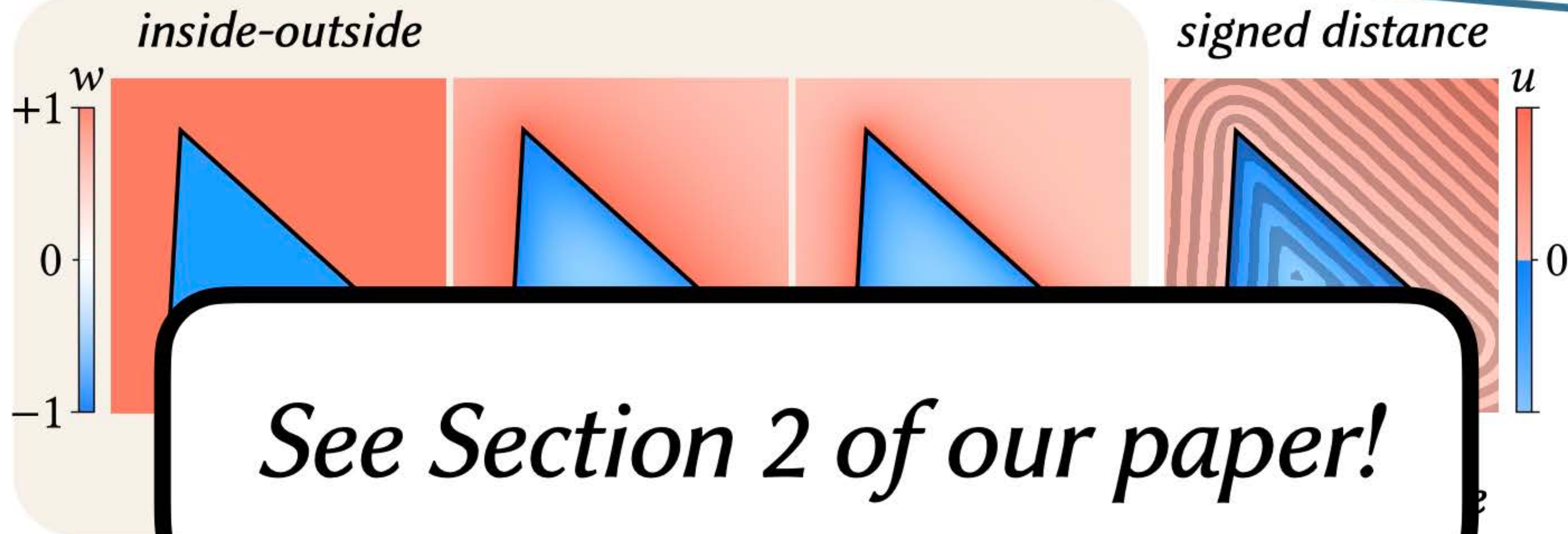
$$\begin{aligned} \|\nabla u(x)\|^2 &= 1, & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega \end{aligned} \quad \begin{aligned} \Delta w(x) - \lambda^2 w(x) &= 0, & x \notin \Omega \\ w^\pm(x) &= \pm 1 & x \in \Omega \end{aligned}$$

$\lambda \rightarrow 0$: winding numbers
(jump harmonic functions) [Feng et al. 2023]

Poisson Surface Reconstruction
[Kazhdan et al. 2006]

$$\frac{\partial G^\lambda(x, z)}{\partial n(z)} dz$$

regularized winding number
[Chen et al. 2024]



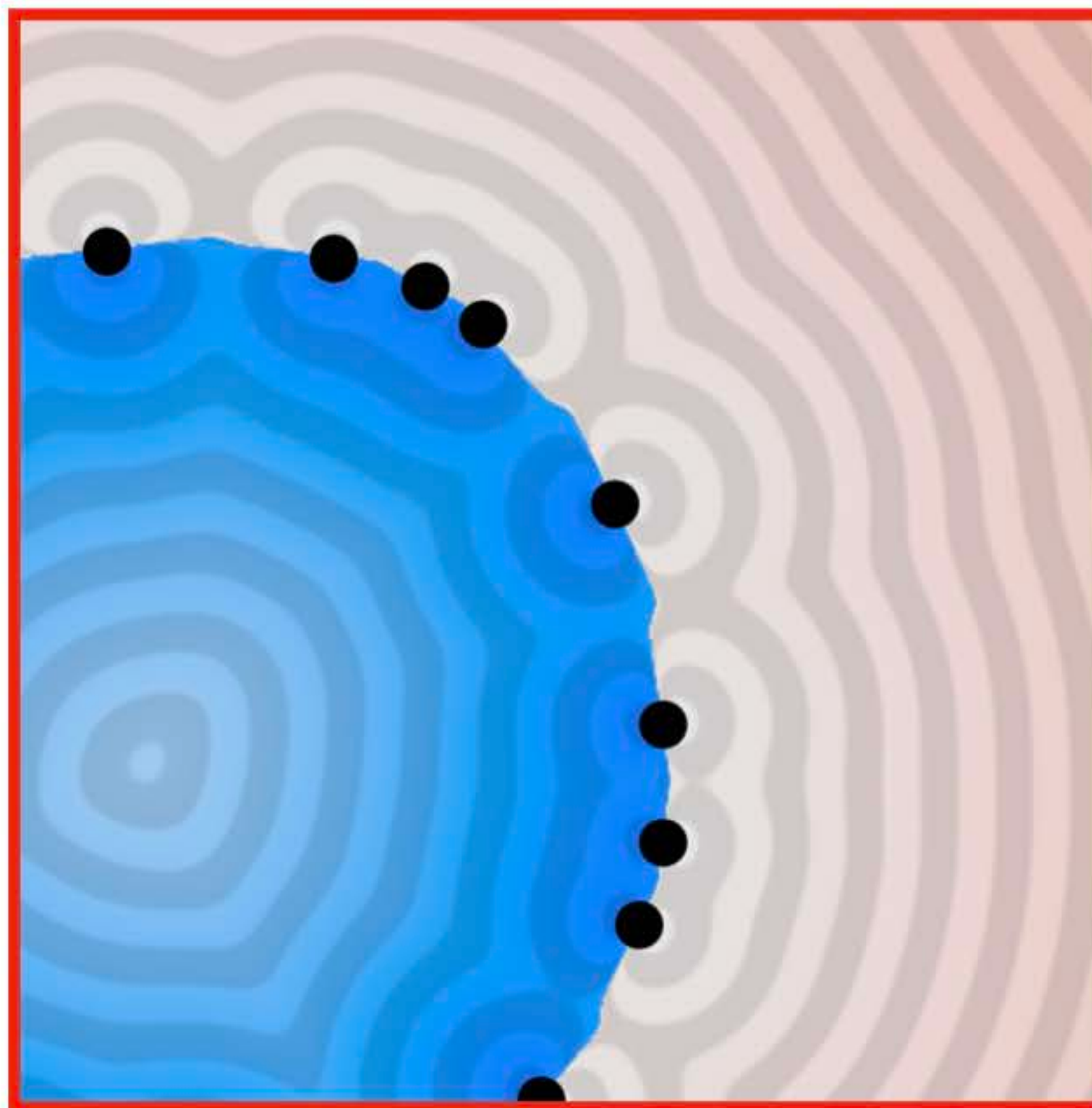
See Section 2 of our paper!

Convolutional formulas break down on real data

input point cloud P



signed distance approximation

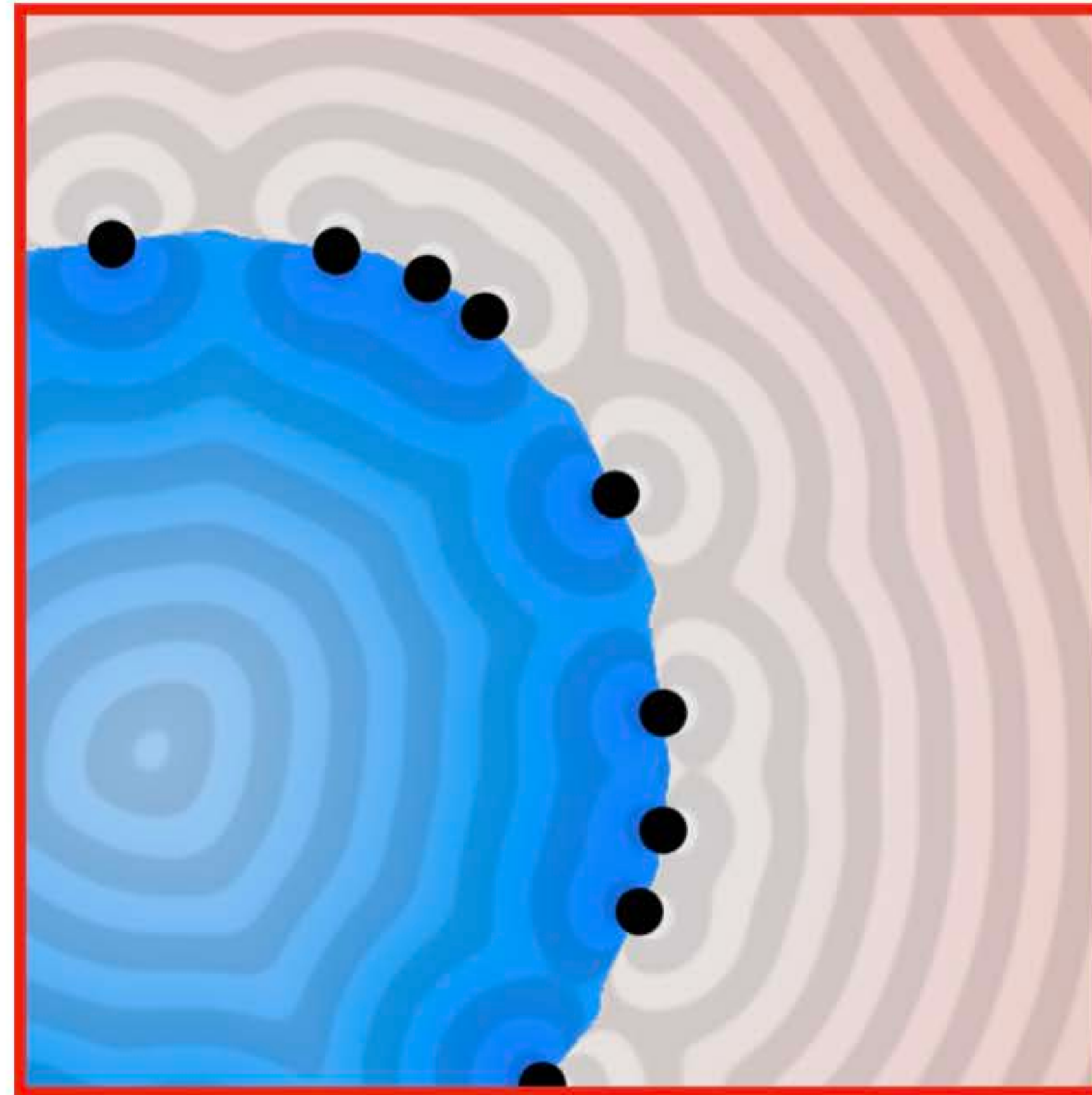


Convolutional formulas break down on real data

input point cloud P



signed distance approximation



We obtain distance to the *sampled* geometry, no matter what!

The exponential is (too) strong

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) \, dz \right] \sim \|x - x^*\|$$
$$x^* := \operatorname{argmin}_{z \in \Omega} \|x - z\|$$

The exponential is (too) strong

essentially a single-point
approximation

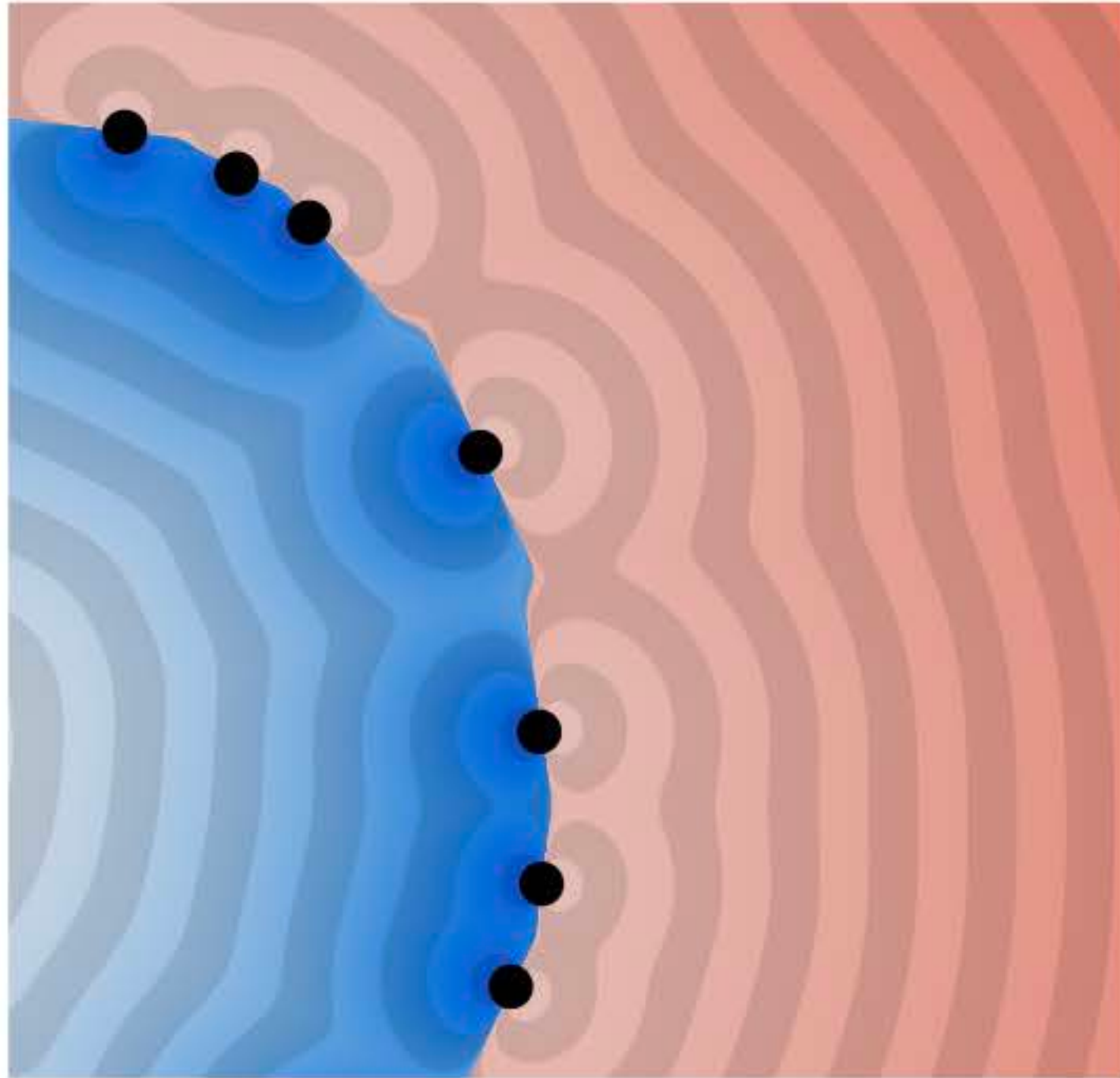
$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right] \sim \|x - x^*\| \text{ as } \lambda \rightarrow \infty$$

$x^* := \operatorname{argmin}_{z \in \Omega} \|x - z\|$

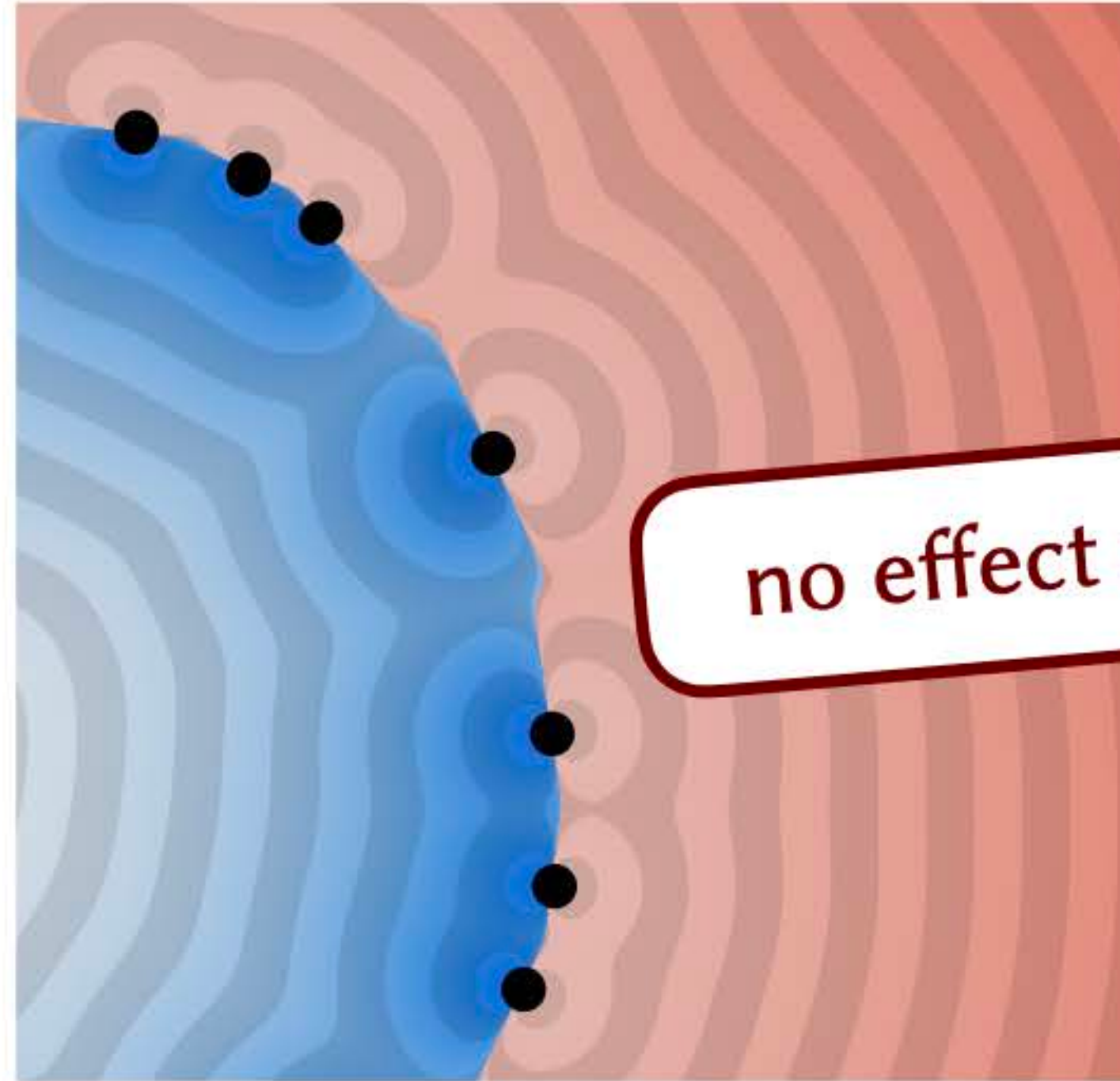
Regularization is futile for distance computation

regularization $\equiv$ choosing h

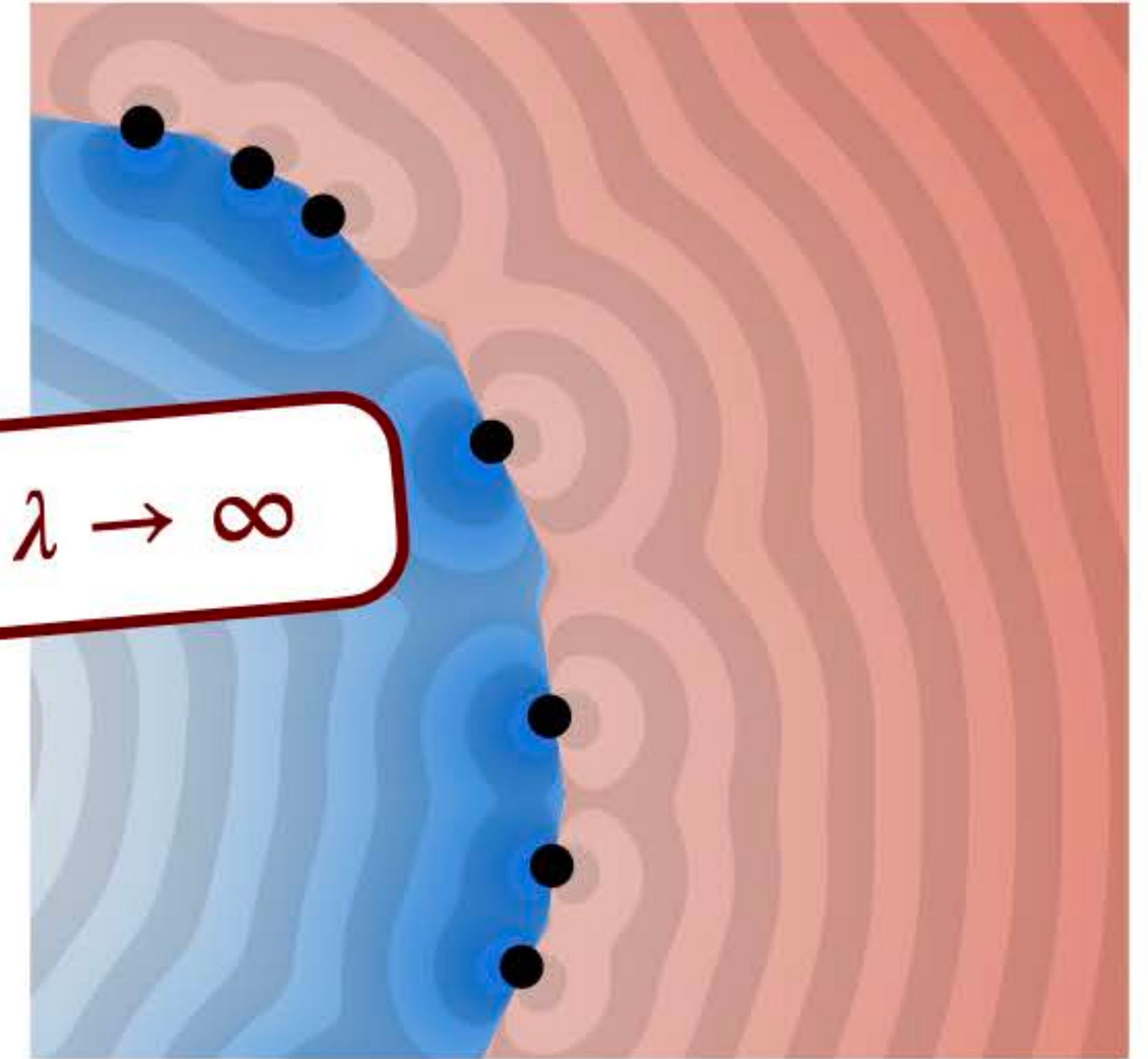
$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right] \sim \|x - x^*\| \quad \text{as } \lambda \rightarrow \infty$$



no regularization



isotropic regularization



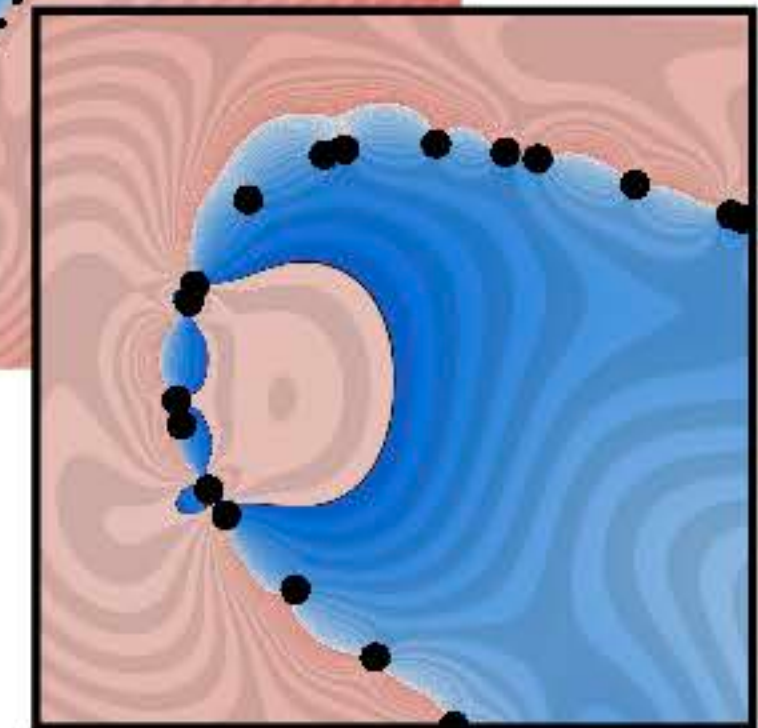
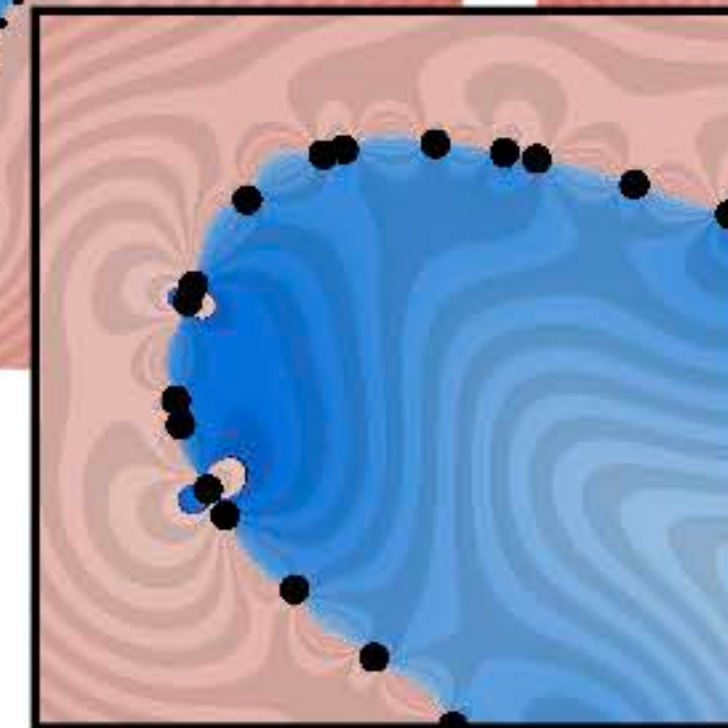
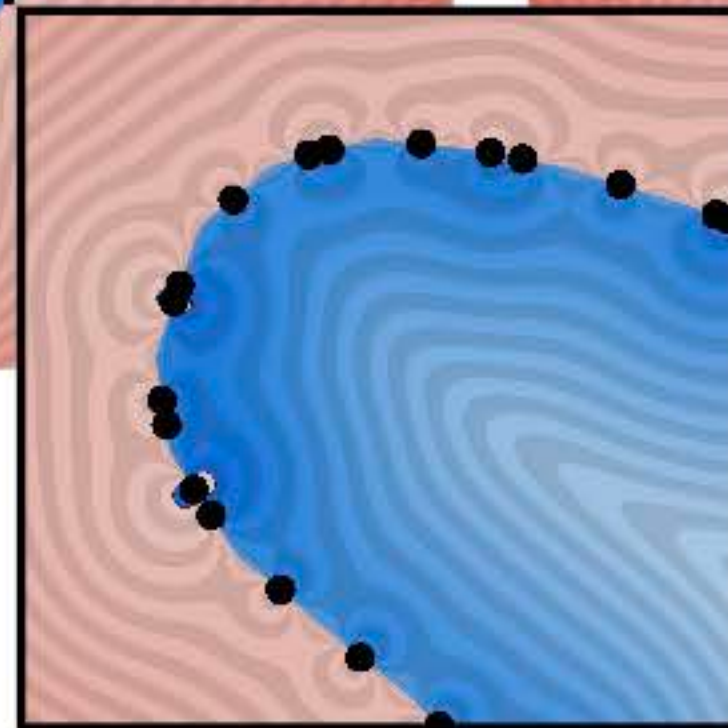
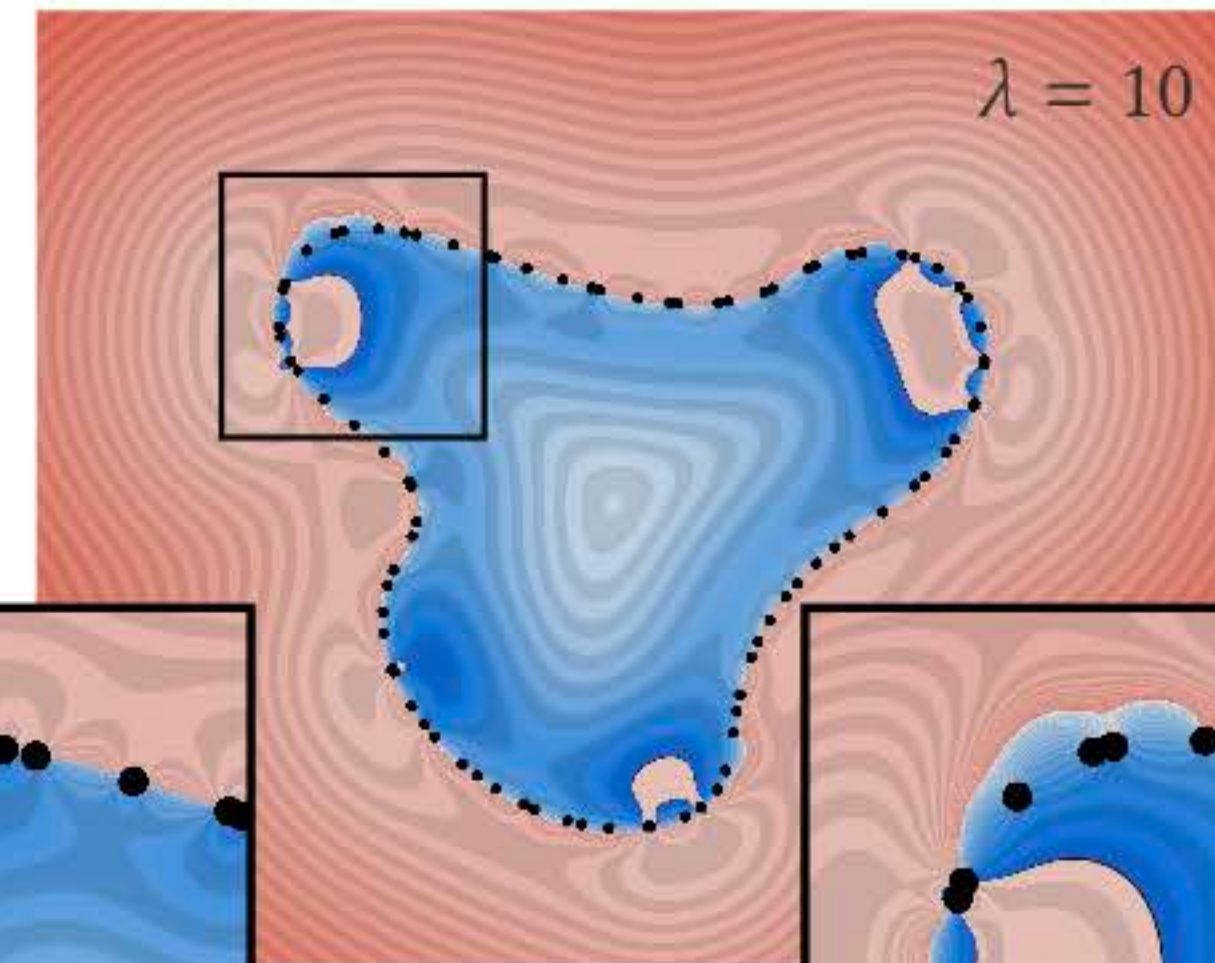
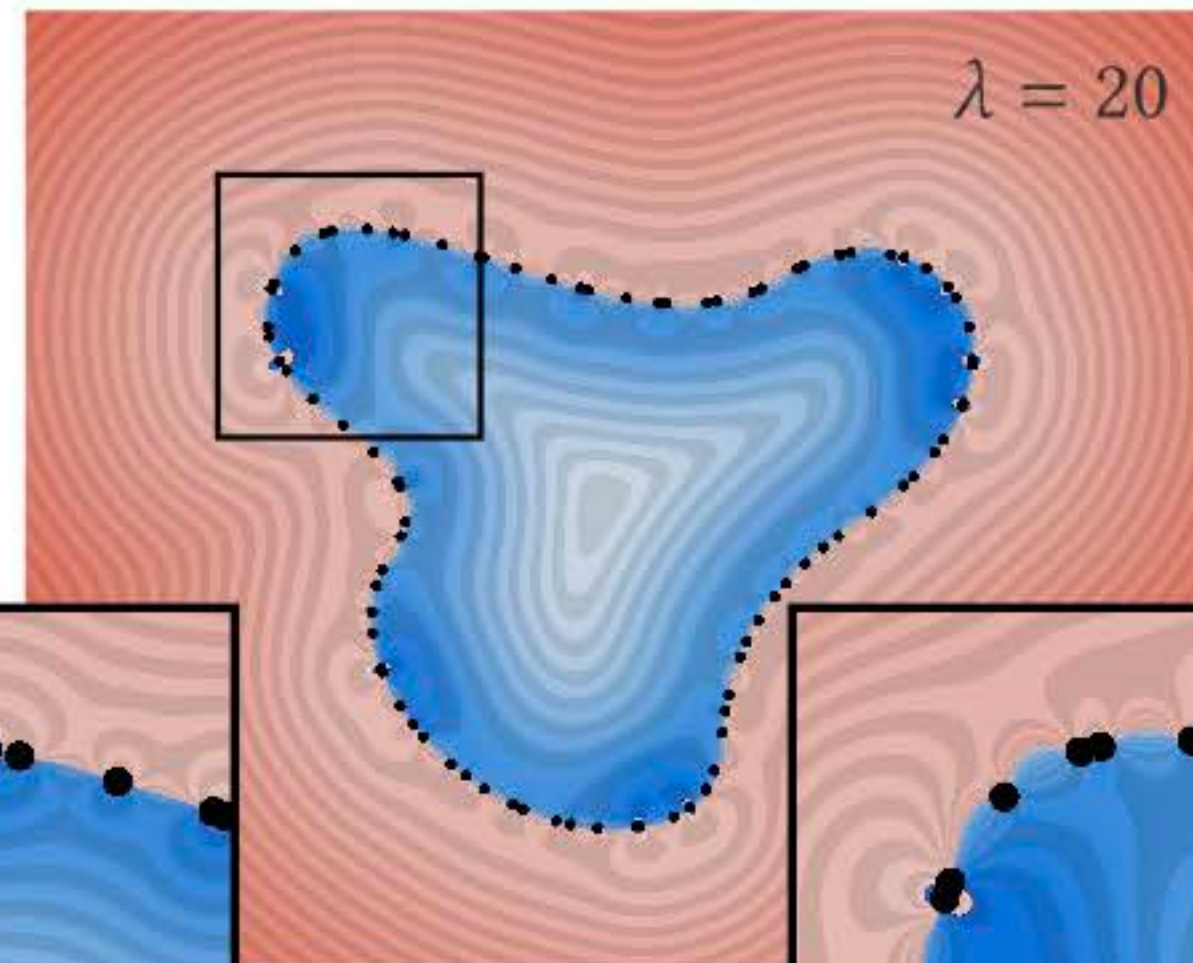
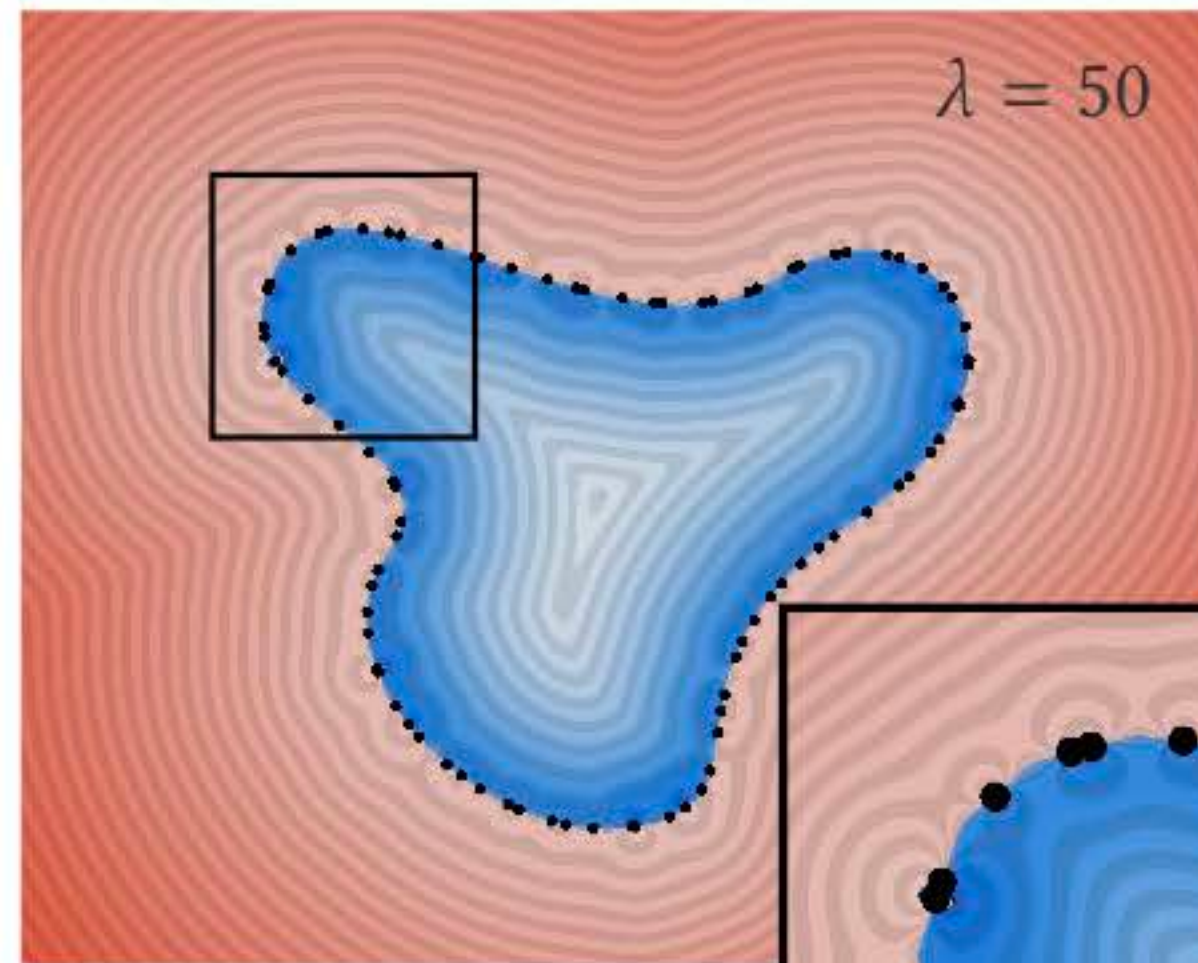
anisotropic regularization

Regression is coupled to distance — can't have both

larger λ
*better eikonality,
overfit level sets*

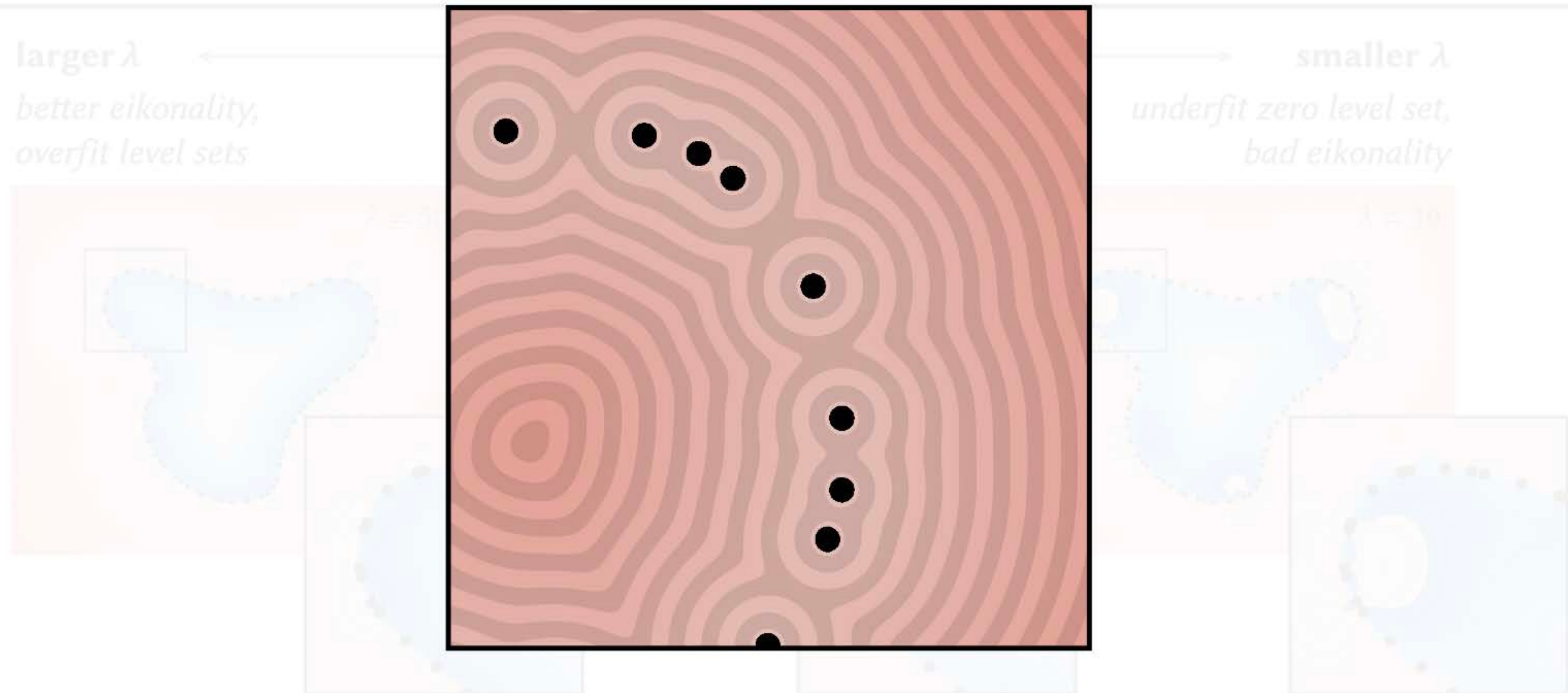
No happy medium

smaller λ
*underfit zero level set,
bad eikonality*



The parameter λ is coupled to both regression quality and distance accuracy, in opposite directions.

Past distance formulas fail for point clouds



[Varadhan 1967, Kreisselmeier & Steinhauser 1980, Gurumoorthy & Rangarajan 2009, Karam et al. 2019, Abgrall 2022, Madan & Levin 2022]

A second convolutional distance formula

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

A second convolutional distance formula

Formula #1:

$$\tilde{d}(x) := -\frac{1}{\lambda} \log \left[\int_{\Omega} h(z) \exp(-\lambda \|x - z\|) dz \right]$$

Formula #2:

a.k.a. kernel density estimator

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

A second convolutional distance formula

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz} \quad \lambda \xrightarrow{\sim} \infty \quad g(x^*)$$

$x^* := \operatorname{argmin}_{z \in \Omega} \|x - z\|$

A second convolutional distance formula

closest-point interpolation

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

$$\lambda \rightarrow \infty$$
$$\sim$$

$$g(x^*)$$

$$x^* := \operatorname{argmin}_{z \in \Omega} \|x - z\|$$

A second convolutional distance formula

closest-point interpolation

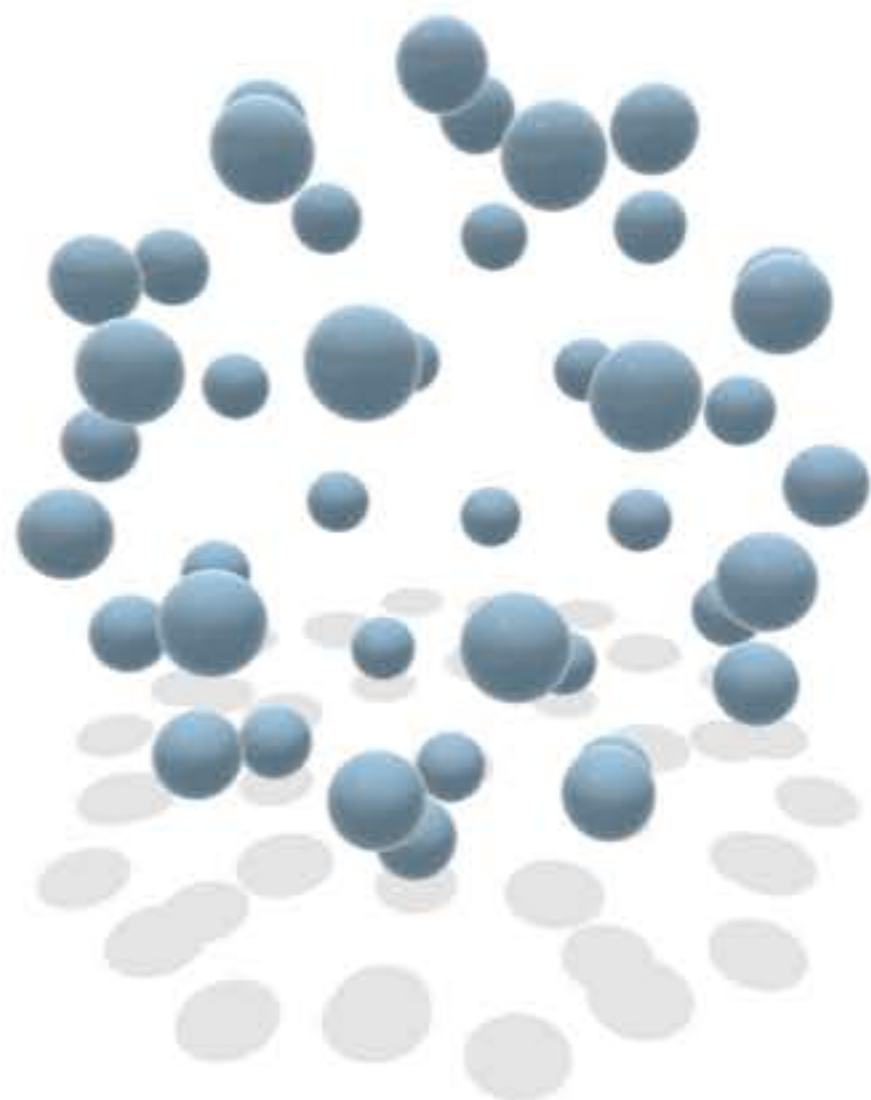
$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz} \quad \lambda \xrightarrow{\sim} \infty \quad g(x^*)$$

A second convolutional distance formula

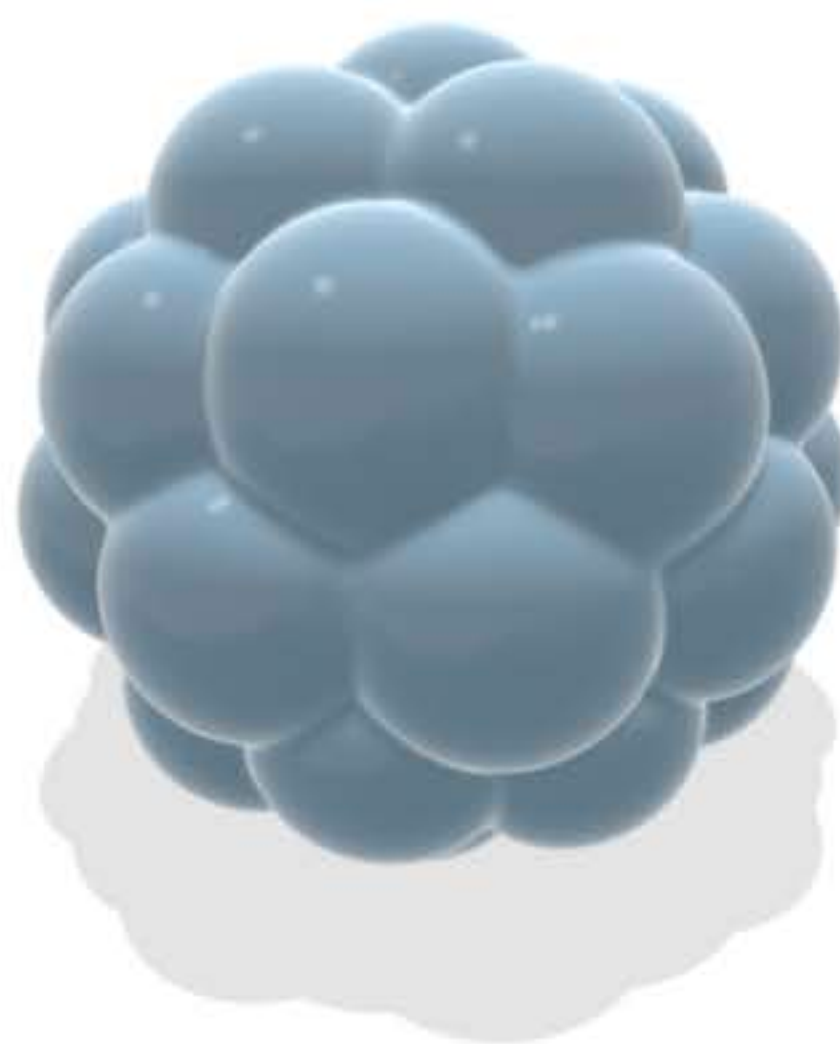
closest-point interpolation

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz} \quad \lambda \xrightarrow{\sim} \infty \quad g(x^*)$$

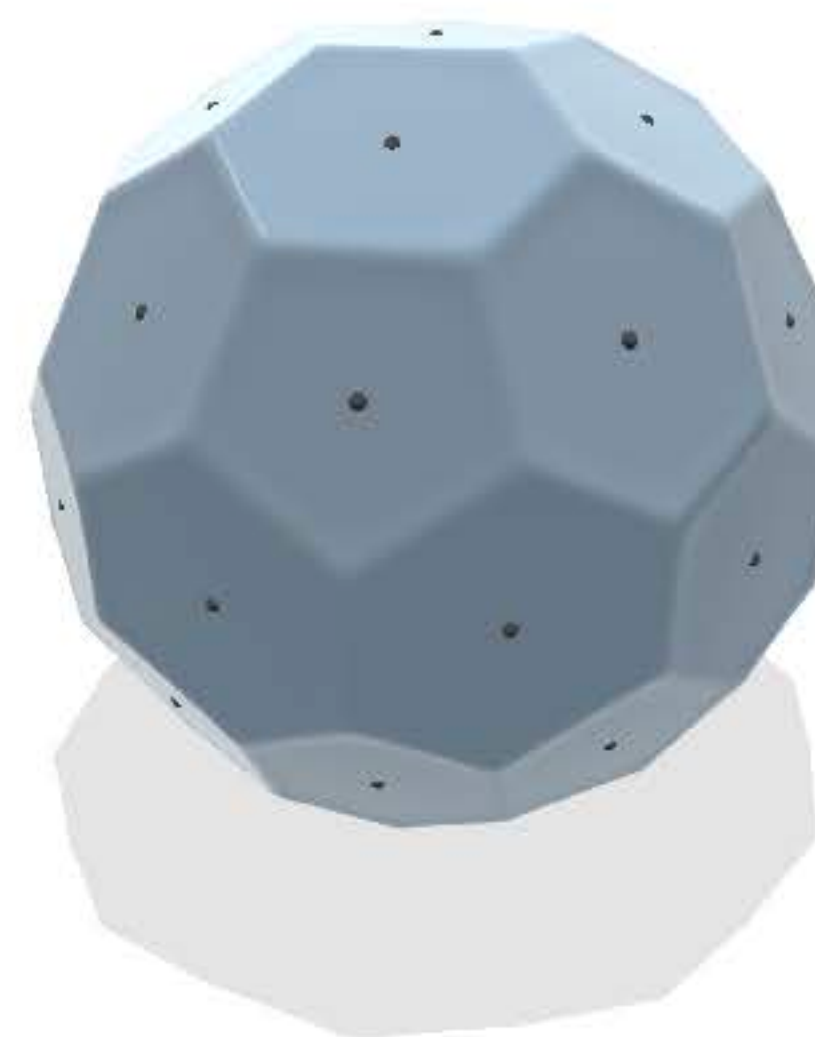
$$g(z) = \|x - z\| - 0.1$$



$$g(z) = \|x - z\| - 1$$



$$g(z) = \langle x - z, n_z \rangle$$



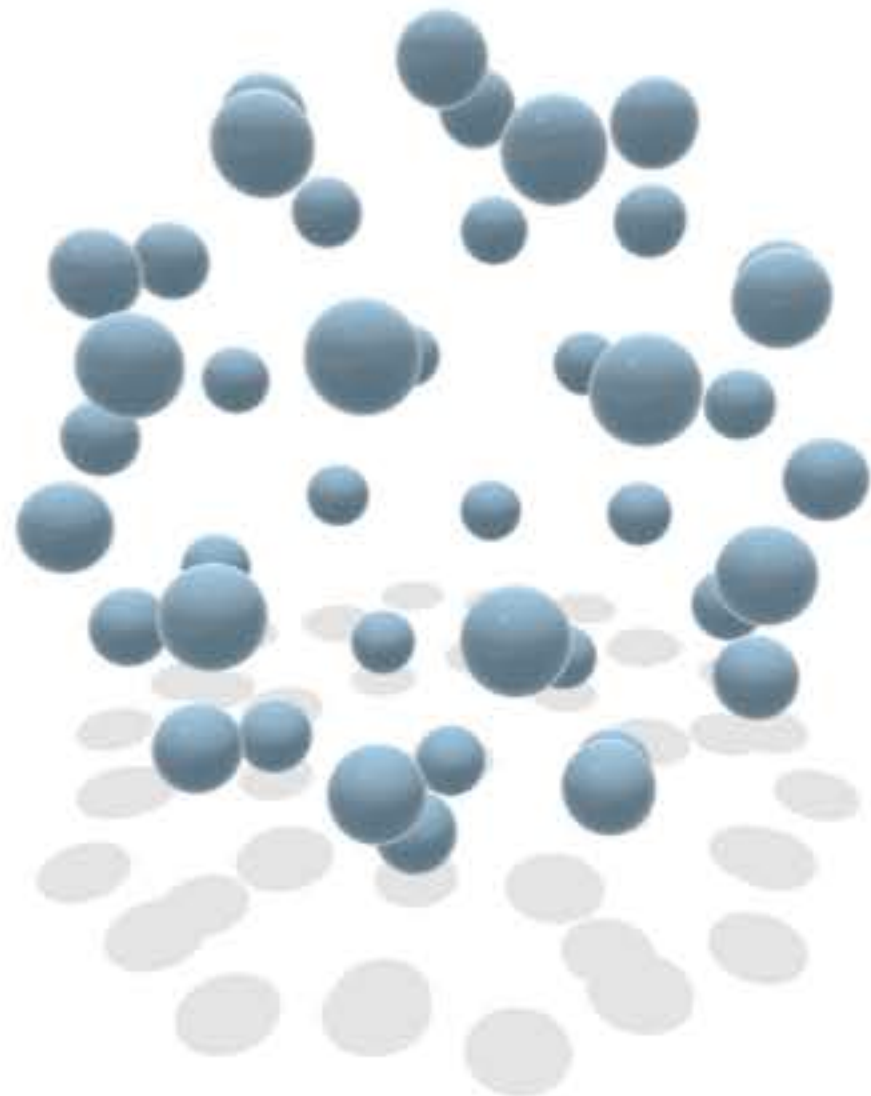
A second convolutional distance formula

*a.k.a. kernel density estimator, partition-of-unity methods, softmax function,
vector heat method, induced exponential function, ...*

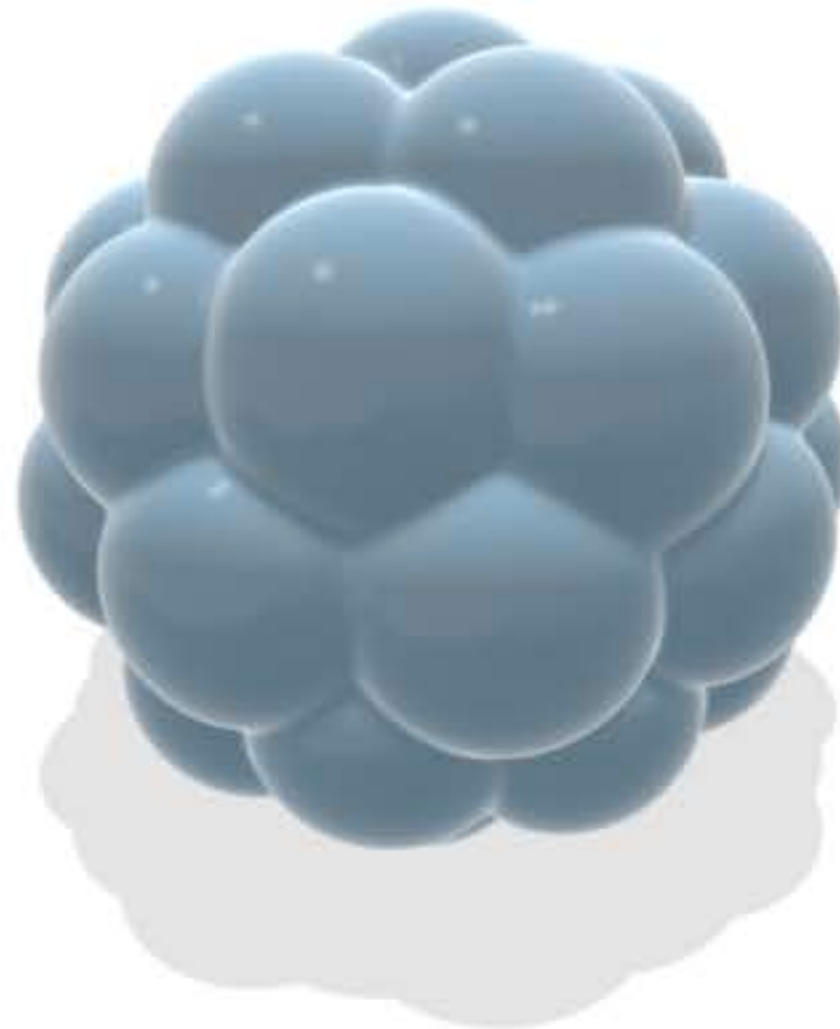
closest-point interpolation

$$\hat{d}(x) := \frac{\int_{\Omega} g(z) \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz} \quad \lambda \xrightarrow{\sim} \infty \quad g(x^*)$$

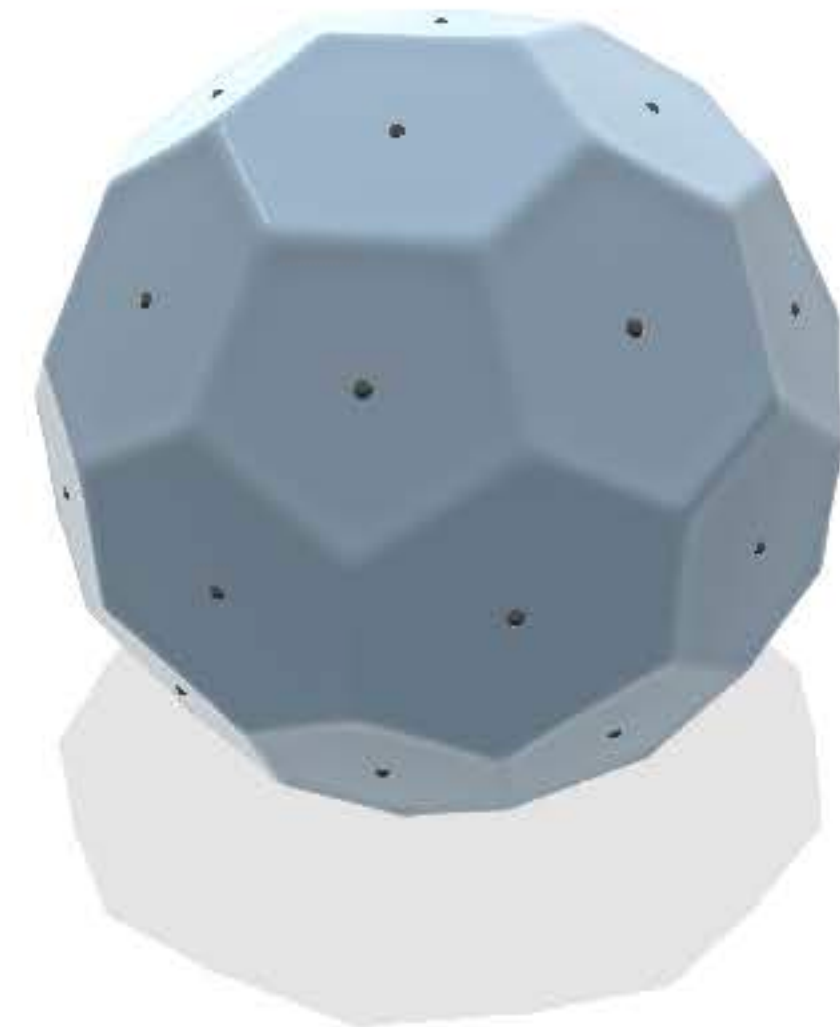
$$g(z) = \|x - z\| - 0.1$$



$$g(z) = \|x - z\| - 1$$



$$g(z) = \langle x - z, n_z \rangle$$



Linear approximations are undesirable...

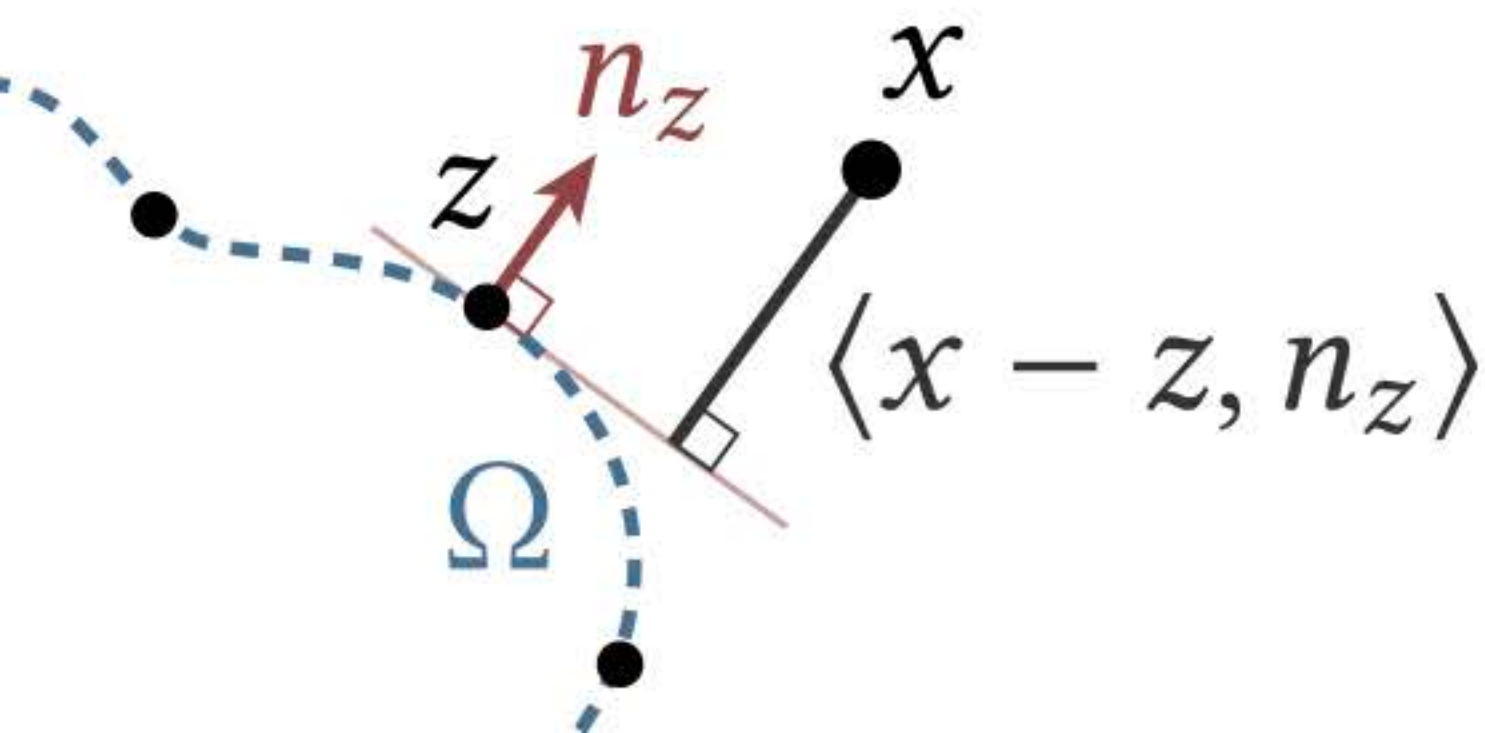
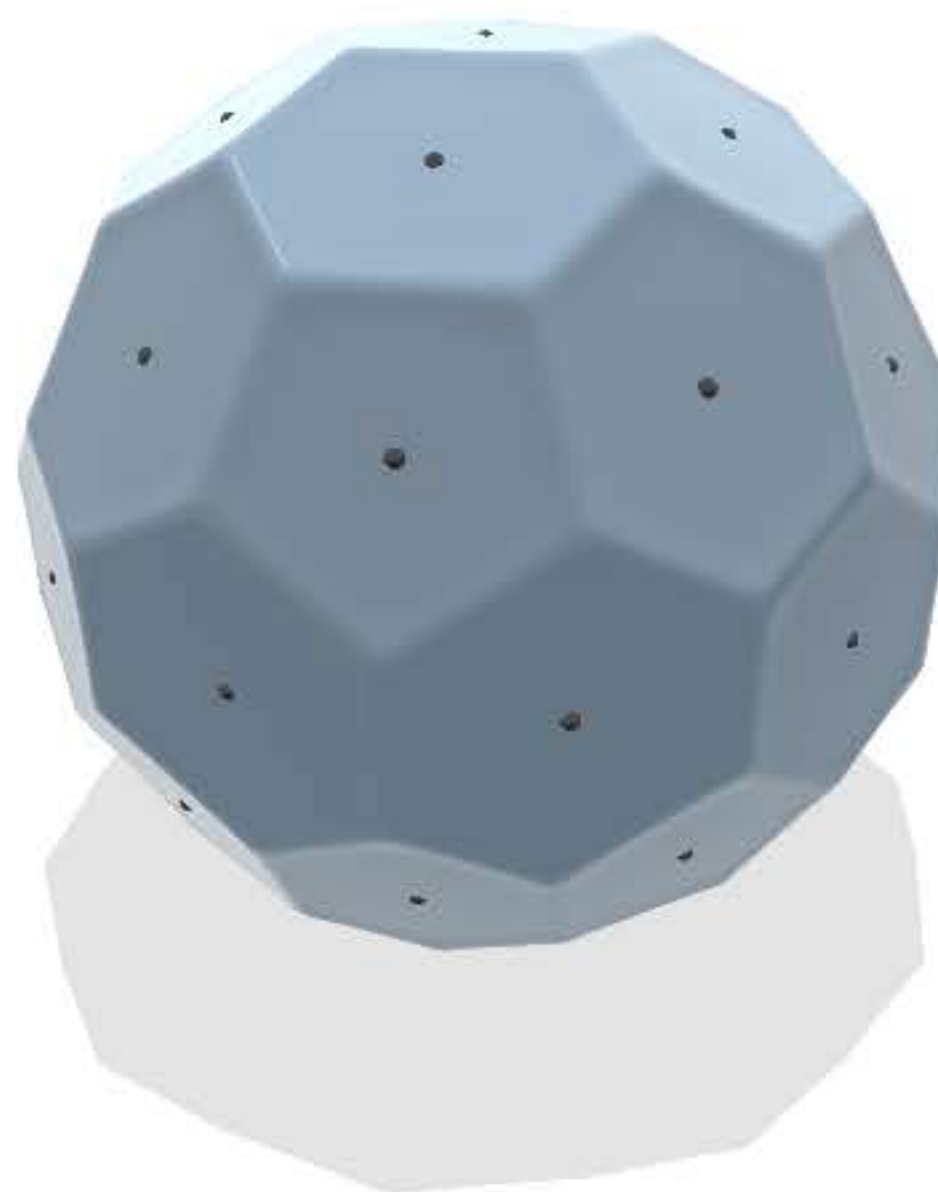
$$\hat{d}(x) = \frac{\int_{\Omega} \overset{\text{"normal" distance}}{\langle x - z, n_z \rangle} \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

“pseudonormal test”

“tangent plane approximation”

“plane test”

[Alexa et al. 2001; Boissonnat & Cazals
2002; Bærentzen & Aanæs 2005, Kolluri
2008; Öztireli et al. 2009; Yang et al. 2025]



Linear approximations are undesirable...

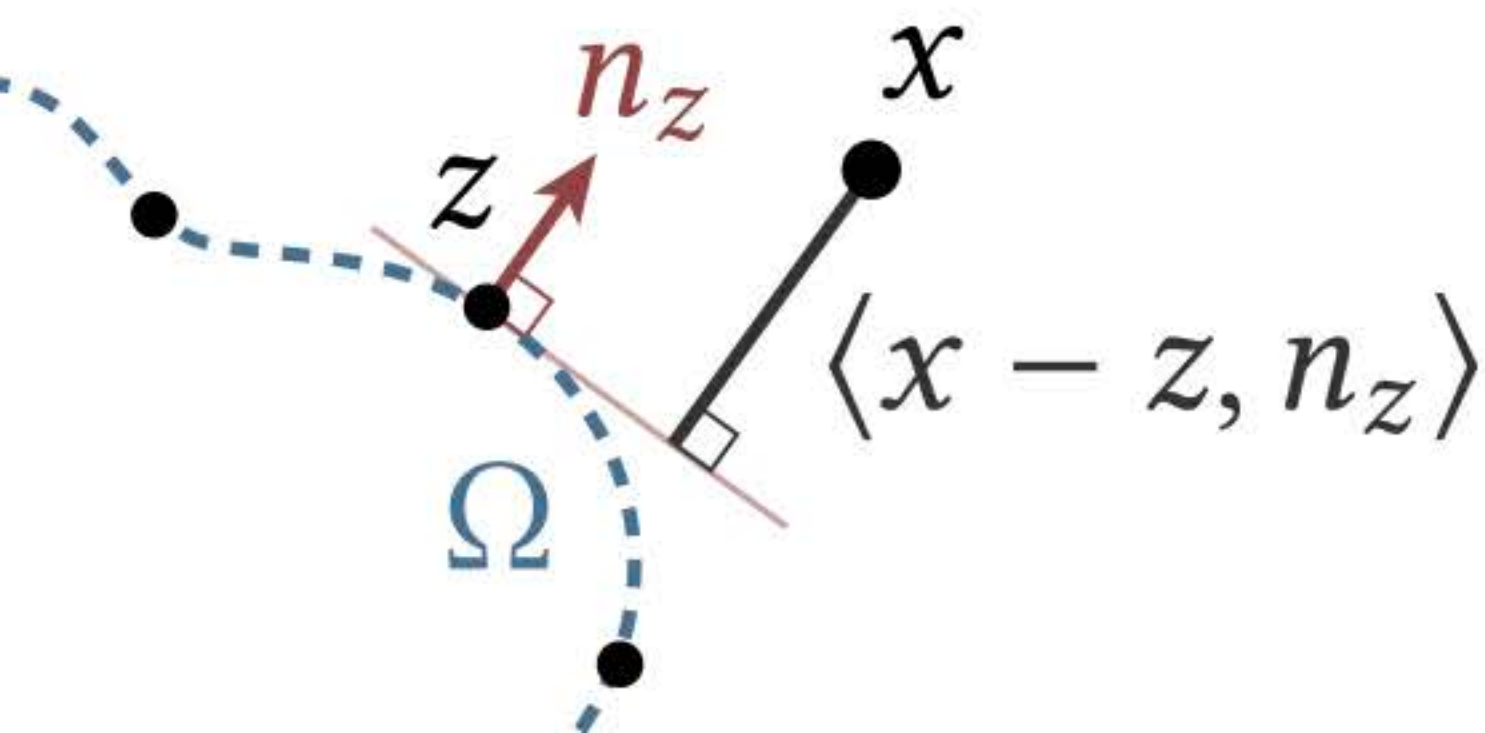
$$\hat{d}(x) = \frac{\int_{\Omega} \overset{\text{"normal" distance}}{\langle x - z, n_z \rangle} \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

“pseudonormal test”

“tangent plane approximation”

“plane test”

[Alexa et al. 2001; Boissonnat & Cazals
2002; Bærentzen & Aanæs 2005, Kolluri
2008; Öztireli et al. 2009; Yang et al. 2025]

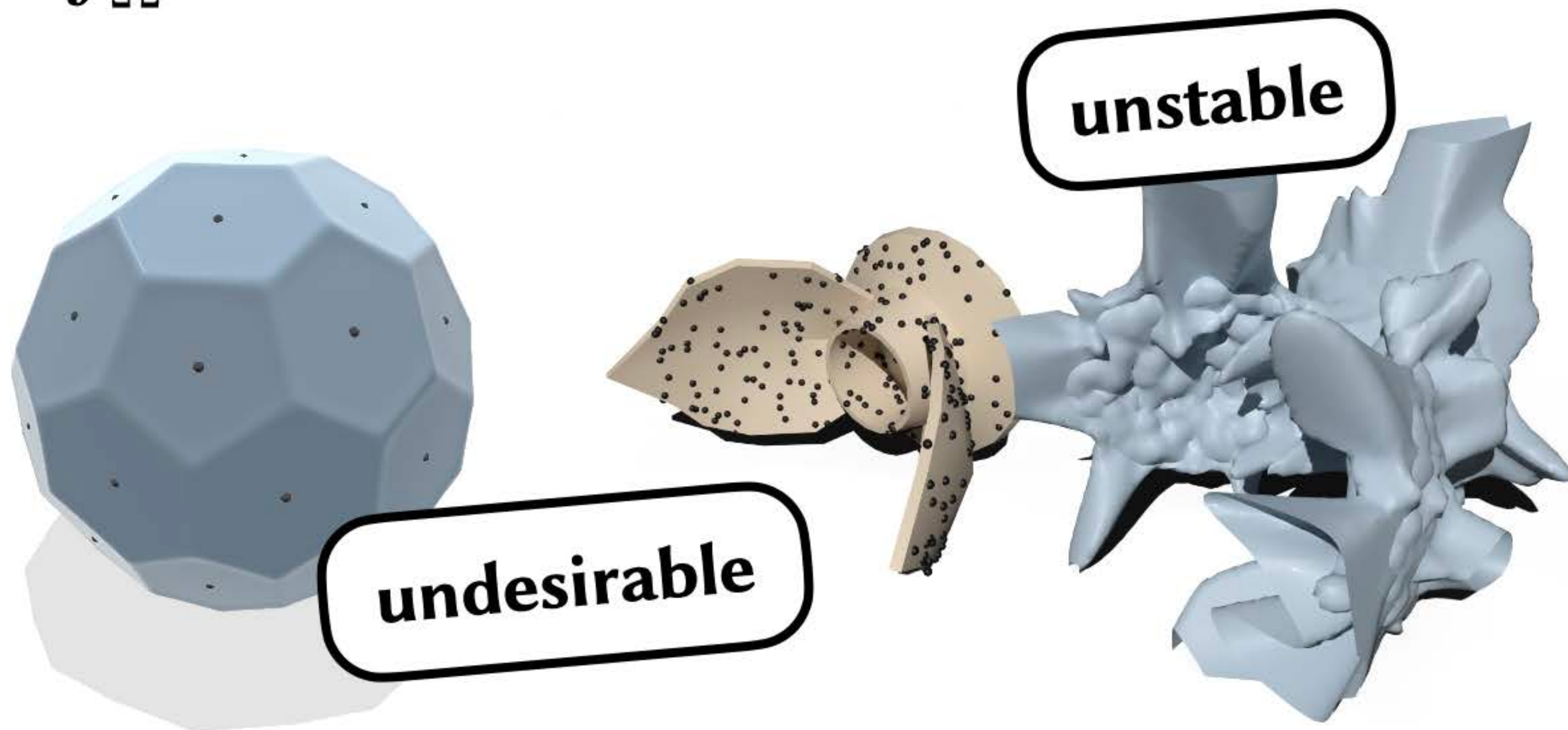
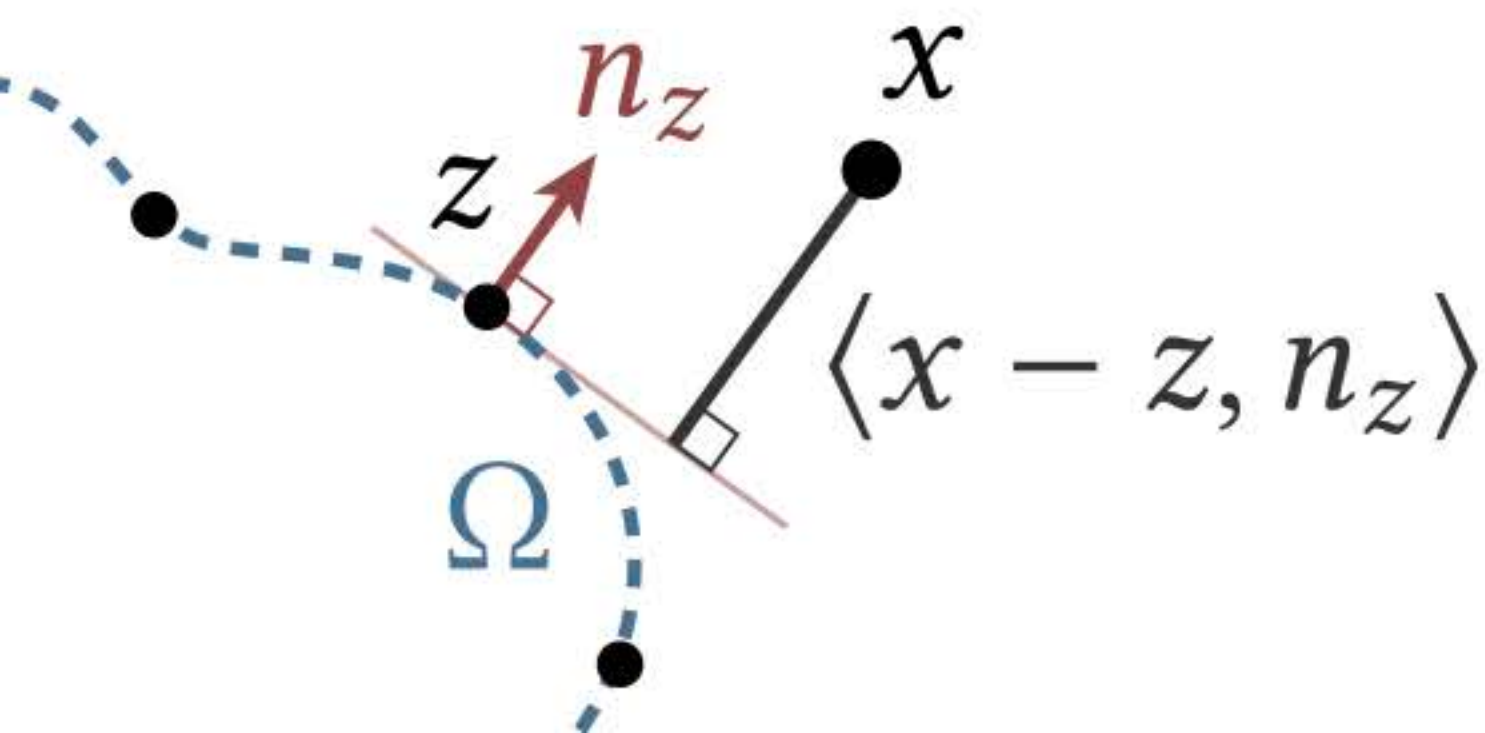


Linear approximations are undesirable...

$$\hat{d}(x) = \frac{\int_{\Omega} \overset{\text{"normal" distance}}{\langle x - z, n_z \rangle} \exp(-\lambda \|x - z\|) dz}{\int_{\Omega} \exp(-\lambda \|x - z\|) dz}$$

“pseudonormal test”
“tangent plane approximation”
“plane test”

[Alexa et al. 2001; Boissonnat & Cazals 2002; Bærentzen & Aanæs 2005, Kolluri 2008; Öztireli et al. 2009; Yang et al. 2025]



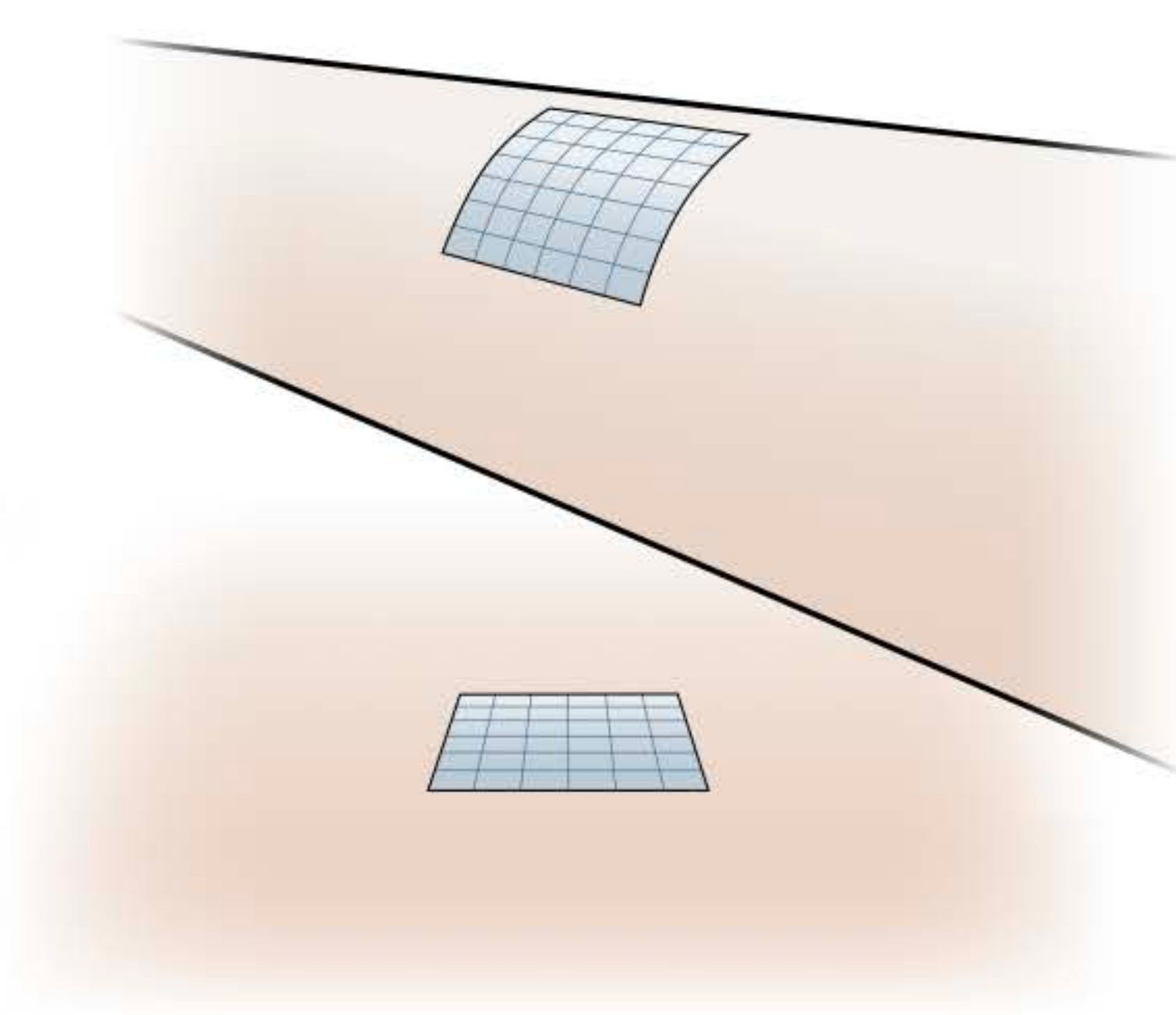
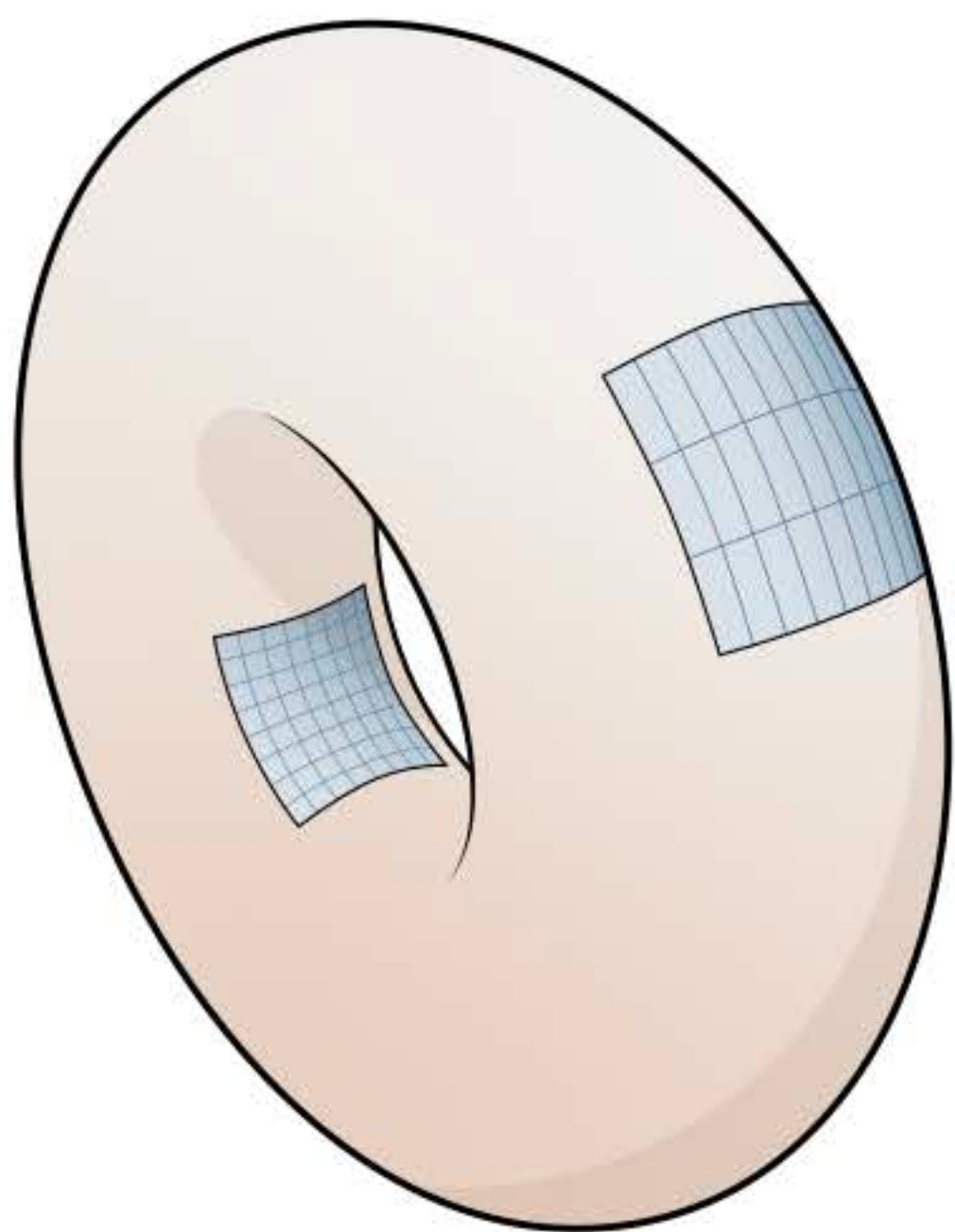
... so let's use second-order surfaces!

Discrete estimator:
$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

... so let's use second-order surfaces!

Discrete estimator:

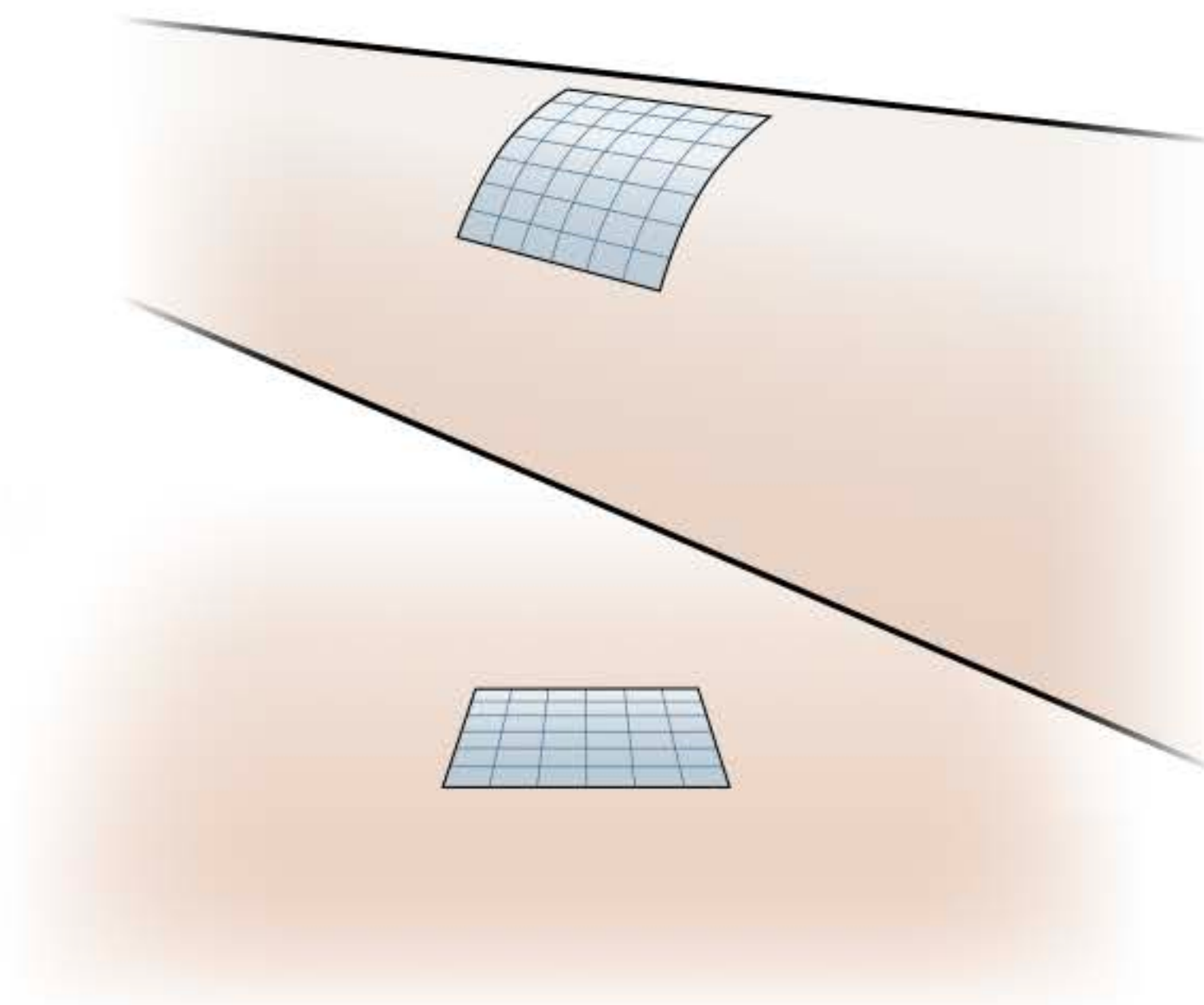
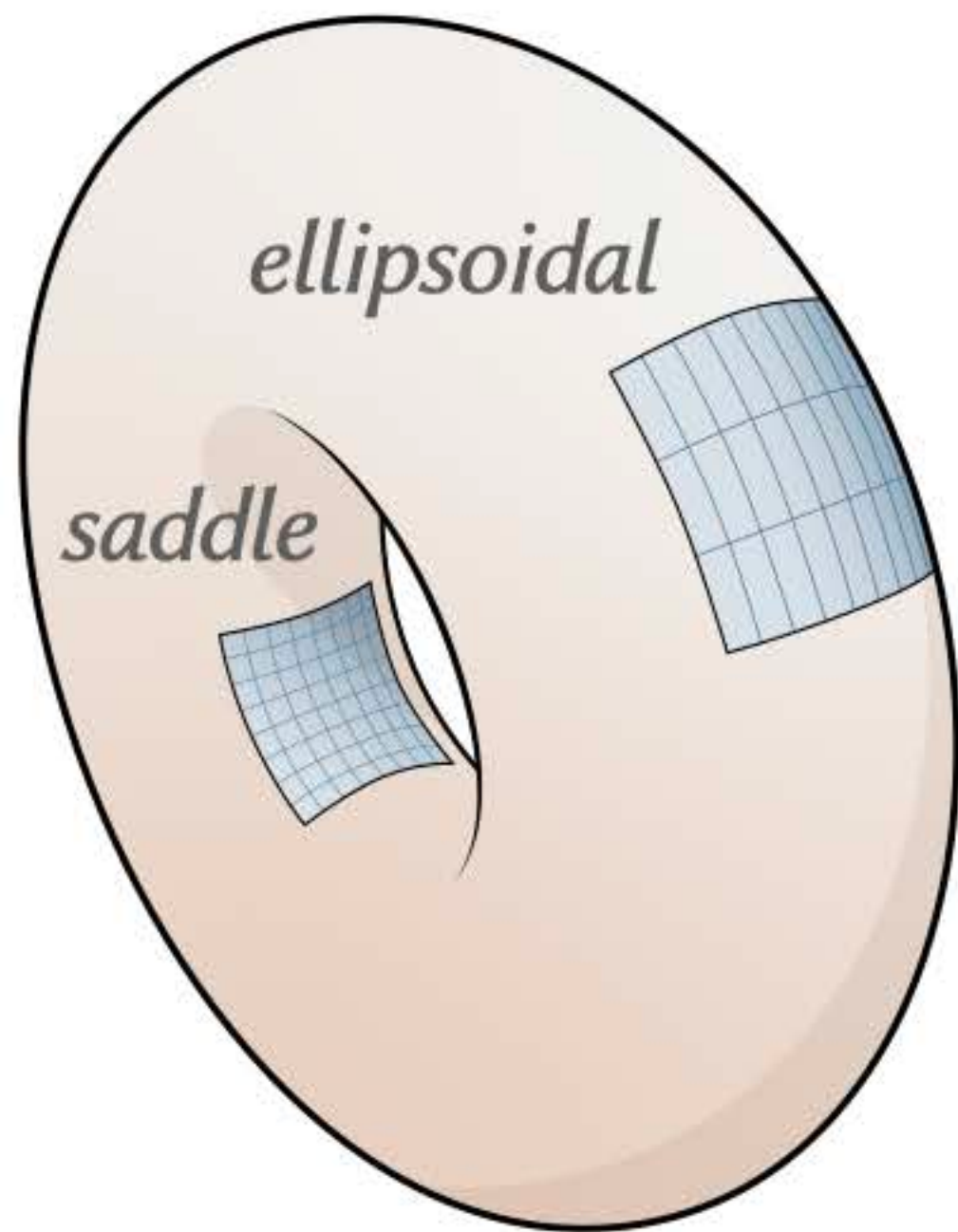
$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{torus distance!}}{g_i(x)} \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$



... so let's use second-order surfaces!

Discrete estimator:

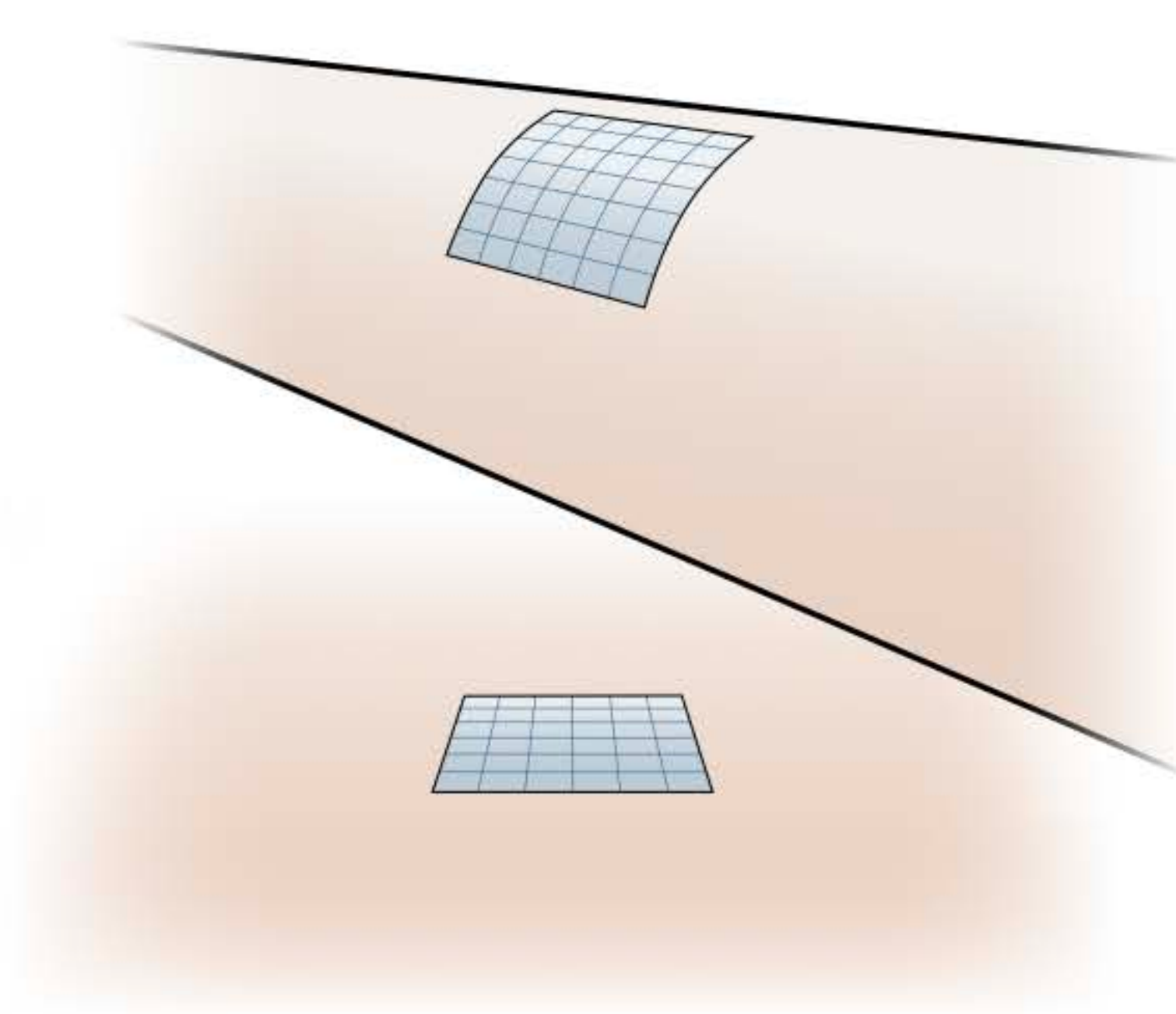
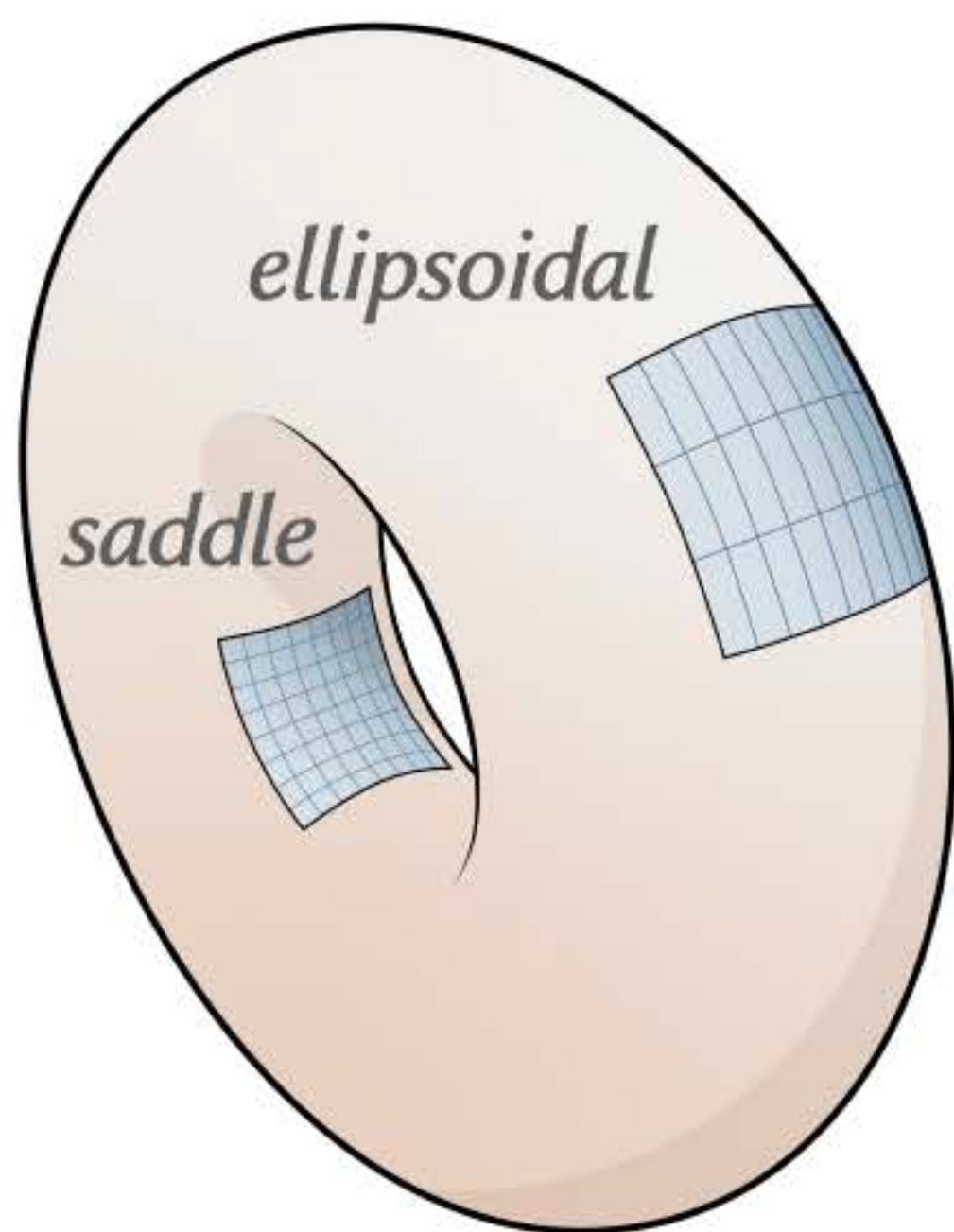
$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{torus distance!}}{g_i(x)} \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$



... so let's use second-order surfaces!

Discrete estimator:

$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{torus distance!}}{g_i(x)} \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$



Fitting geometric proxies to point clouds (planes, cylinders, spheres, quadrics, splines, polynomial patches...)

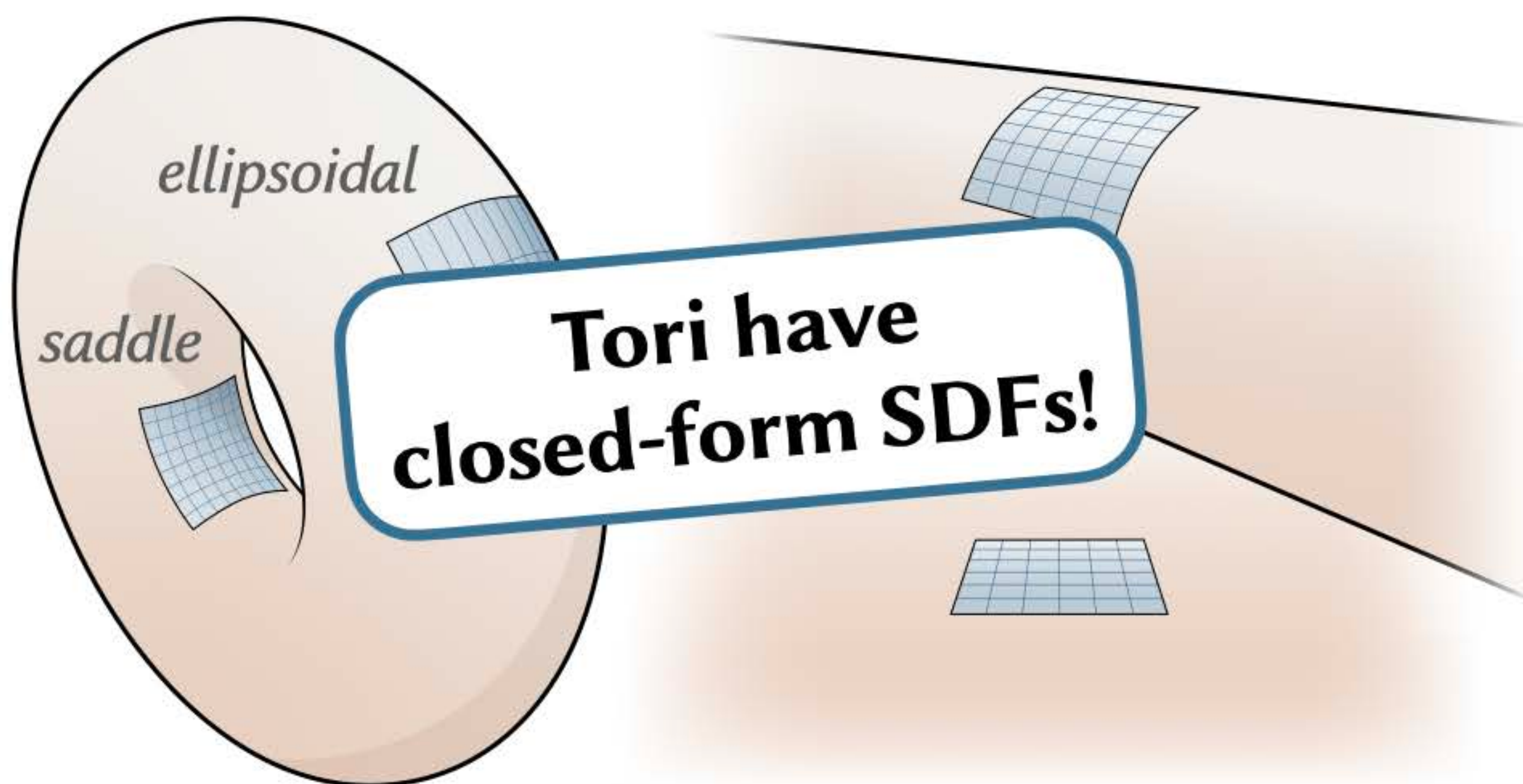
Faugeras 1983; Besl & Jain 1988; Cao et al. 1994; Kaveti et al. 1996; Yang & Lee 1999; Cohen-Steiner et al. 2004; Cazals & Pouget 2005; Wu & Kobbelt 2005; Simari & Singh 2005; Yan et al. 2006; Attene et al. 2006; Xia & Ju 2025...

Tori: Lukács et al. 1998; Liu et al. 2009; Eberly 2020
— but not used to construct SDFs!

... so let's use second-order surfaces!

Discrete estimator:

$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{torus distance!}}{g_i(x)} \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

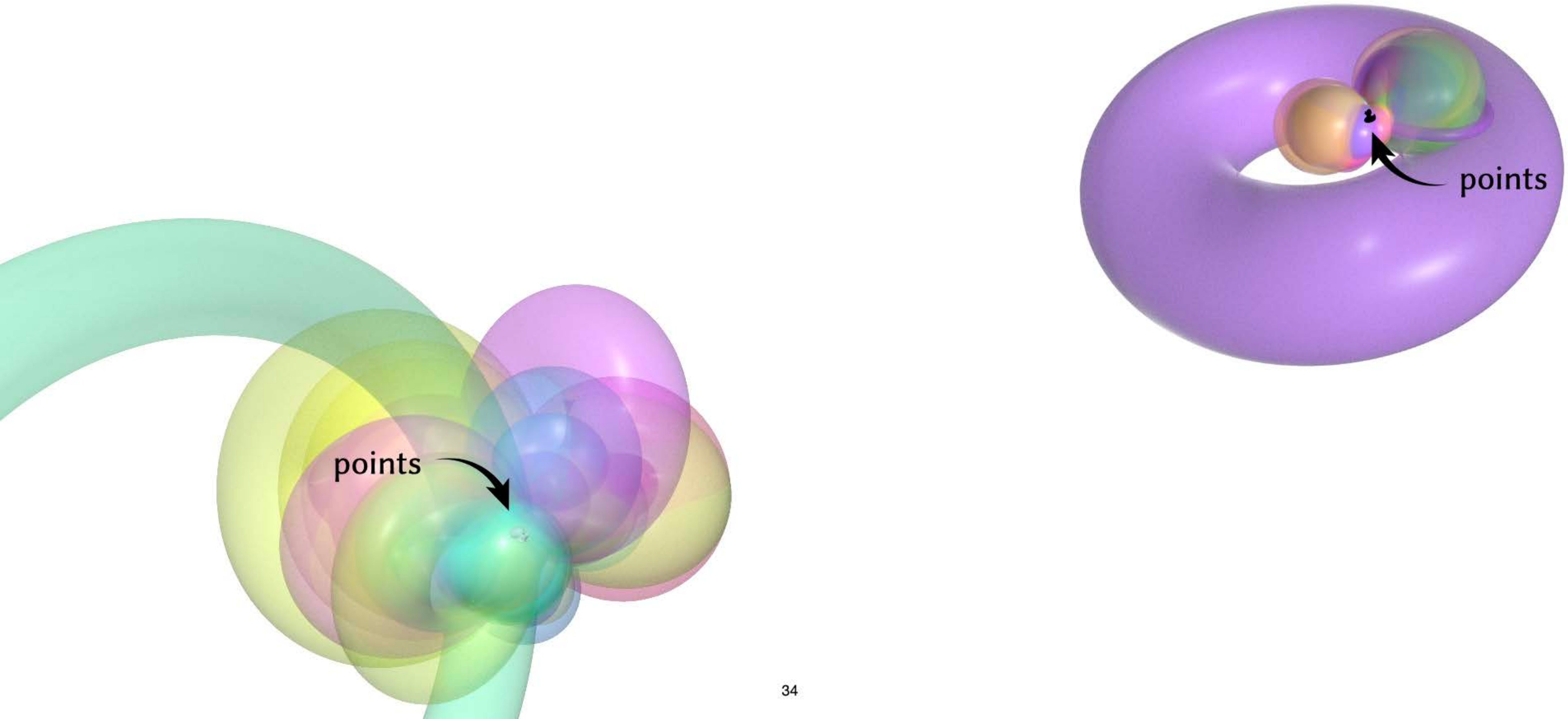


Fitting geometric proxies to point clouds (planes, cylinders, spheres, quadrics, splines, polynomial patches...)

Faugeras 1983; Besl & Jain 1988; Cao et al. 1994; Kaveti et al. 1996; Yang & Lee 1999; Cohen-Steiner et al. 2004; Cazals & Pouget 2005; Wu & Kobbelt 2005; Simari & Singh 2005; Yan et al. 2006; Attene et al. 2006; Xia & Ju 2025...

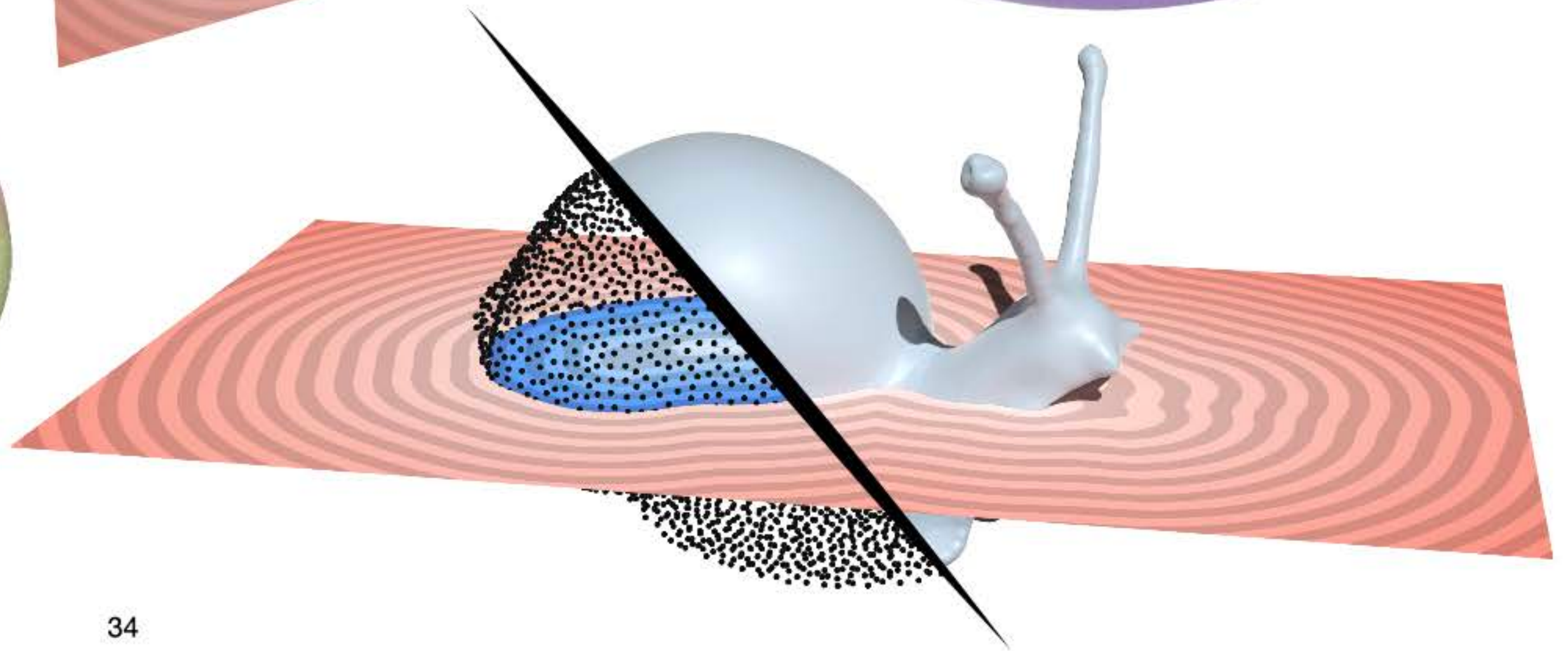
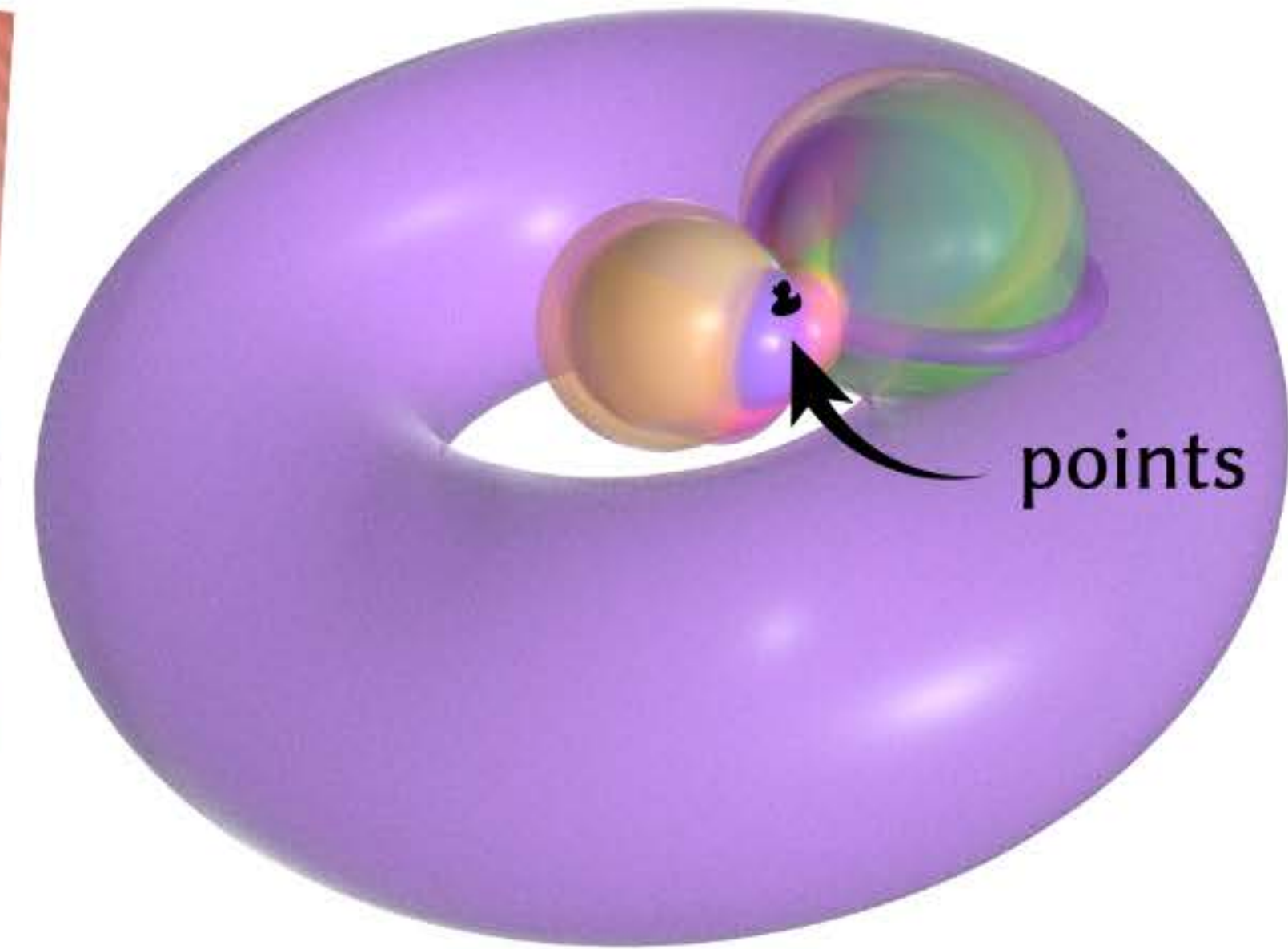
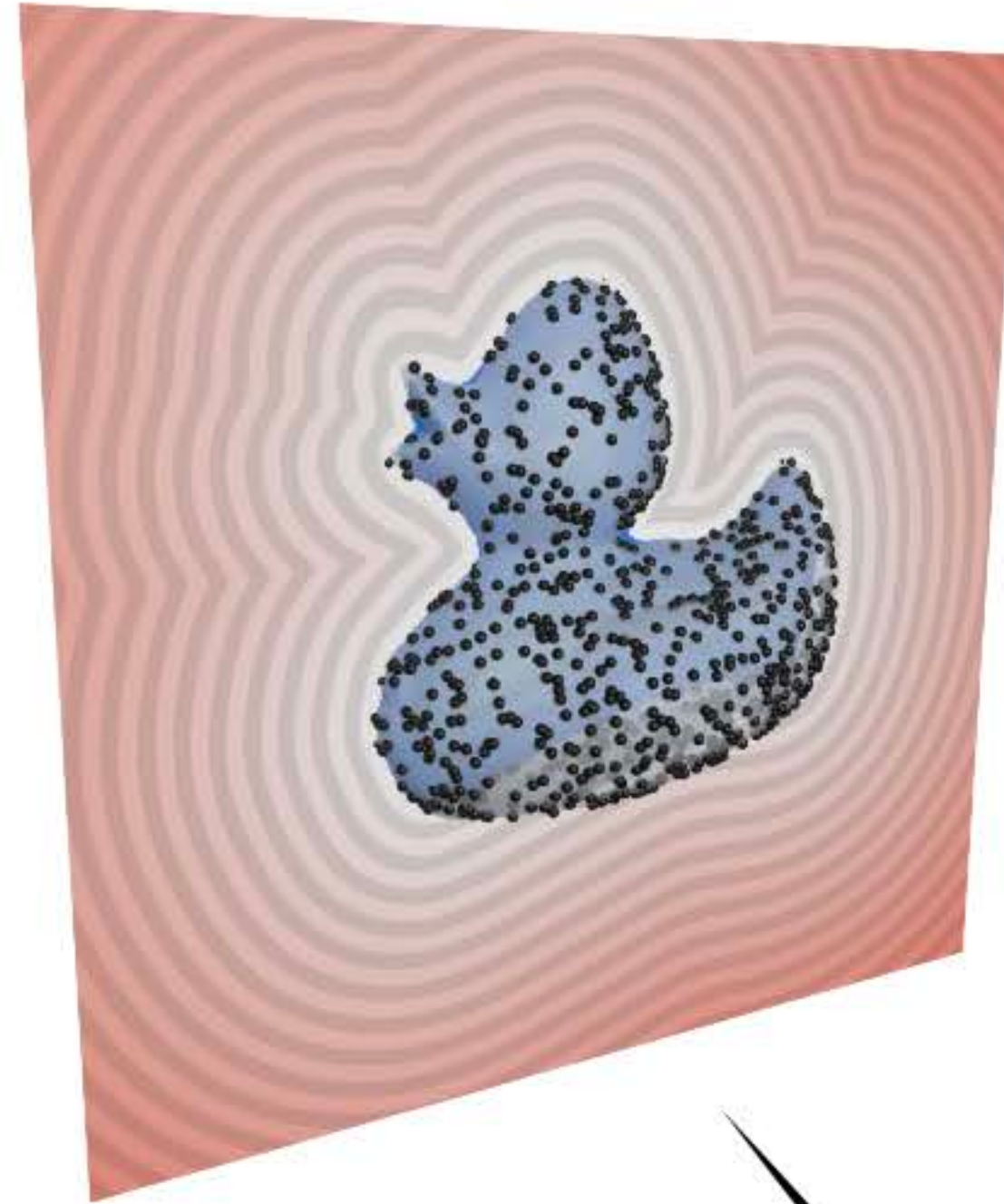
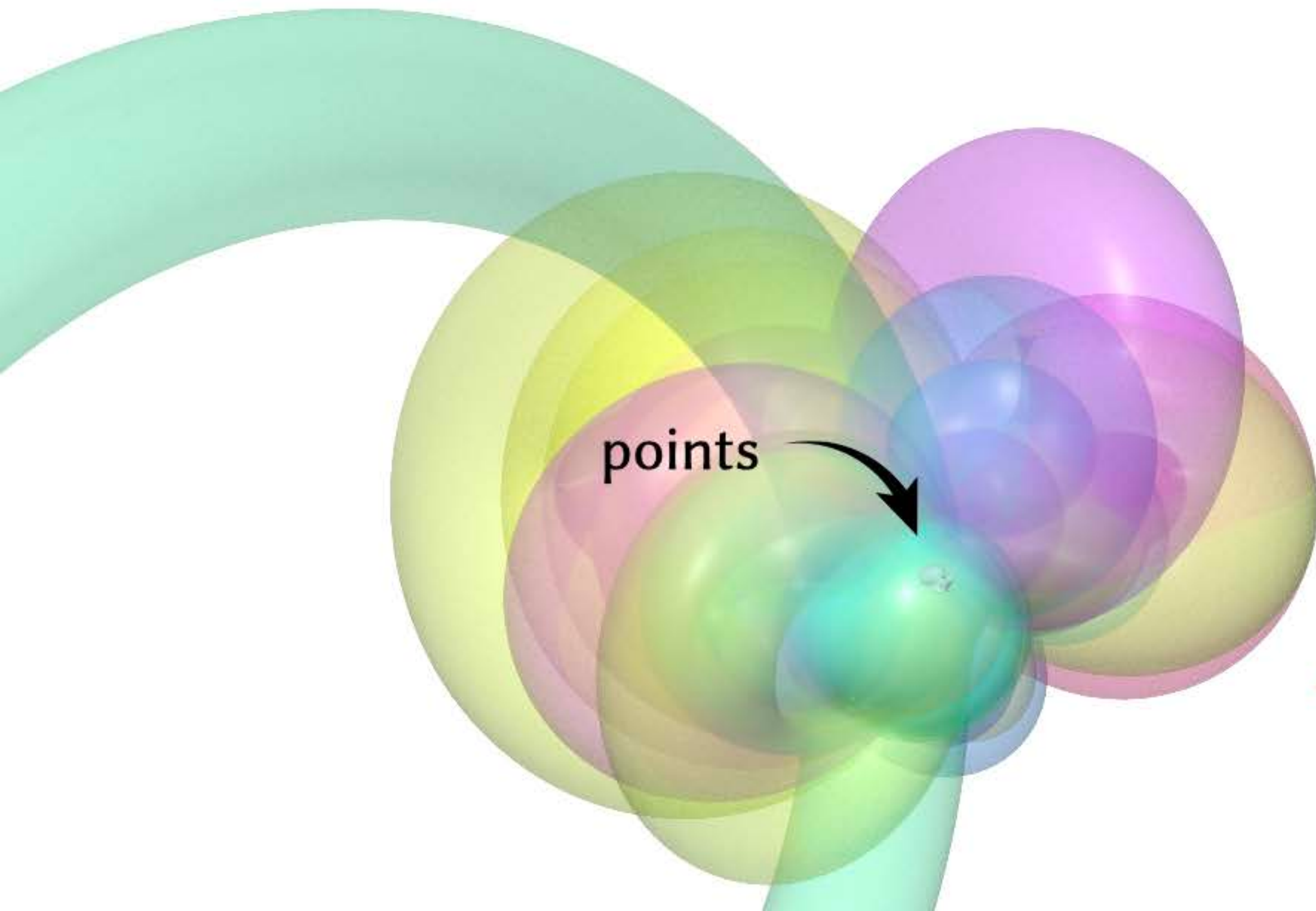
Tori: Lukács et al. 1998; Liu et al. 2009; Eberly 2020
— but not used to construct SDFs!

The tori look crazy... but they work together



The tori look crazy... but they work together

$$\phi(x) = \frac{\sum_{i=1}^{|P|} \overset{\text{i-th torus SDF}}{g_i(x)} \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

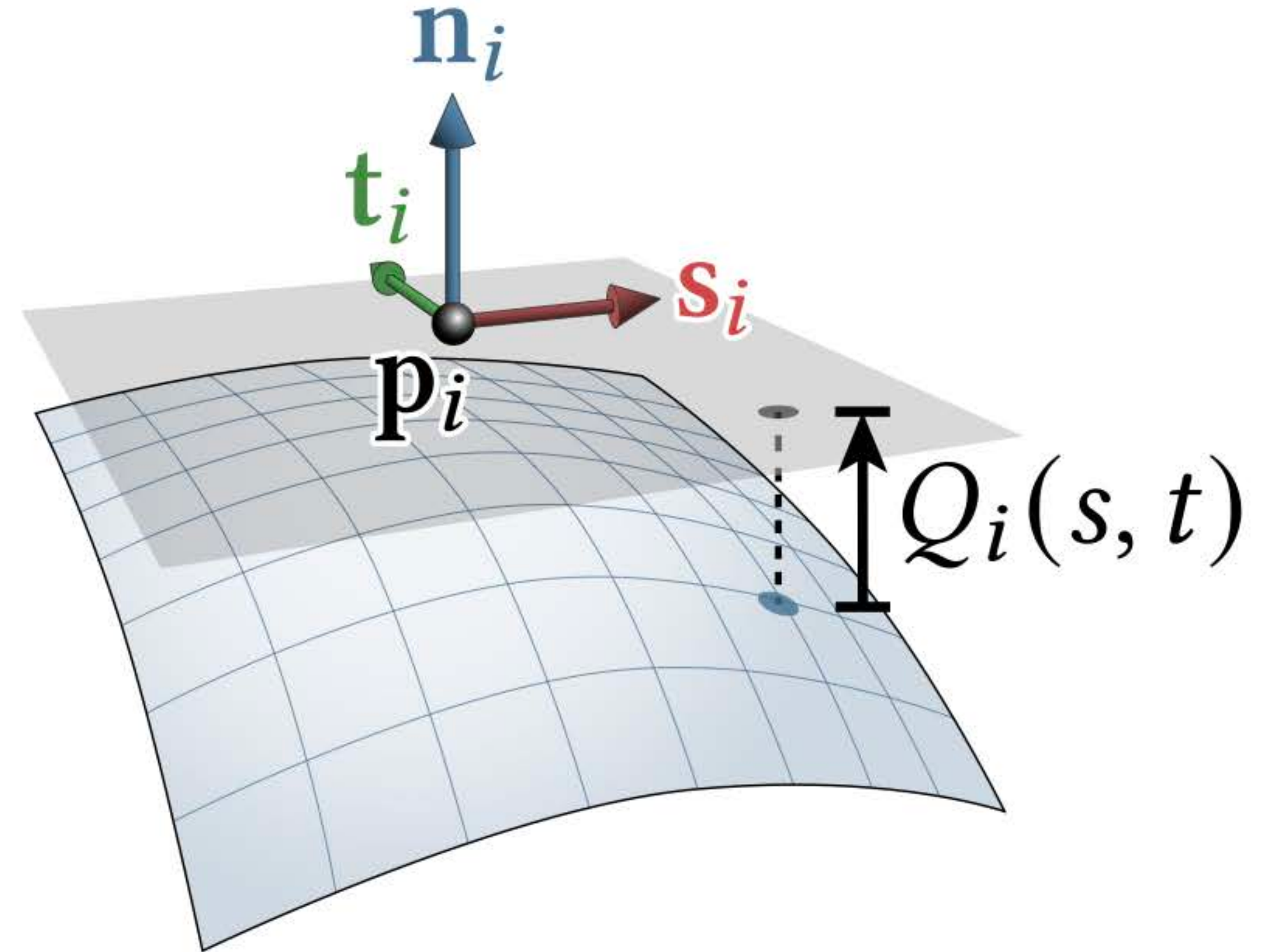


The surface, one polynomial patch at a time

The surface, one polynomial patch at a time

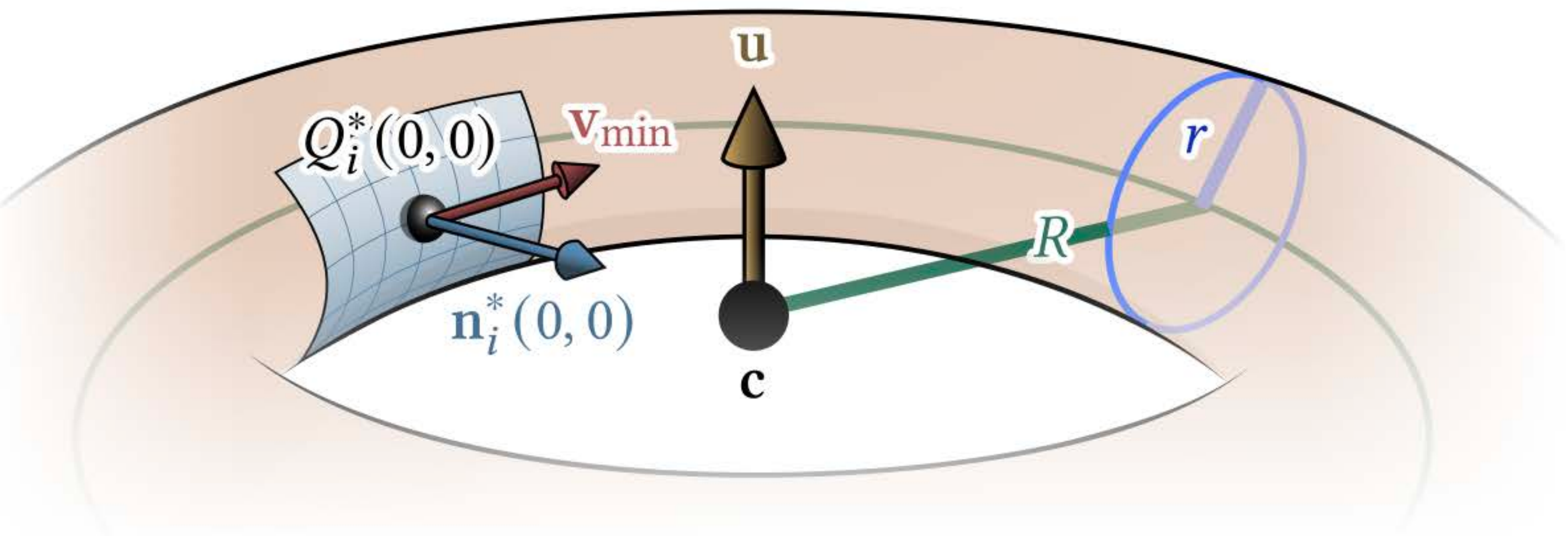
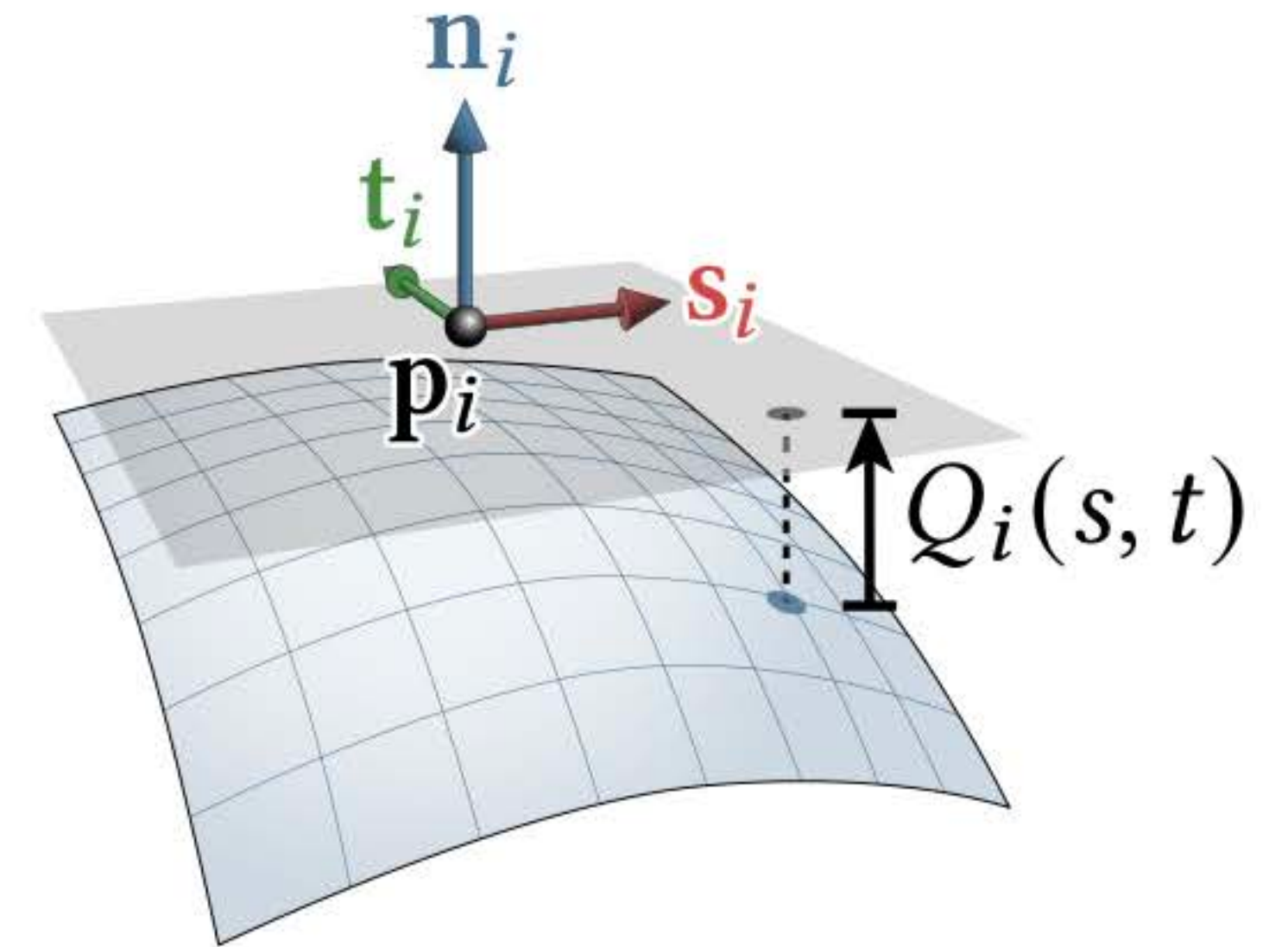
Quadratic polynomial
around each point:

$$Q_i(s, t) = \sum_{n=0}^2 \sum_{m=0}^2 a_{n,m} s^n t^m$$



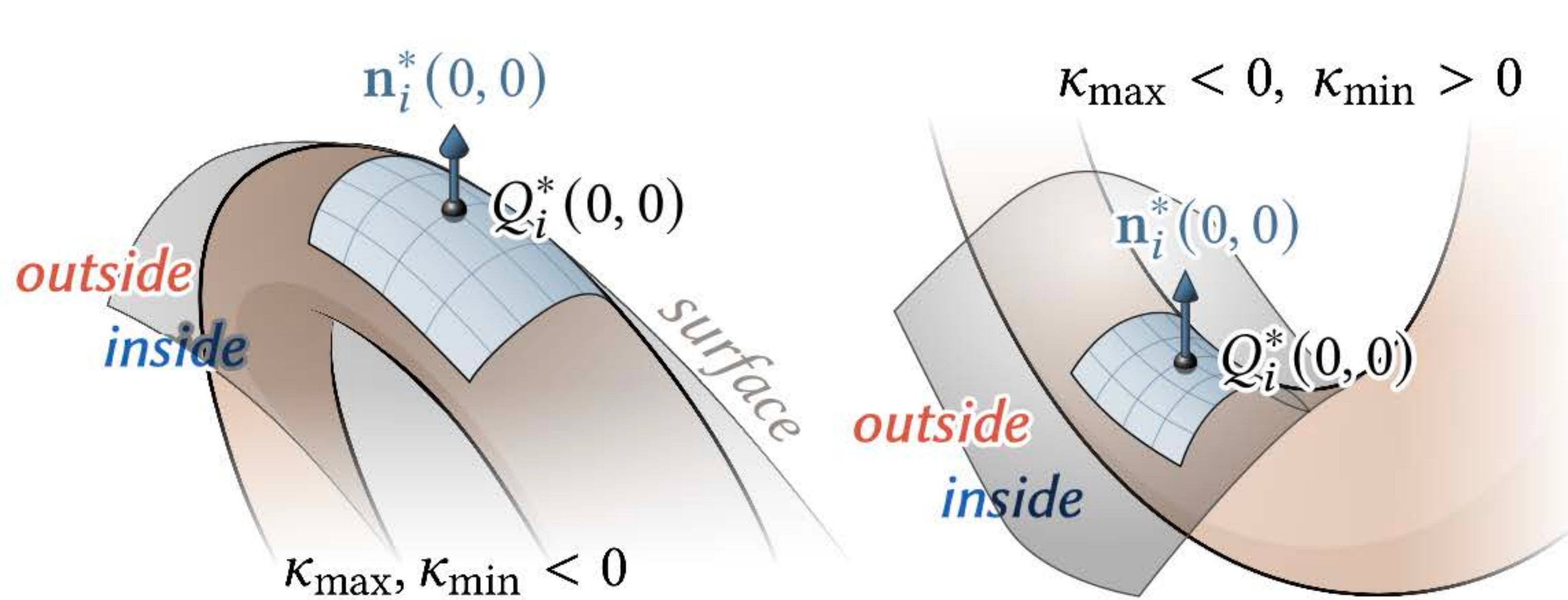
The surface, one polynomial patch at a time

$$Q_i(s, t) = \sum_{n=0}^2 \sum_{m=0}^2 a_{n,m} s^n t^m$$

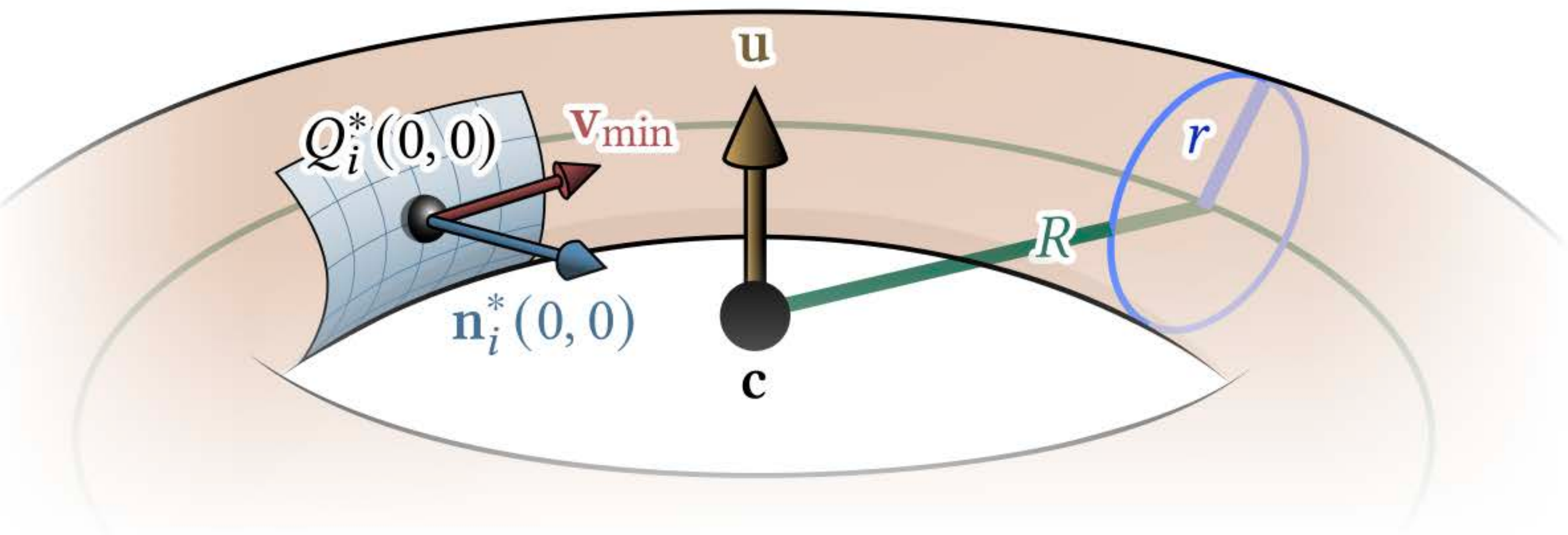
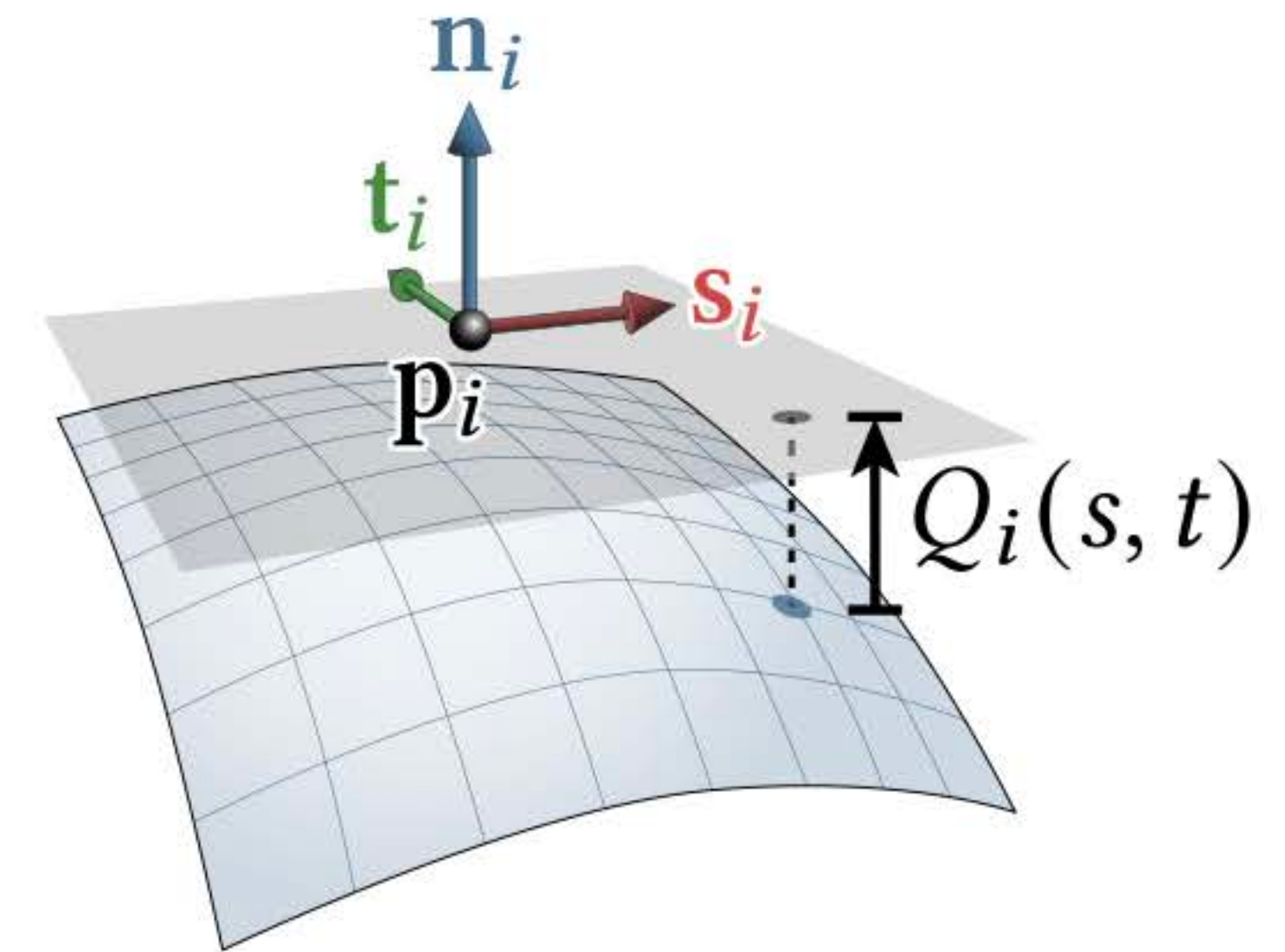


A torus is determined by the six coefficients $a_{0,0}, a_{0,1}, a_{1,0}, a_{1,1}, a_{0,2}, a_{2,0}$

The surface, one polynomial patch at a time

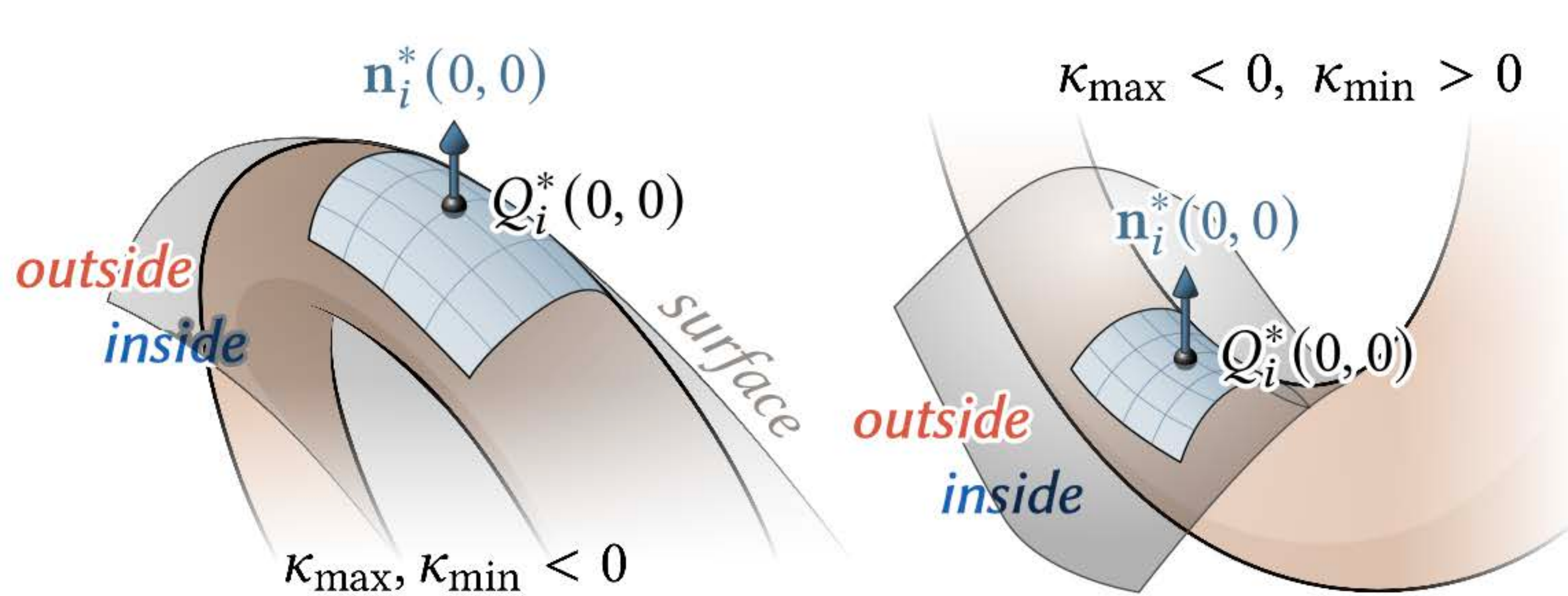


$$Q_i(s, t) = \sum_{n=0}^2 \sum_{m=0}^2 a_{n,m} s^n t^m$$



A torus is determined by the six coefficients $a_{0,0}, a_{0,1}, a_{1,0}, a_{1,1}, a_{0,2}, a_{2,0}$

The surface, one polynomial patch at a time

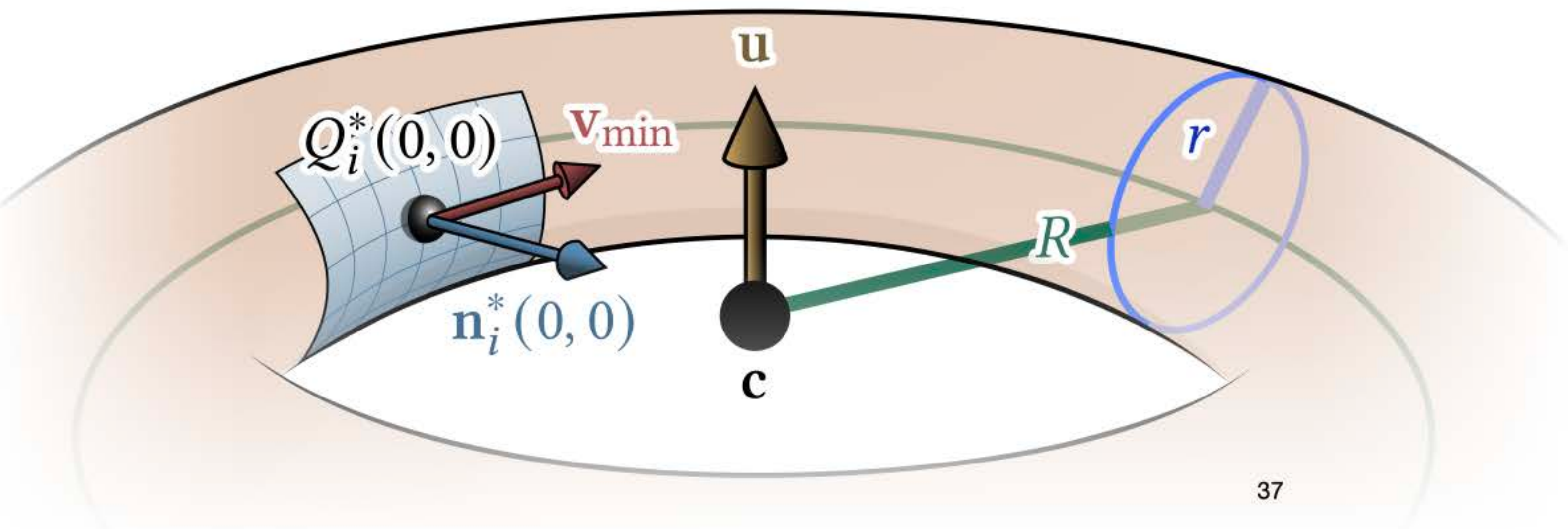


$$Q_i(s, t) = \sum_{n=0}^2 \sum_{m=0}^2 a_{n,m} s^n t^m$$

SDF of a torus:

$$\phi_{\mathbb{T}} = \|\mathbf{d}\| - r,$$

$$\mathbf{d} := (\|(\mathbf{x} - \mathbf{c}) \times \mathbf{u}\| - R, \langle \mathbf{x} - \mathbf{c}, \mathbf{u} \rangle)$$



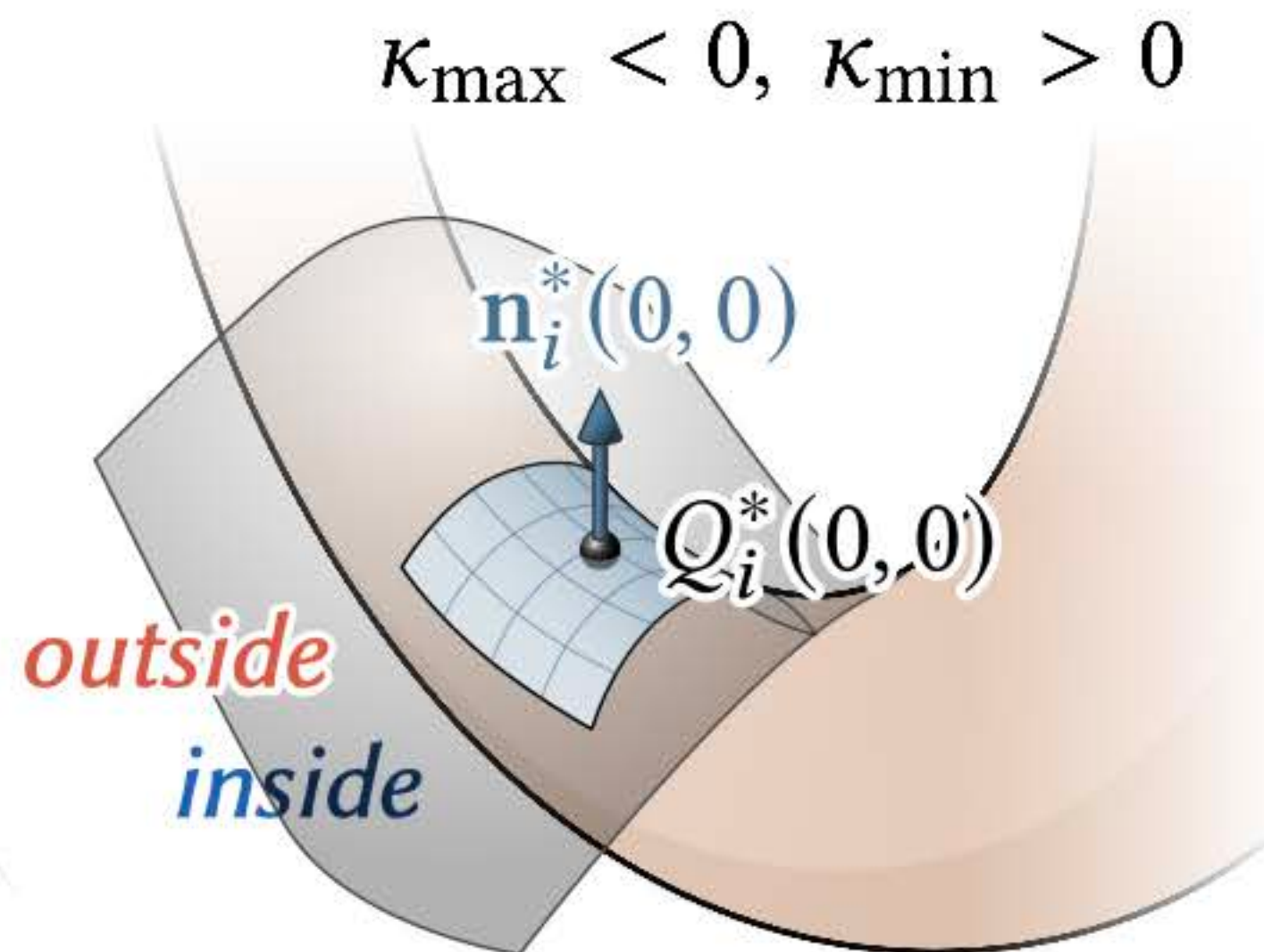
A torus is determined by the six coefficients $a_{0,0}, a_{0,1}, a_{1,0}, a_{1,1}, a_{0,2}, a_{2,0}$

The surface, one polynomial patch at a time

Fitting a torus to a given surface is easy...

$$\kappa_{\max}, \kappa_{\min} < 0$$

$$\kappa_{\max} < 0, \kappa_{\min} > 0$$

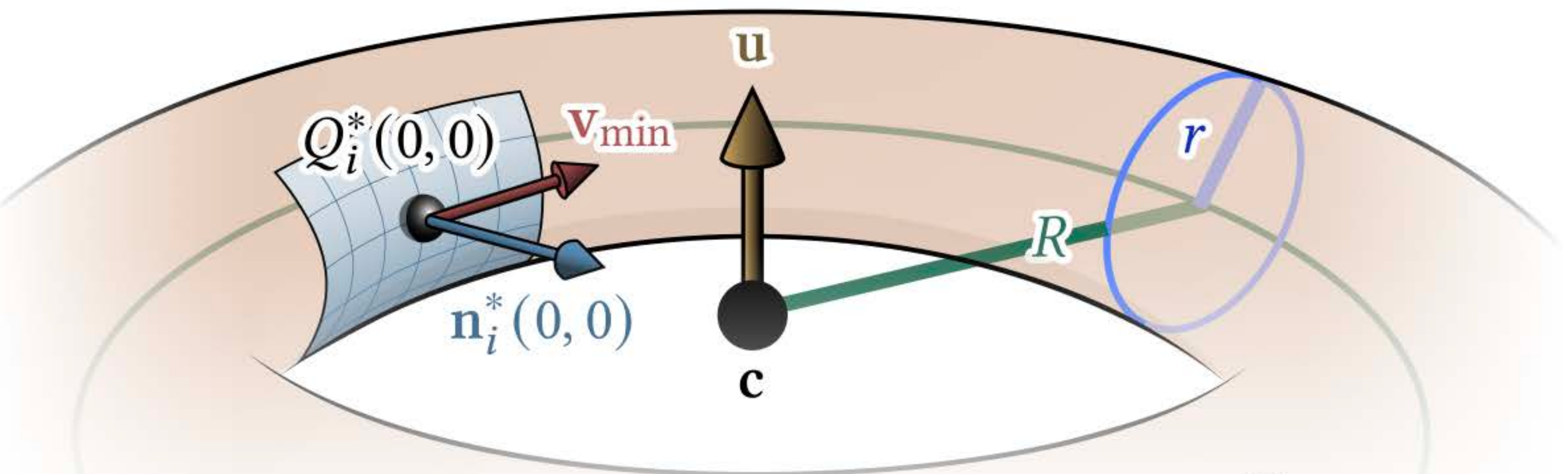


$$Q_i(s, t) = \sum_{n=0}^2 \sum_{m=0}^2 a_{n,m} s^n t^m$$

SDF of a torus:

$$\phi_{\mathbb{T}} = \|\mathbf{d}\| - r,$$

$$\mathbf{d} := (\|(\mathbf{x} - \mathbf{c}) \times \mathbf{u}\| - R, \langle \mathbf{x} - \mathbf{c}, \mathbf{u} \rangle)$$



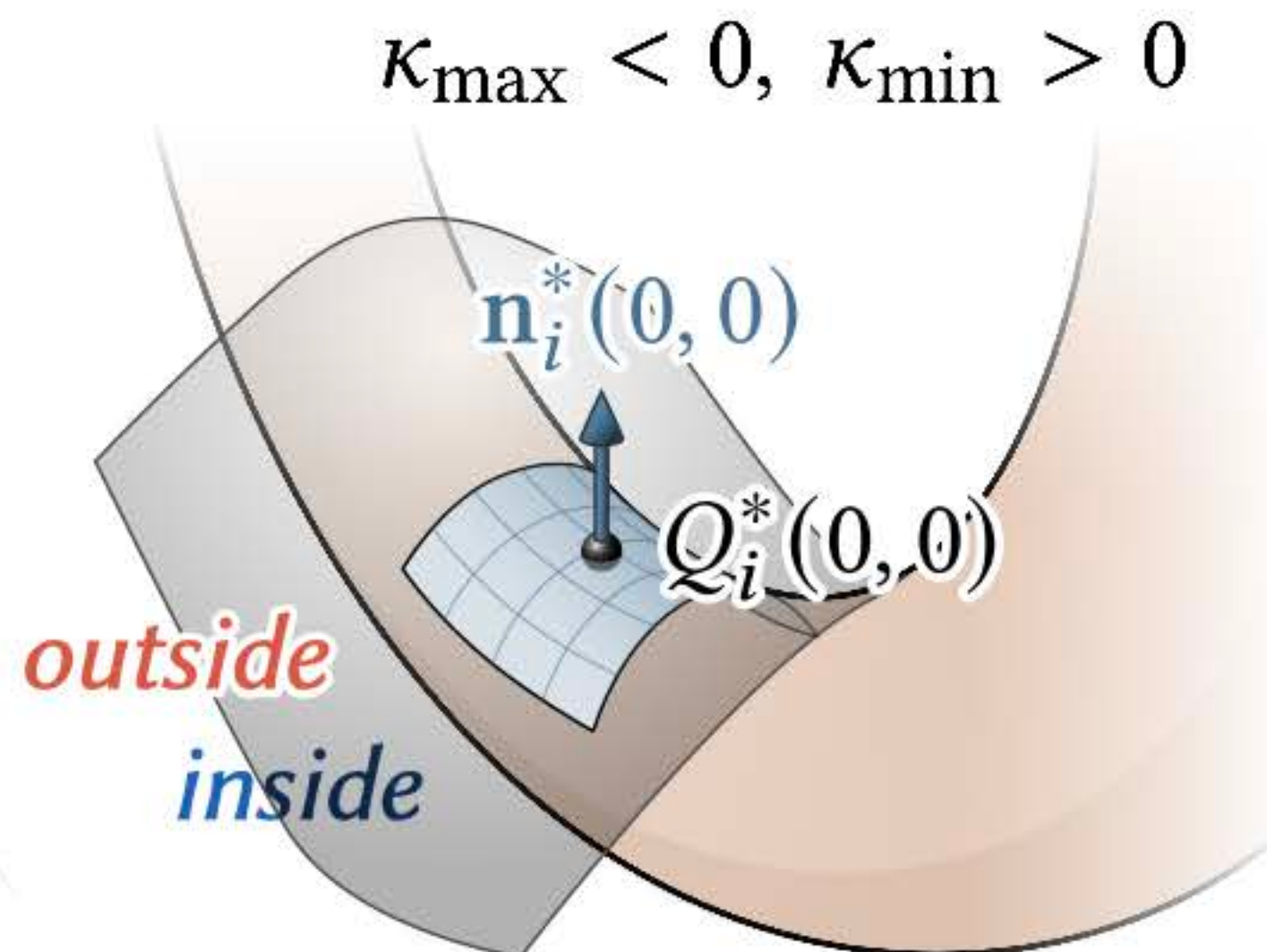
A torus is determined by the six coefficients $a_{0,0}, a_{0,1}, a_{1,0}, a_{1,1}, a_{0,2}, a_{2,0}$

The surface, one polynomial patch at a time

Fitting a torus to a given surface is easy...

$$\kappa_{\max}, \kappa_{\min} < 0$$

$$\kappa_{\max} < 0, \kappa_{\min} > 0$$

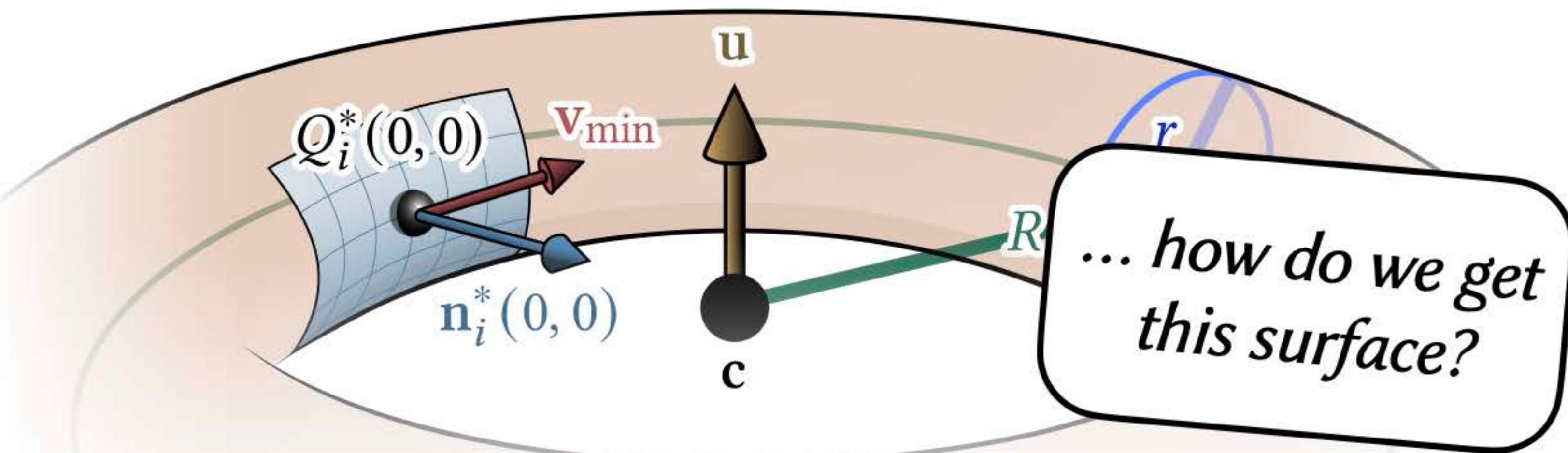


$$Q_i(s, t) = \sum_{n=0}^2 \sum_{m=0}^2 a_{n,m} s^n t^m$$

SDF of a torus:

$$\phi_{\mathbb{T}} = \|\mathbf{d}\| - r,$$

$$\mathbf{d} := (\|(\mathbf{x} - \mathbf{c}) \times \mathbf{u}\| - R, \langle \mathbf{x} - \mathbf{c}, \mathbf{u} \rangle)$$



... how do we get this surface?

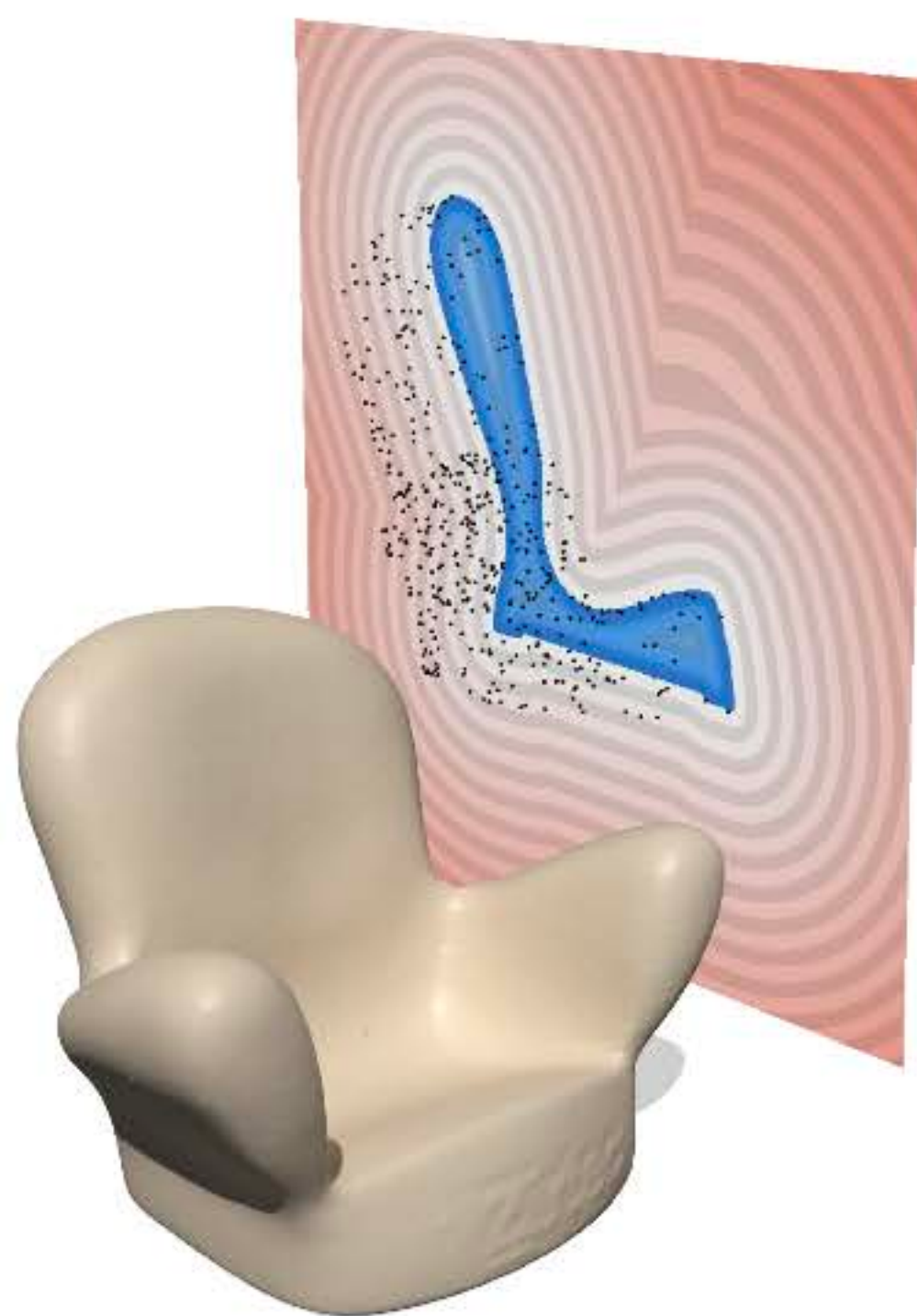
A torus is determined by the six coefficients $a_{0,0}, a_{0,1}, a_{1,0}, a_{1,1}, a_{0,2}, a_{2,0}$

Inferring the surface, even locally, is hard

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are dependent on brittle, domain-specific heuristics for neighborhood size and shape.

ground truth



neighborhood estimation with Gaussians

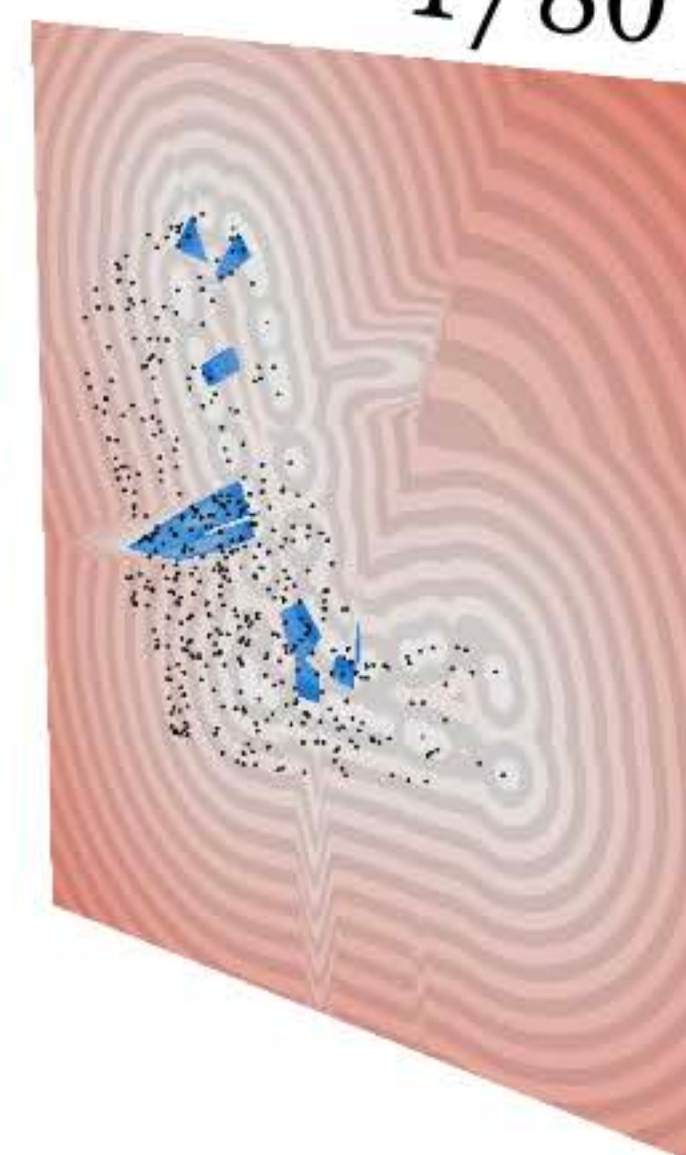
$$\sigma = 1/10$$



$$\sigma = 1/40$$



$$\sigma = 1/80$$



(polynomial estimation with weighted least squares)

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are dependent on brittle, domain-specific heuristics for neighborhood size and shape.

A **huge** bag of possible approaches:

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are dependent on brittle, domain-specific heuristics for neighborhood size and shape.

A **huge** bag of possible approaches:

- Endless kernel-based reconstruction methods (*MLS*, *partition-of-unity*, *meshless interpolation...*)

[Lancaster & Salkauskas 1981; Levin 1998; Turk & O'Brien 1999; Carr et al. 2001; Boissonnat & Cazals 2002; Amenta & Kil 2004; Shen et al. 2004; Dey & Sun 2005; Ohtake et al. 2005; Fleishman et al. 2005; Adamson & Alexa 2006; Kolluri 2008; Wang et al. 2008; Öztireli et al. 2009; Zagorchev & Goshtasby 2012; Fuhrmann & Goesele 2014; Wei et al. 2023...]

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are dependent on brittle, domain-specific heuristics for neighborhood size and shape.

A **huge** bag of possible approaches:

- Endless kernel-based reconstruction methods (*MLS*, *partition-of-unity*, *meshless interpolation...*)
- Anisotropic kernels (*PCA*, *quadric error metric*, *etc.*) → *brittle*
- Hierarchy → *brittle*

[Lancaster & Salkauskas 1981; Levin 1998; Turk & O'Brien 1999; Carr et al. 2001; Boissonnat & Cazals 2002; Amenta & Kil 2004; Shen et al. 2004; Dey & Sun 2005; Ohtake et al. 2005; Fleishman et al. 2005; Adamson & Alexa 2006; Kolluri 2008; Wang et al. 2008; Öztireli et al. 2009; Zagorchev & Goshtasby 2012; Fuhrmann & Goesele 2014; Wei et al. 2023...]

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are dependent on brittle, domain-specific heuristics for neighborhood size and shape.

A **huge** bag of possible approaches:

- Endless kernel-based reconstruction methods (*MLS*, *partition-of-unity*, *meshless interpolation...*)
- Anisotropic kernels (*PCA*, *quadric error metric*, *etc.*) → *brittle*
- Hierarchy → *brittle*
- Iterative/nonlinear solvers → *slow*
- Robust statistics → *slow*

[Lancaster & Salkauskas 1981; Levin 1998; Turk & O'Brien 1999; Carr et al. 2001; Boissonnat & Cazals 2002; Amenta & Kil 2004; Shen et al. 2004; Dey & Sun 2005; Ohtake et al. 2005; Fleishman et al. 2005; Adamson & Alexa 2006; Kolluri 2008; Wang et al. 2008; Öztireli et al. 2009; Zagorchev & Goshtasby 2012; Fuhrmann & Goesele 2014; Wei et al. 2023...]

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are dependent on brittle, domain-specific heuristics for neighborhood size and shape.

A **huge** bag of possible approaches:

- Endless kernel-based reconstruction methods (*MLS*, *partition-of-unity*, *meshless interpolation...*)
- Anisotropic kernels (*PCA*, *quadric error metric*, *etc.*) → *brittle*
- Hierarchy → *brittle*
- Iterative/nonlinear solvers → *slow*
- Robust statistics → *slow*
- ... *any one method is unlikely to generalize!*

[Lancaster & Salkauskas 1981; Levin 1998; Turk & O'Brien 1999; Carr et al. 2001; Boissonnat & Cazals 2002; Amenta & Kil 2004; Shen et al. 2004; Dey & Sun 2005; Ohtake et al. 2005; Fleishman et al. 2005; Adamson & Alexa 2006; Kolluri 2008; Wang et al. 2008; Öztireli et al. 2009; Zagorchev & Goshtasby 2012; Fuhrmann & Goesele 2014; Wei et al. 2023...]

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are extremely dependent on brittle, domain-specific heuristics for neighborhood size and shape.

A huge bag of possible approaches:

- Endless kernel-based reconstruction methods (*MLS*, *partition of unity*, *meshless interpolation*...)
Just use a *fixed-size* neighborhood with a *data-driven* approach.
- Anisotropic kernels (*PCA*, *quadric error metric*, etc.) $\rightarrow$ brittle
- Hierarchy $\rightarrow$ brittle
- Iterative/nonlinear solvers $\rightarrow$ slow
- Robust statistics $\rightarrow$ slow
- ...

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are extremely dependent on brittle, domain-specific heuristics for neighborhood size and shape.

Just use a *fixed-size* neighborhood with a *data-driven* approach.

A huge bag of possible approaches:

local function

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

- Endless kernel-based reconstruction methods (MLS, partition-of-unity, meshless, ...)
- Anisotropic kernels (PCA, quadratic error metrics) $\rightarrow$ brittle
- Hierarchy $\rightarrow$ brittle
- Dominated by local surface behavior
- Robust statistics $\rightarrow$ slow
- ...

Inferring the surface, even locally, is hard

Classic methods like *jet-fitting*, *moving least squares* are extremely dependent on brittle, domain-specific heuristics for neighborhood size and shape.

Just use a *fixed-size* neighborhood with a *data-driven* approach.

A huge bag of possible approaches:

local function

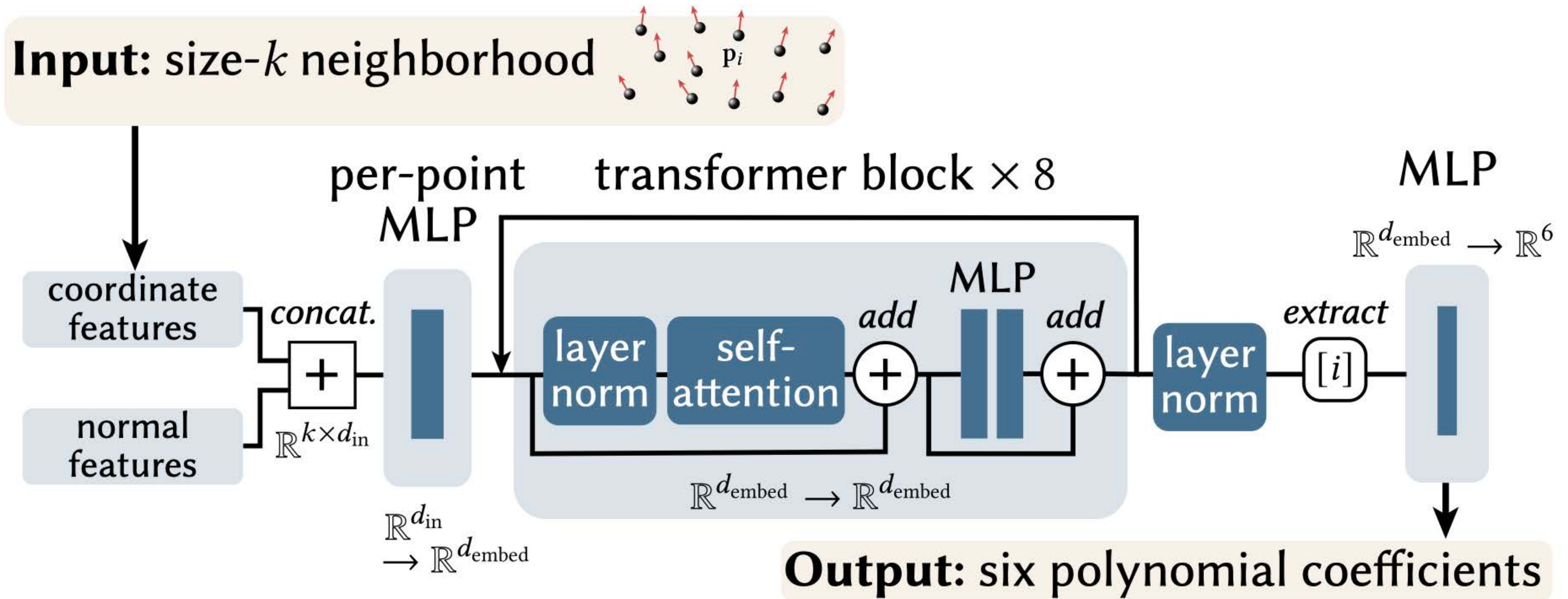
$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

- Endless kernel-based reconstruction methods (MLS, partition-of-unity, meshless, ...)
- Anisotropic kernels (PCA, quadratic error metrics) $\rightarrow$ brittle
- Hierarchy $\rightarrow$ brittle
- Dominated by local surface behavior
- Easy to generate data for small surface patches

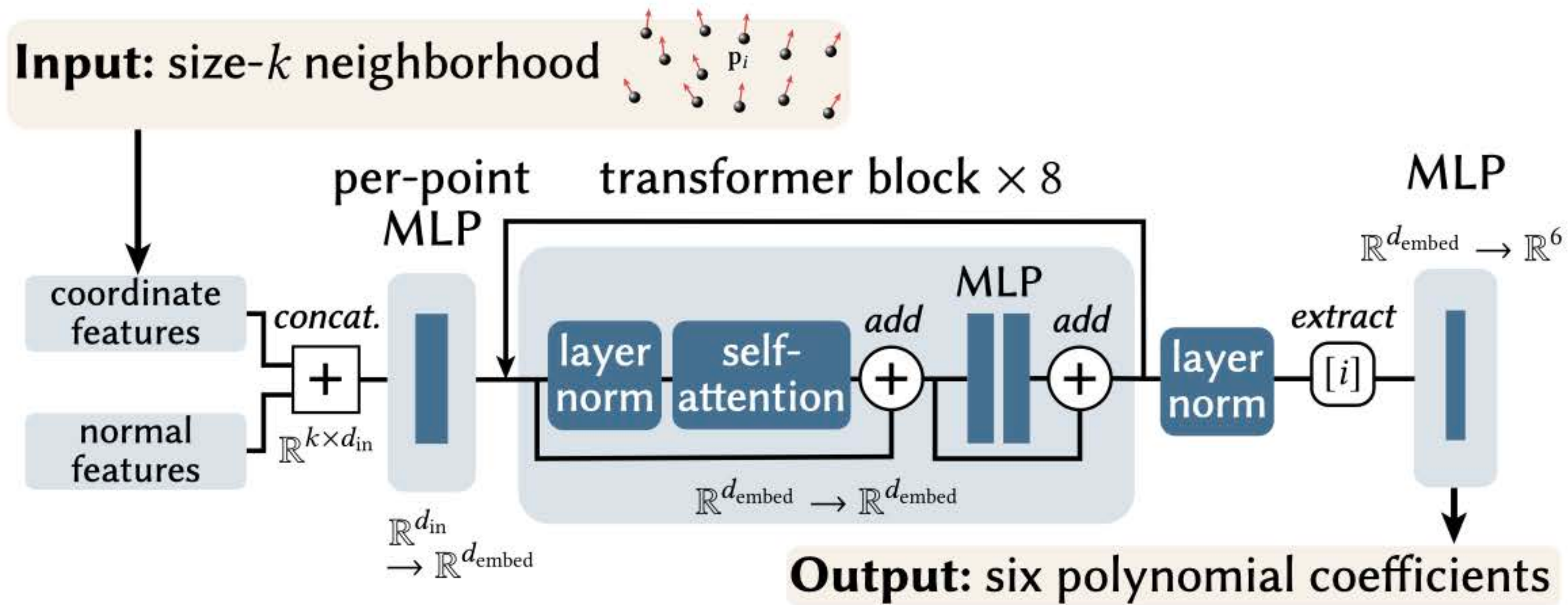
• ...

Network architecture

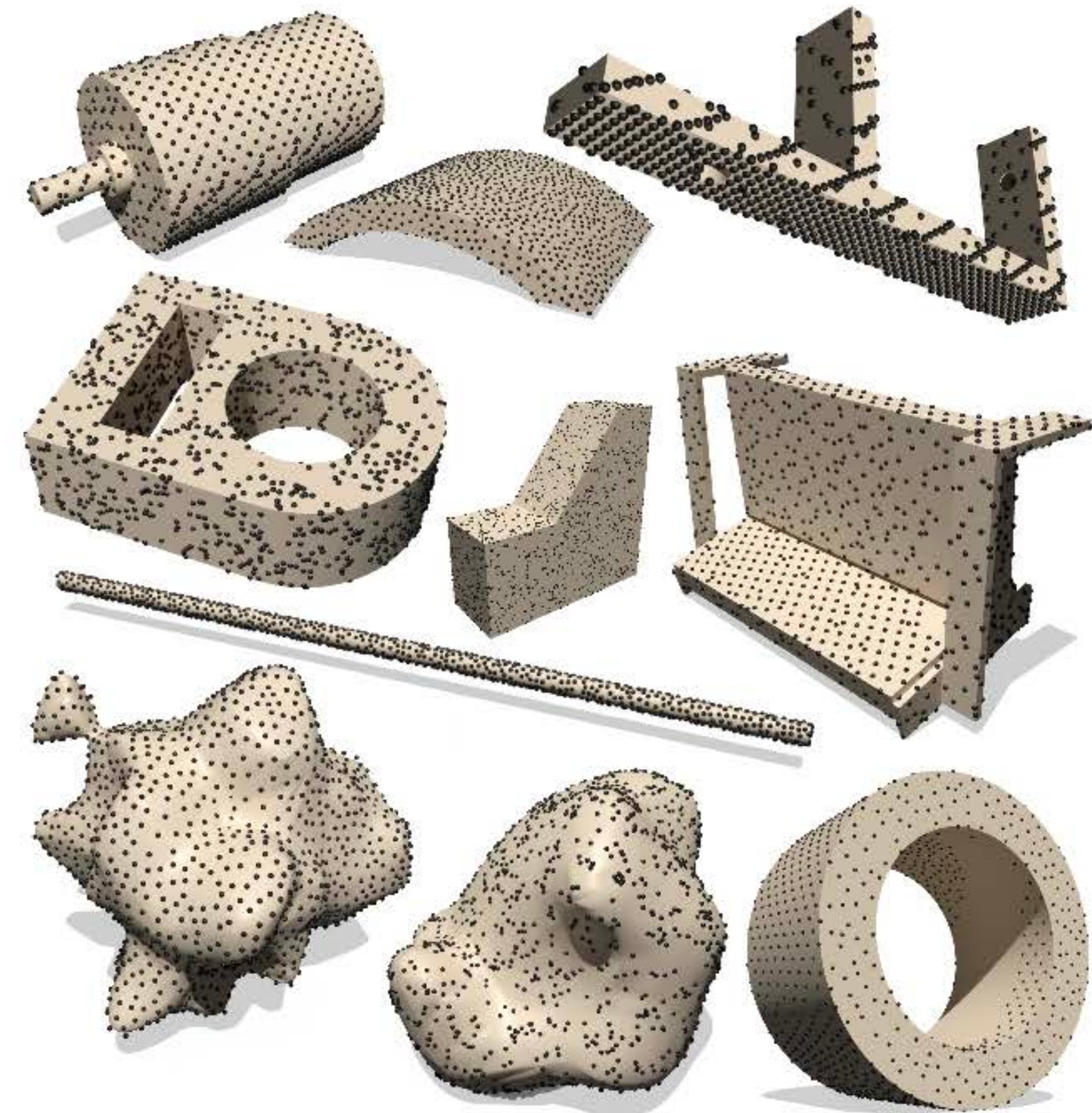
Network architecture



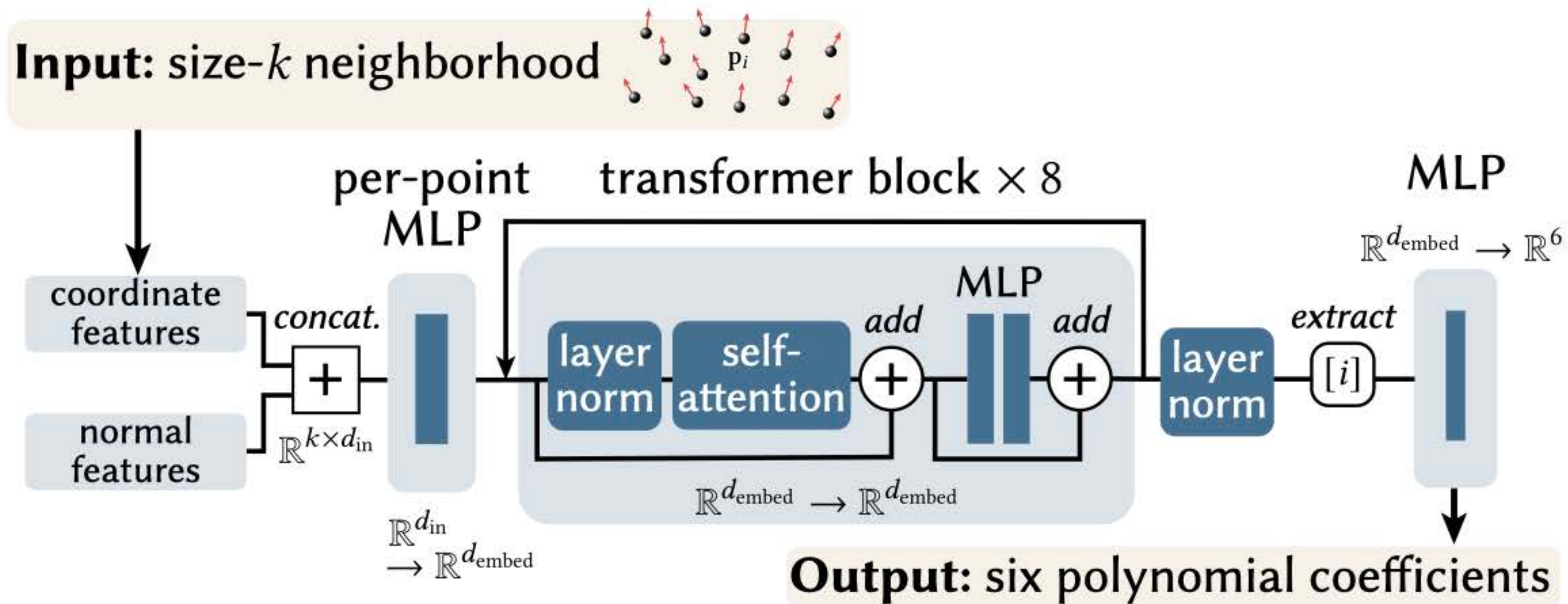
Network architecture



Training data

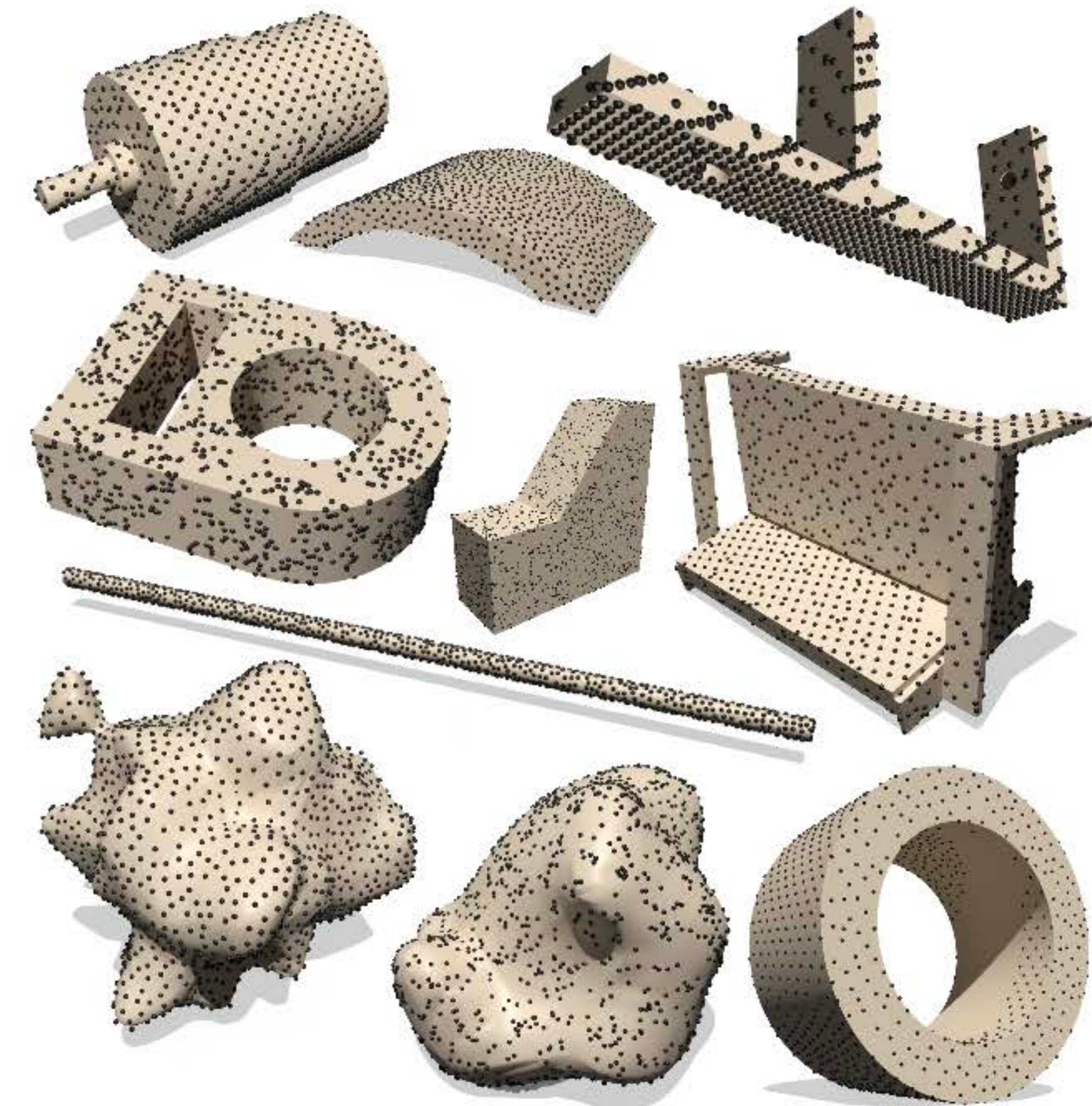


Network architecture

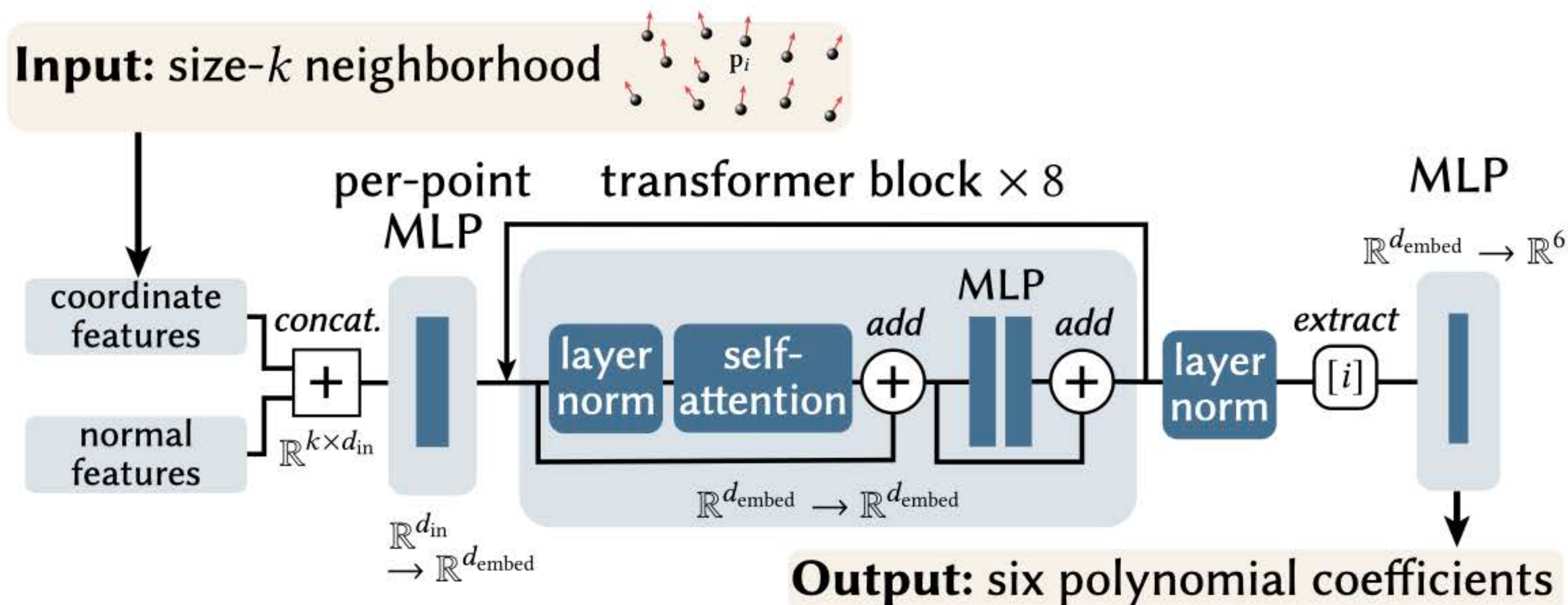


- 10,000 models from ABC dataset [Koch et al. 2019], 2000 procedural

Training data

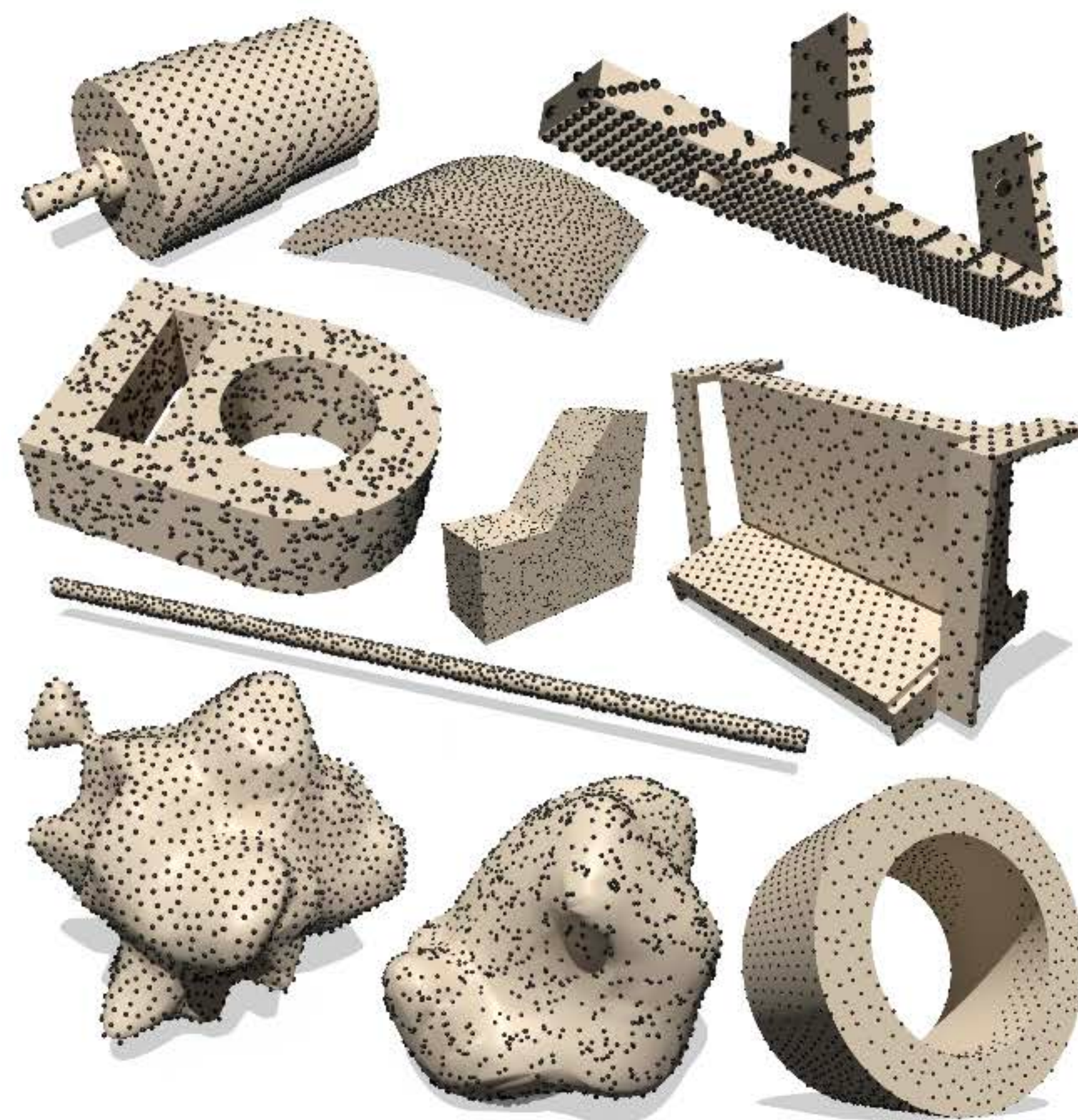


Network architecture

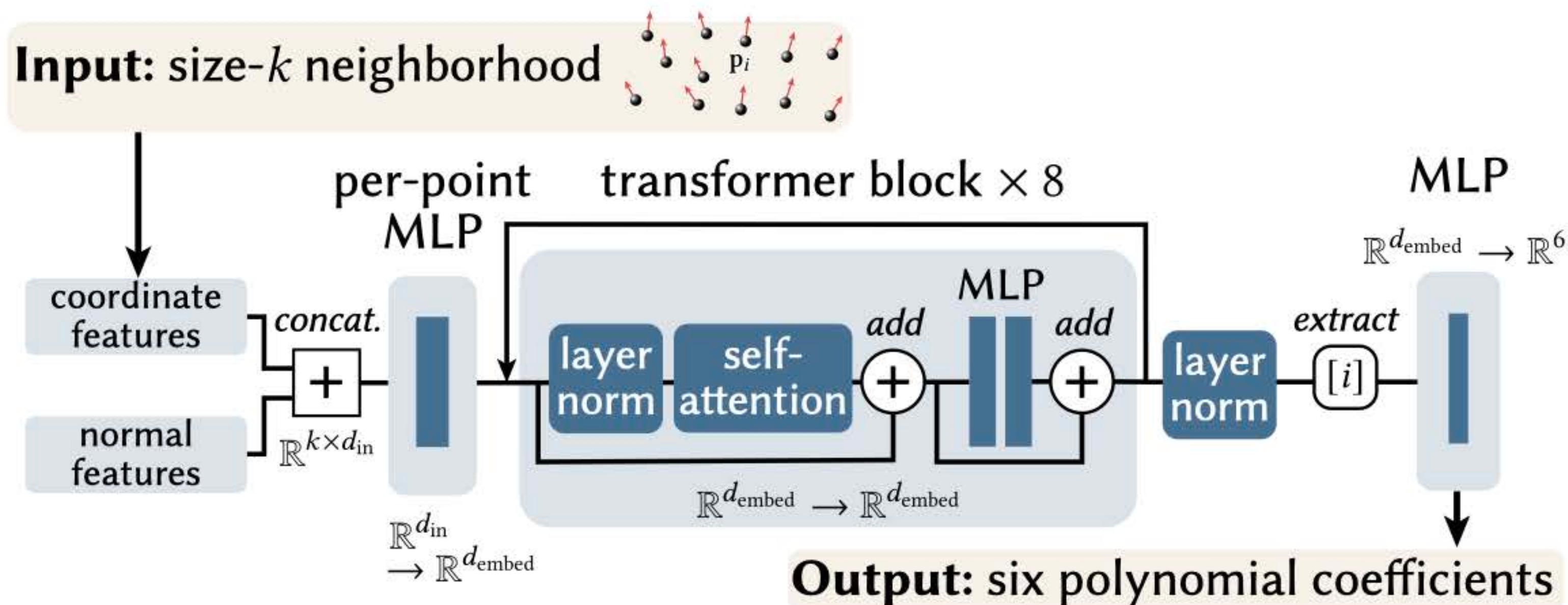


- 10,000 models from ABC dataset [Koch et al. 2019], 2000 procedural
- Normalize and sample into size-2048 point clouds:
 - Uniform sampling
 - Farthest-point sampling
 - Simulated scanning

Training data

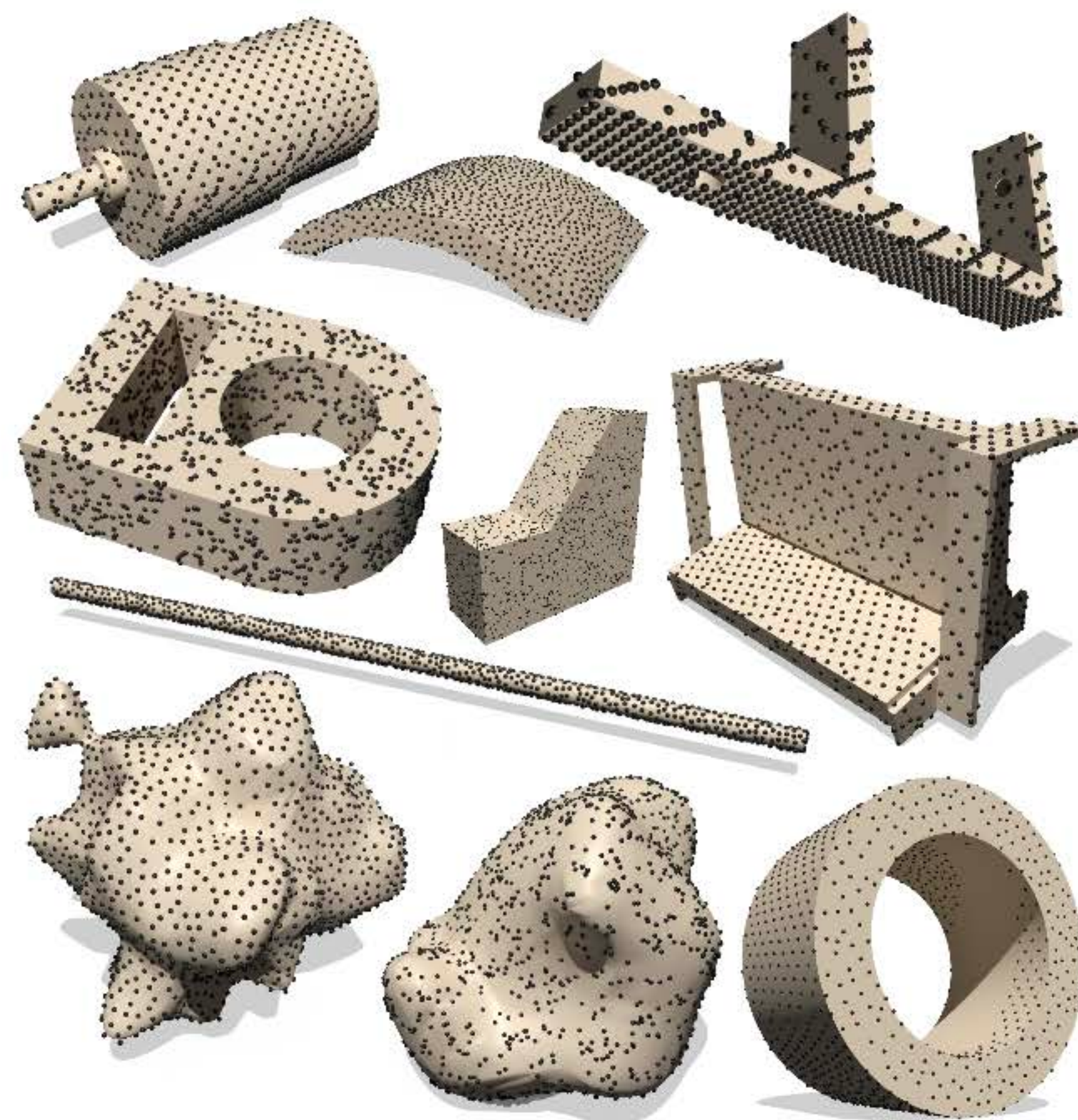


Network architecture

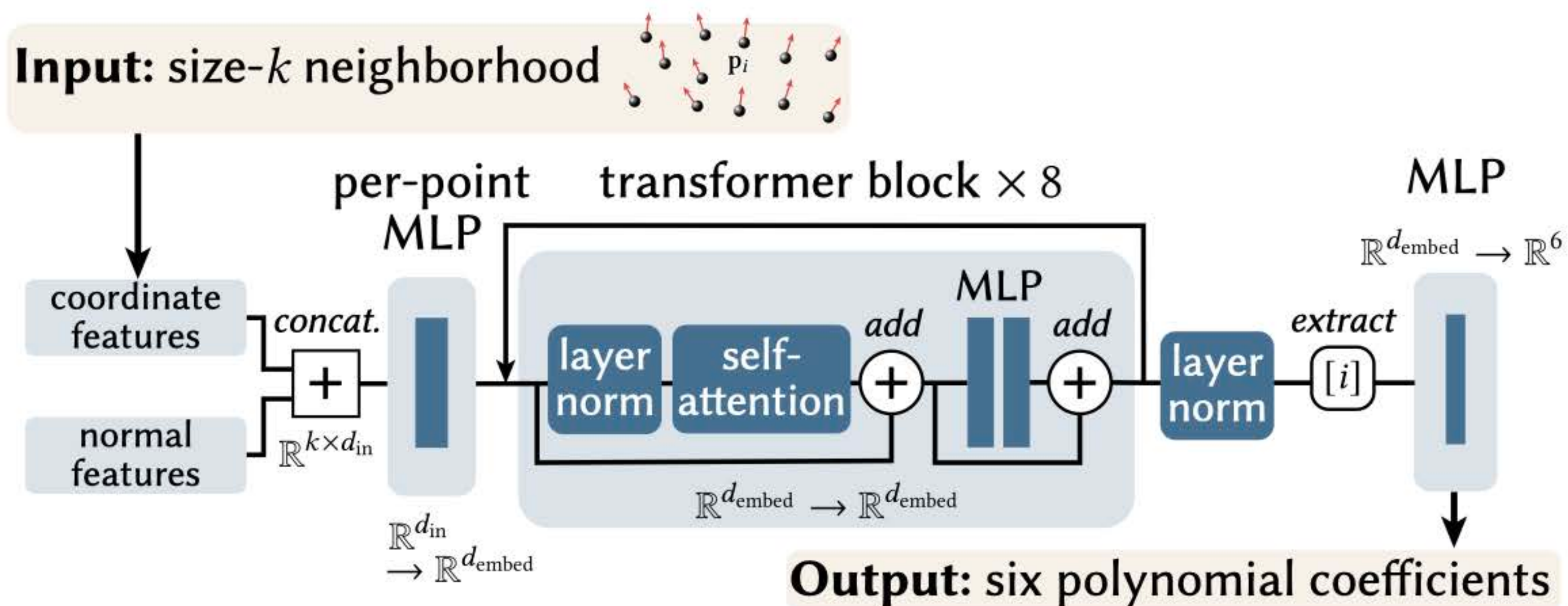


- 10,000 models from ABC dataset [Koch et al. 2019], 2000 procedural
- Normalize and sample into size-2048 point clouds:
 - Uniform sampling
 - Farthest-point sampling
 - Simulated scanning
- Input neighborhood of size $k = 64$
- >40 million training examples

Training data

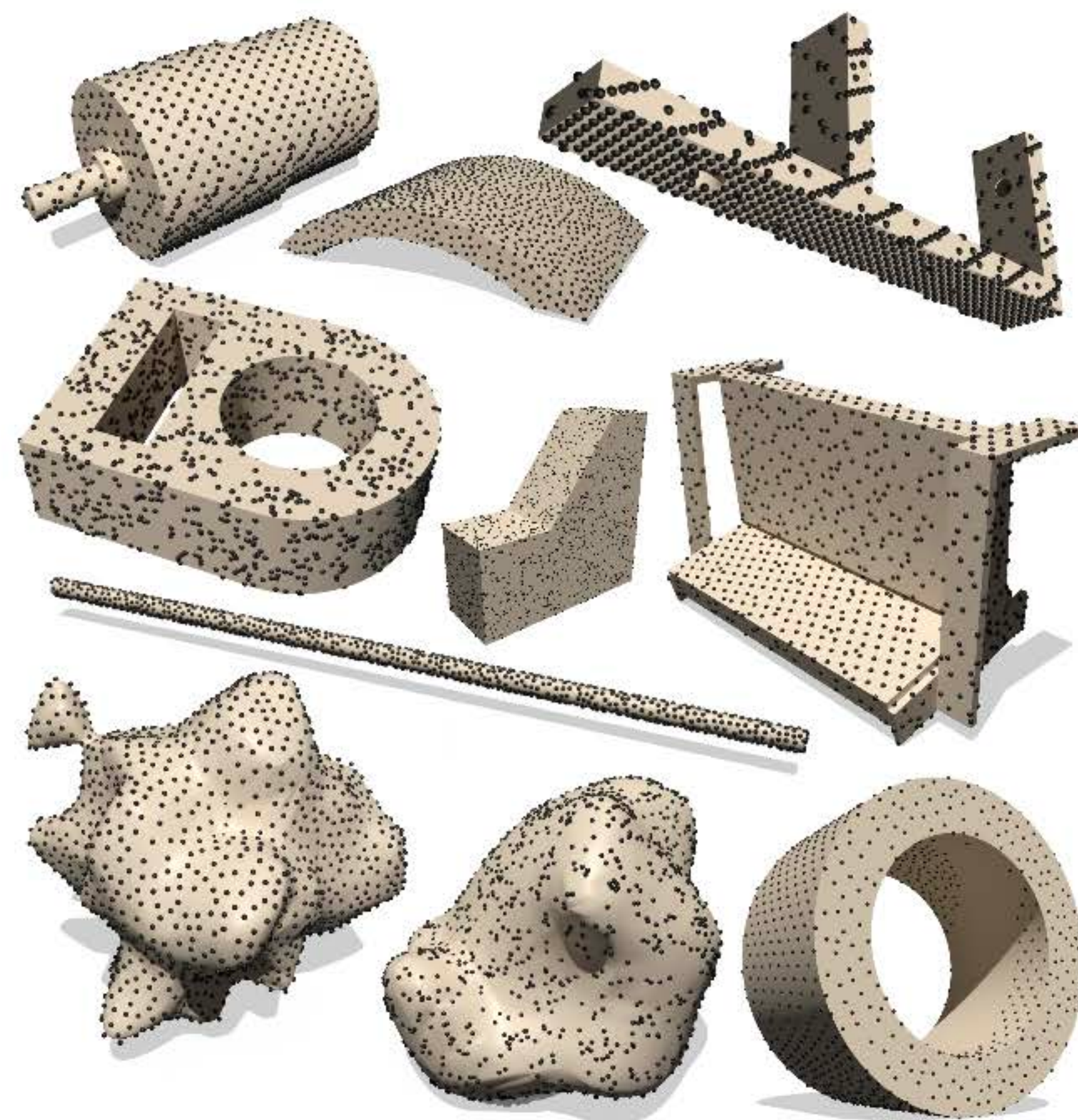


Network architecture

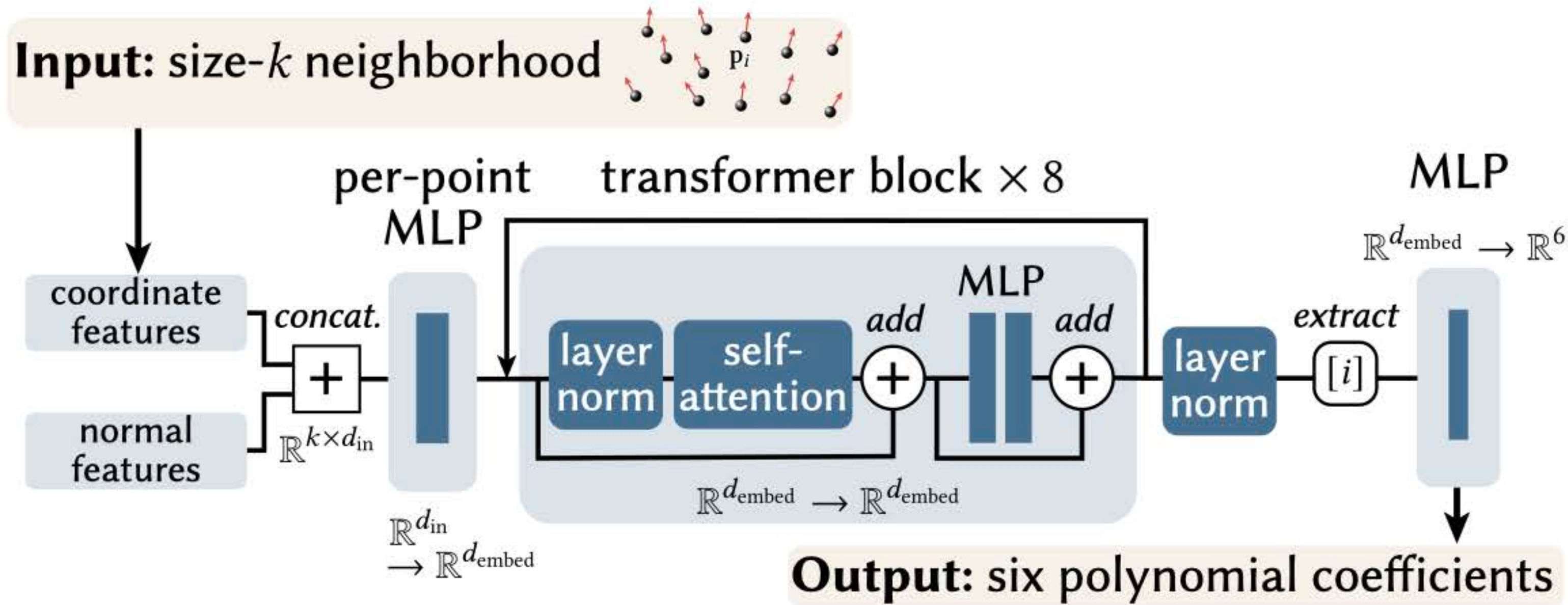


Loss function

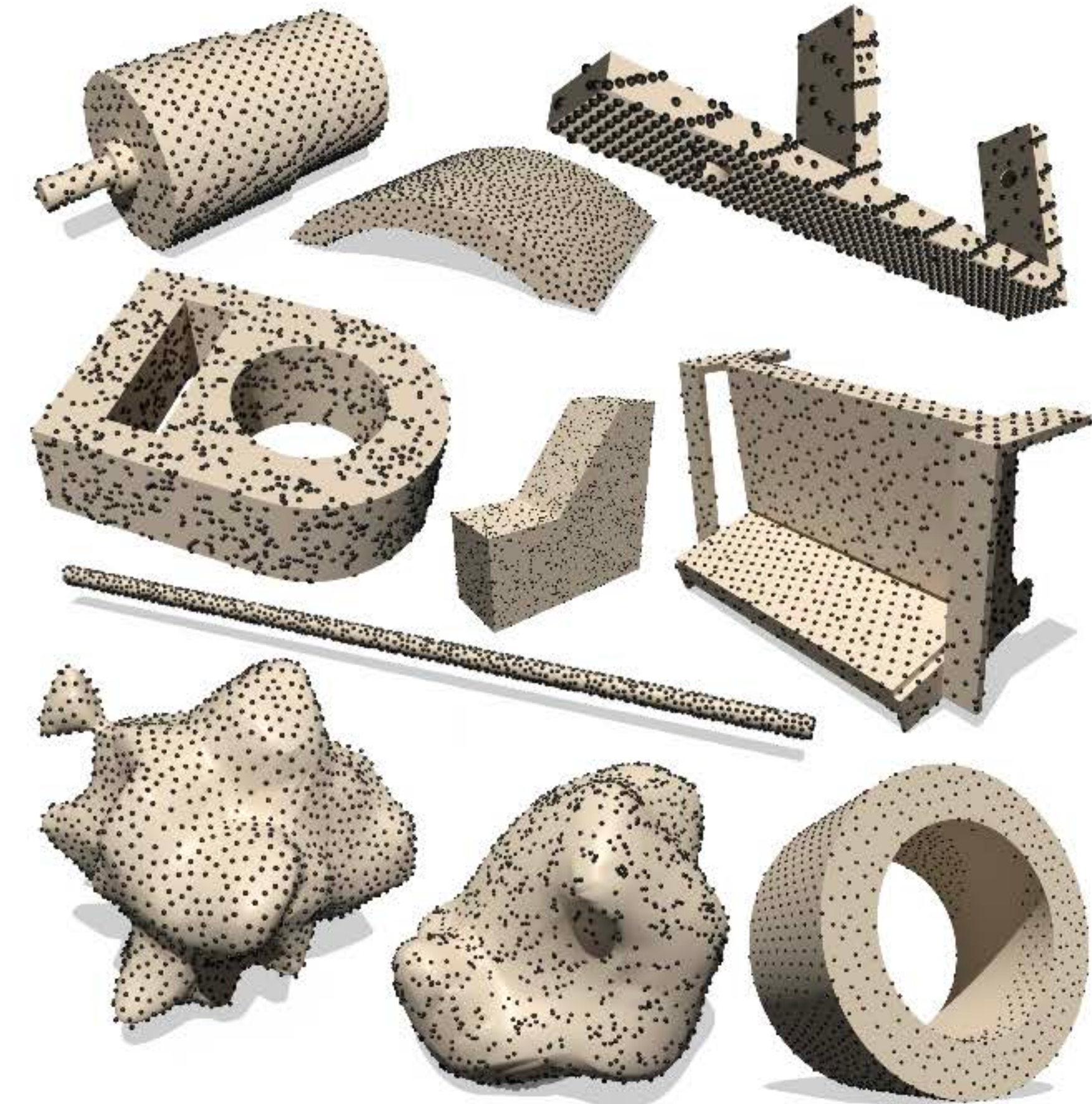
Training data



Network architecture



Training data

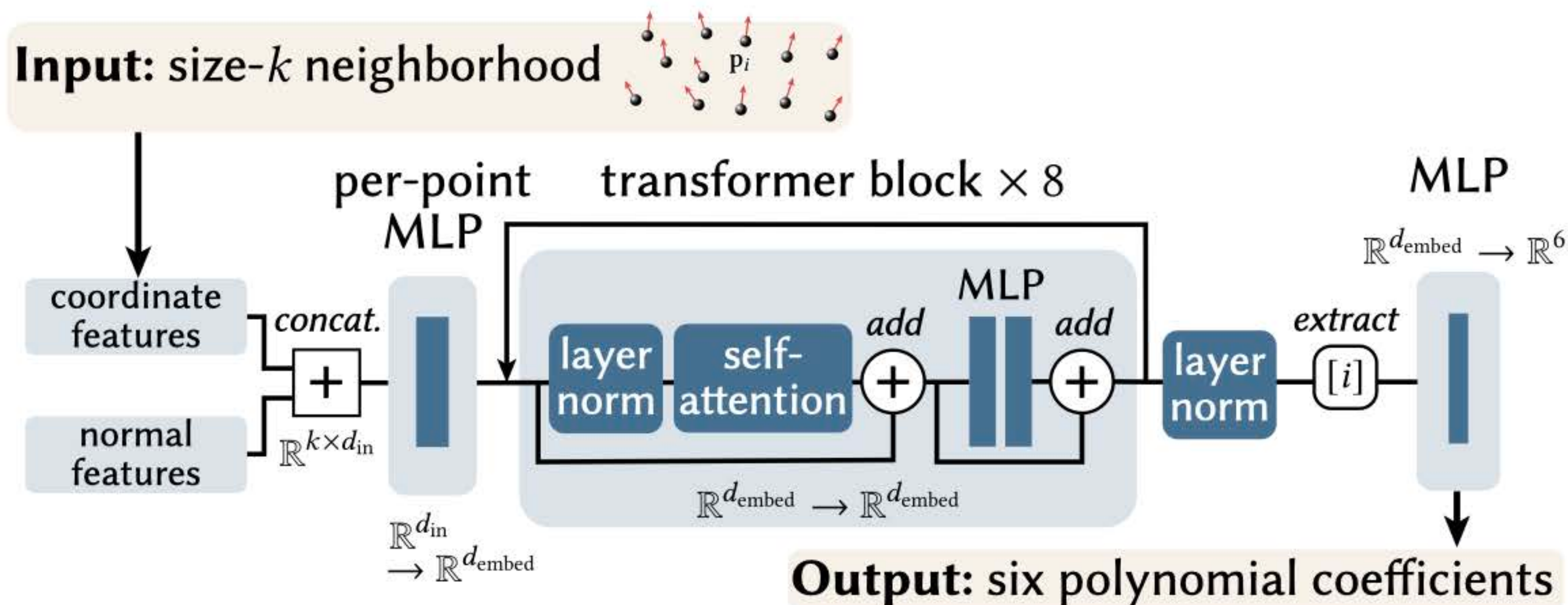


Loss function

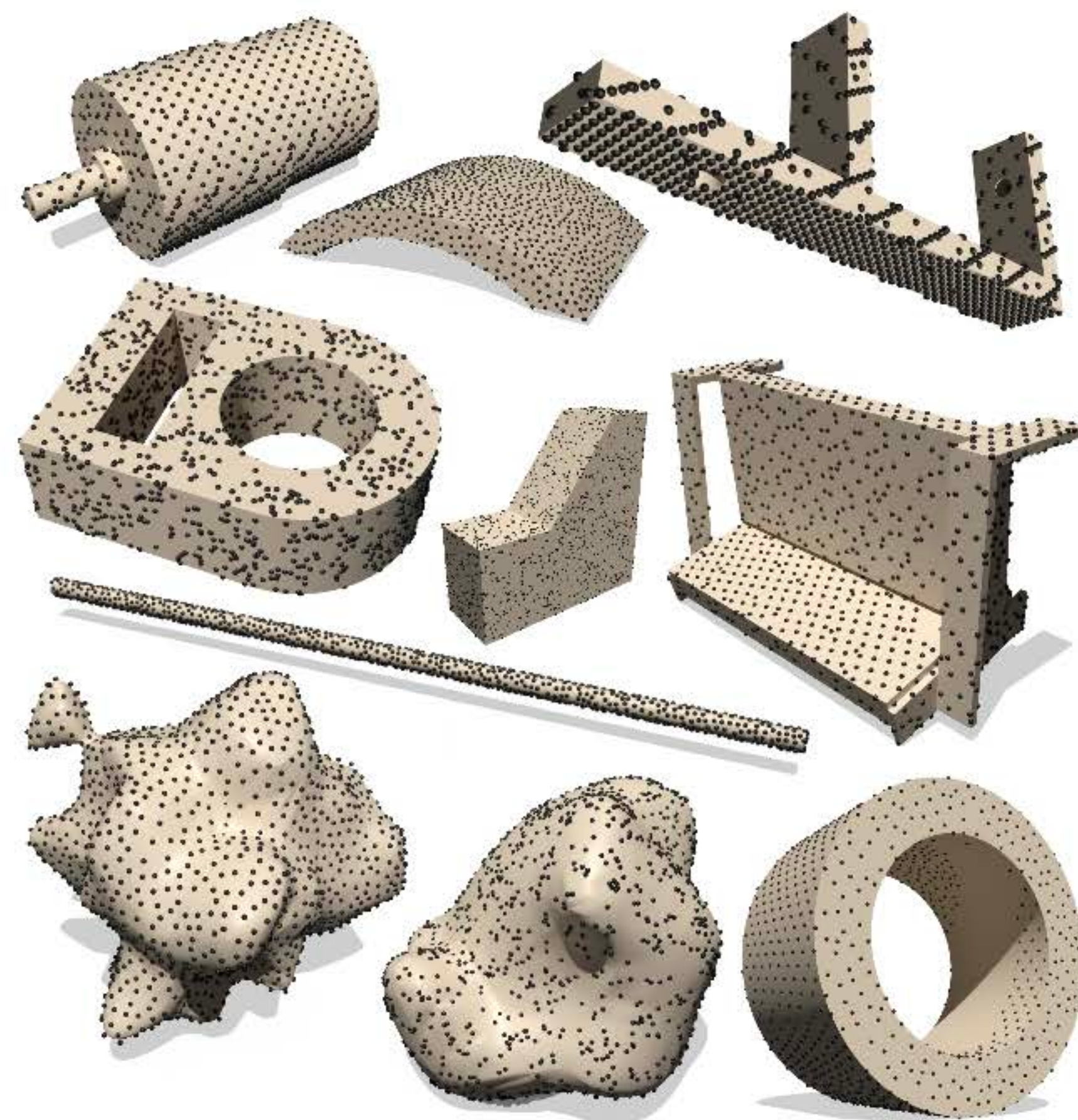
$$\mathcal{L} := \mathcal{L}_{\text{distance}}$$

$$\text{where } \mathcal{L}_{\text{distance}} := \frac{1}{Q} \sum_{j=1}^Q |\phi(\mathbf{q}_j) - \phi_{\text{true}}(\mathbf{q}_i)|$$

Network architecture



Training data



Loss function

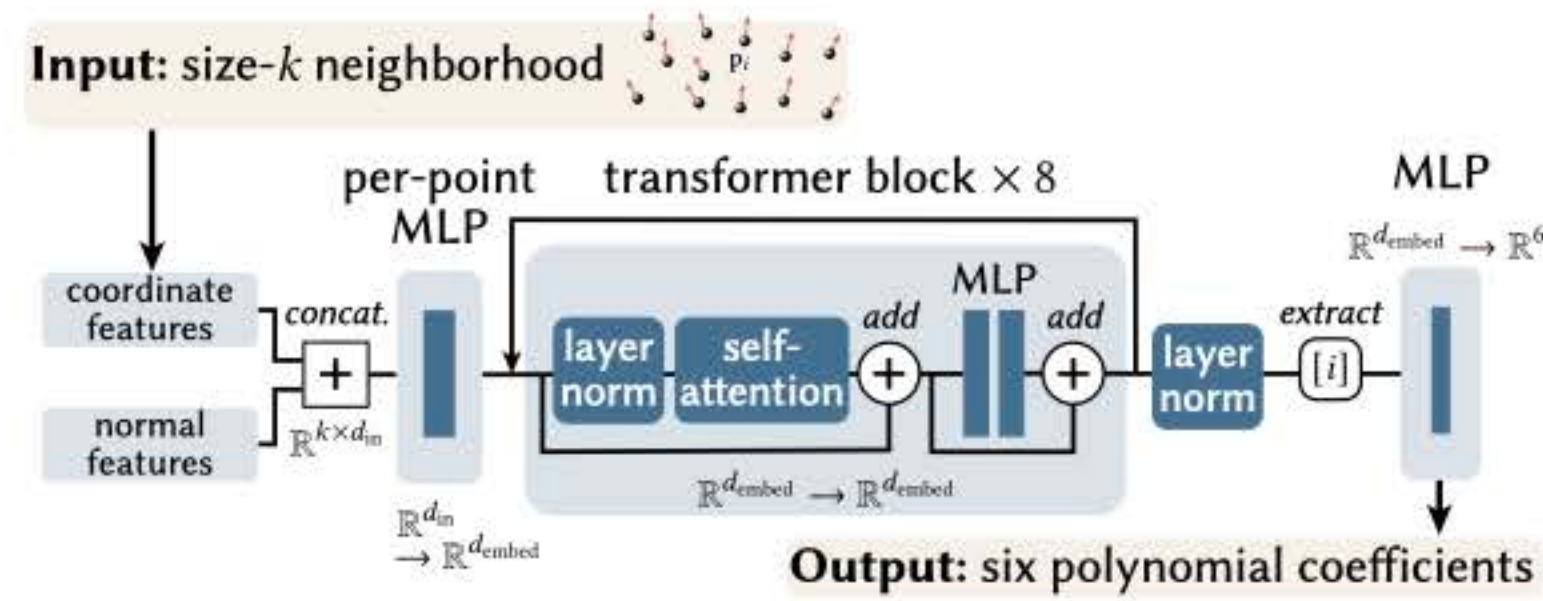
$$\mathcal{L} := \mathcal{L}_{\text{distance}} + \mathcal{L}_{\text{blend}}$$

$$\text{where } \mathcal{L}_{\text{distance}} := \frac{1}{Q} \sum_{j=1}^Q |\phi(\mathbf{q}_j) - \phi_{\text{true}}(\mathbf{q}_i)|$$

$$\mathcal{L}_{\text{blend}} := \frac{1}{Q} \sum_{j=1}^Q |1 - \|\nabla \phi(\mathbf{q}_j)\||$$

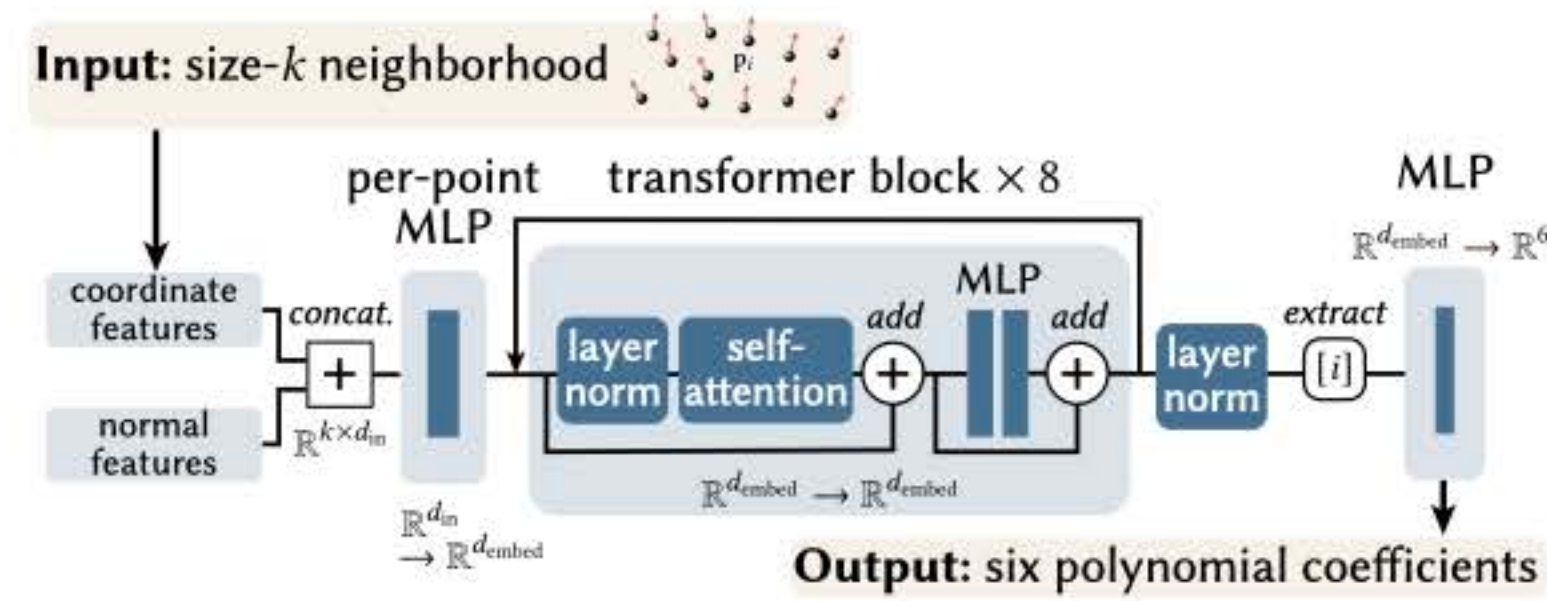
The *points as tori (PAT)* algorithm

Pre-training



The *points as tori (PAT)* algorithm

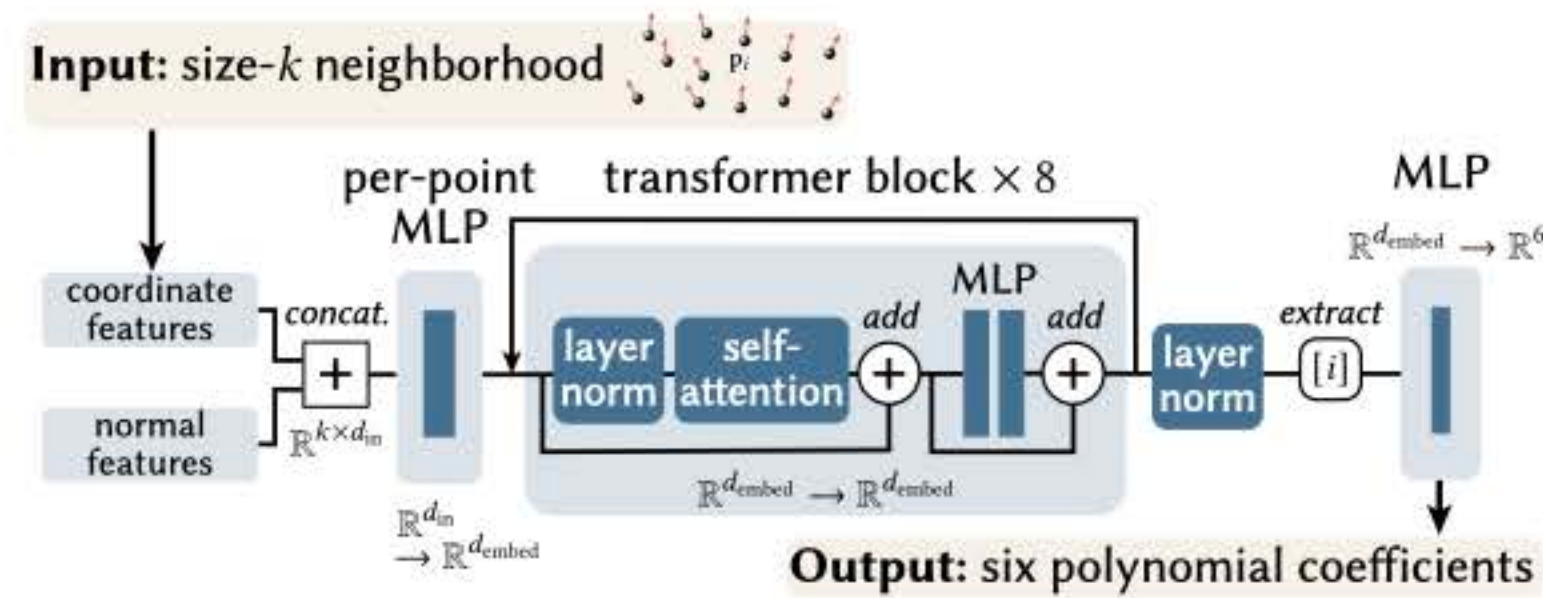
Pre-training



~45 hours on a single NVIDIA RTX 3090

The *points as tori (PAT)* algorithm

Pre-training

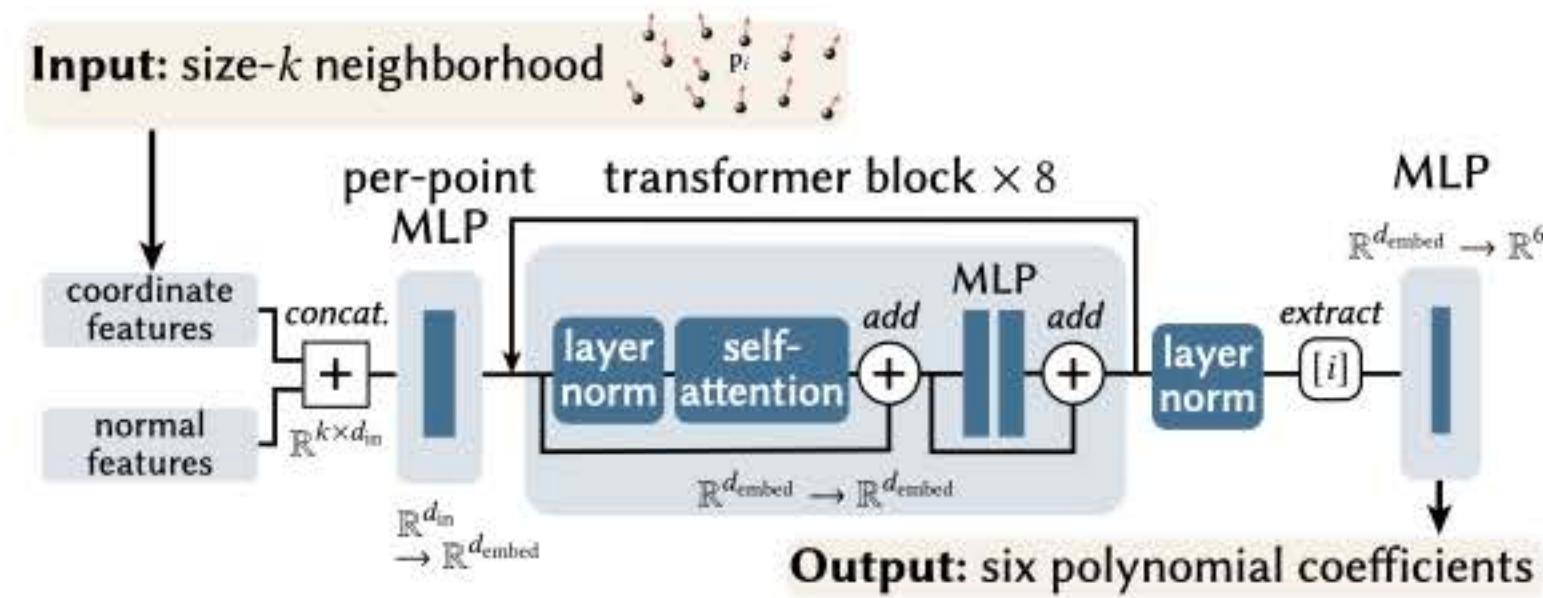


~45 hours on a single NVIDIA RTX 3090

Given an input point cloud $P = \{\mathbf{p}_i, \mathbf{n}_i\}_{i=1}^{|P|} \dots$

The *points as tori (PAT)* algorithm

Pre-training



~45 hours on a single NVIDIA RTX 3090

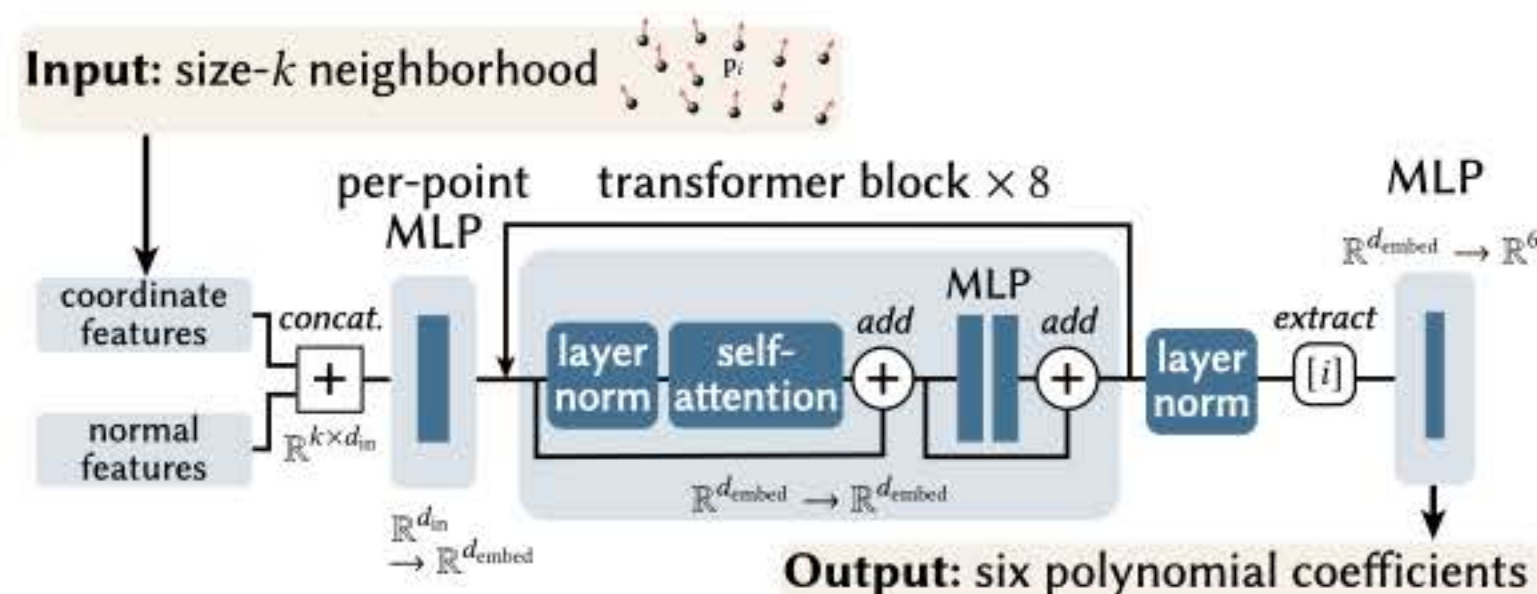
Given an input point cloud $P = \{\mathbf{p}_i, \mathbf{n}_i\}_{i=1}^{|P|} \dots$

One-time precompute:

Apply pre-trained network to each point neighborhood (parallelizable)

The *points as tori (PAT)* algorithm

Pre-training



~45 hours on a single NVIDIA RTX 3090

Given an input point cloud $P = \{\mathbf{p}_i, \mathbf{n}_i\}_{i=1}^{|P|} \dots$

One-time precompute:

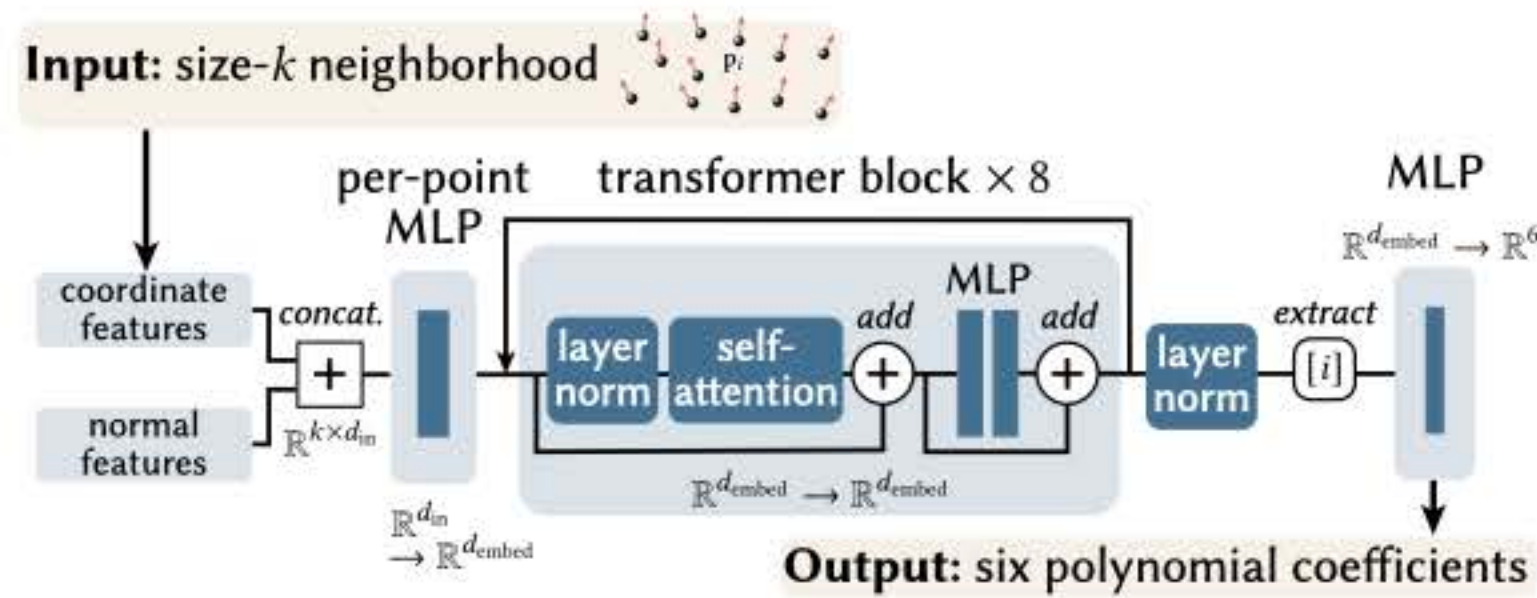
Apply pre-trained network to each point neighborhood (parallelizable)

Inference time:

Evaluate the SDF at $\mathbf{x} \in \mathbb{R}^3$ as
$$\phi(\mathbf{x}) = \frac{\sum_{i=1}^{|P|} g_i(\mathbf{x}) \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}$$

The *points as tori* (PAT) algorithm

Pre-training



~45 hours on a single NVIDIA RTX 3090

Given an input point cloud $P = \{\mathbf{p}_i, \mathbf{n}_i\}_{i=1}^{|P|} \dots$

One-time precompute:

Apply pre-trained network to each point neighborhood (parallelizable)

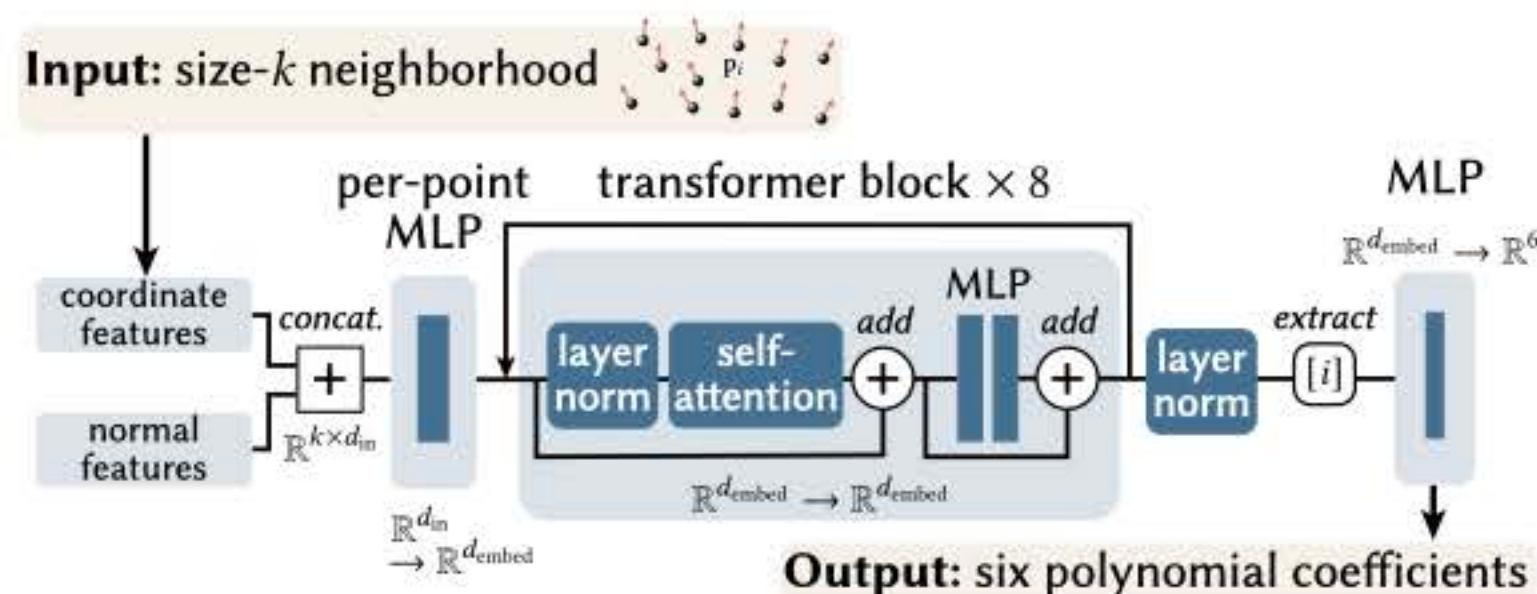
Inference time:

Evaluate the SDF at $\mathbf{x} \in \mathbb{R}^3$ as
$$\phi(\mathbf{x}) = \frac{\sum_{i=1}^{|P|} g_i(\mathbf{x}) \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}$$

$10^{-4} - 10^{-3}$ seconds for a single query to 1M-10M+ points

The *points as tori (PAT)* algorithm

Pre-training



~45 hours on a single NVIDIA RTX 3090

Given an input point cloud $P = \{\mathbf{p}_i, \mathbf{n}_i\}_{i=1}^{|P|} \dots$

One-time precompute:

Apply pre-trained network to each point neighborhood (parallelizable)

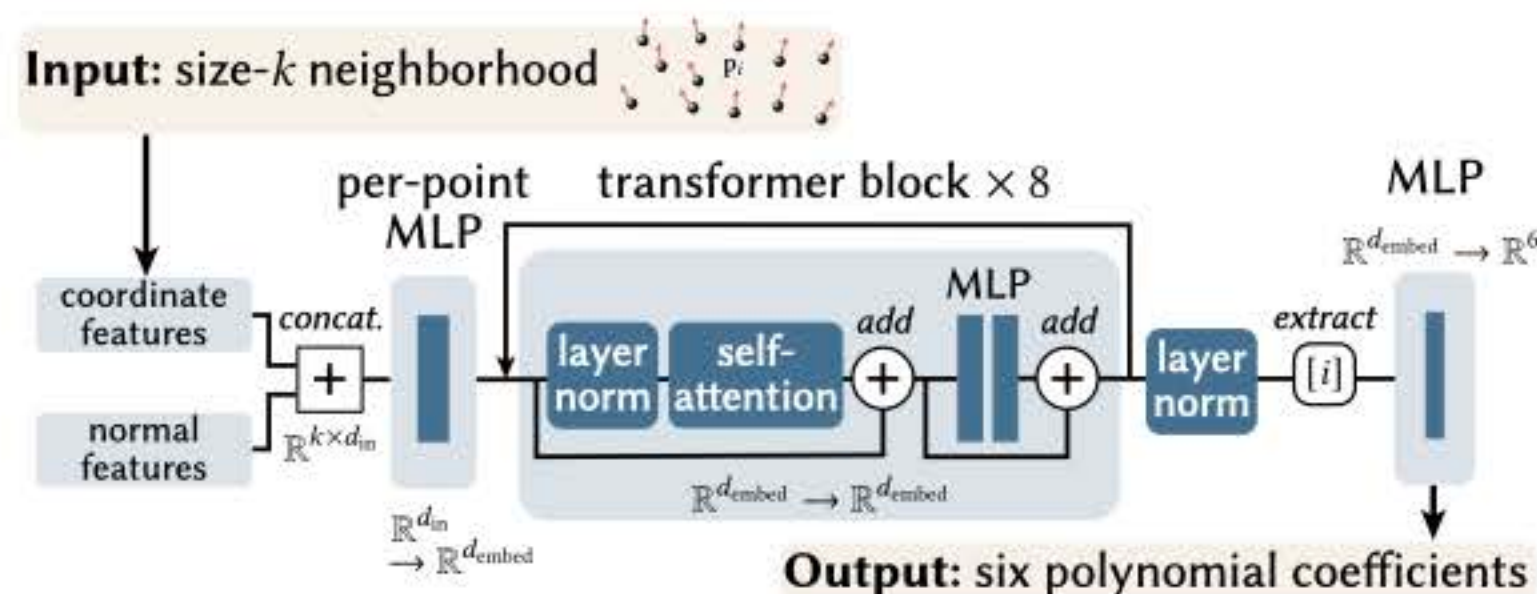
Inference time:

Evaluate the SDF at $\mathbf{x} \in \mathbb{R}^3$ as
$$\phi(\mathbf{x}) = \frac{\sum_{i=1}^{|P|} g_i(\mathbf{x}) \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}$$

*$10^{-4} - 10^{-3}$ seconds for a single query to 1M-10M+ points
multiple queries trivially parallelizable*

The *points as tori (PAT)* algorithm

Pre-training



~45 hours on a single NVIDIA RTX 3090

Given an input point cloud $P = \{\mathbf{p}_i, \mathbf{n}_i\}_{i=1}^{|P|} \dots$

One-time precompute:

Apply pre-trained network to each point neighborhood (parallelizable)

Inference time:

Evaluate the SDF at $\mathbf{x} \in \mathbb{R}^3$ as $\phi(\mathbf{x}) = \frac{\sum_{i=1}^{|P|} g_i(\mathbf{x}) \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|\mathbf{x} - \mathbf{p}_i\|)}$

simple
formula
for λ

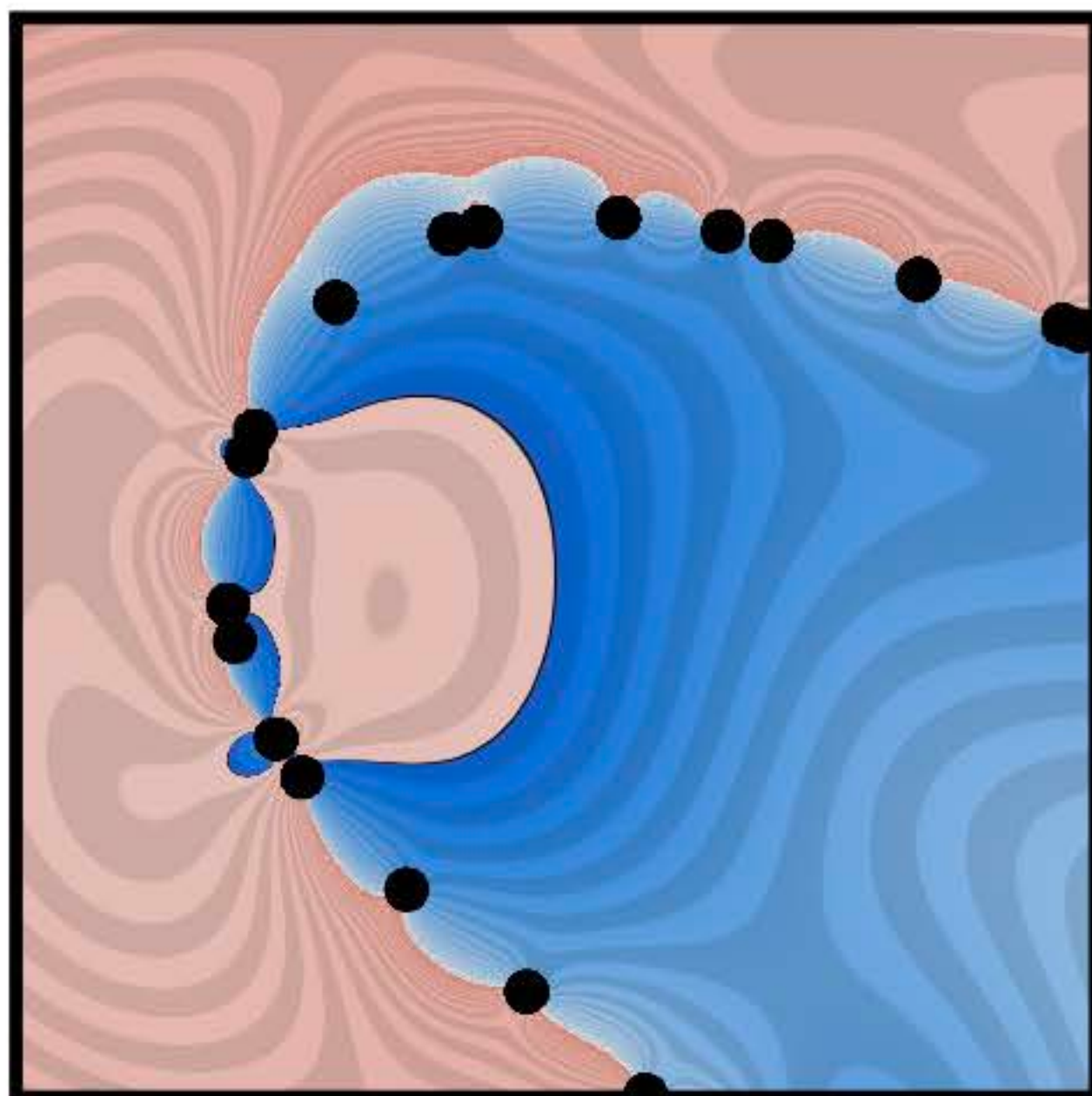
*$10^{-4} - 10^{-3}$ seconds for a single query to 1M-10M+ points
multiple queries trivially parallelizable*

Better than naive convolutional distance

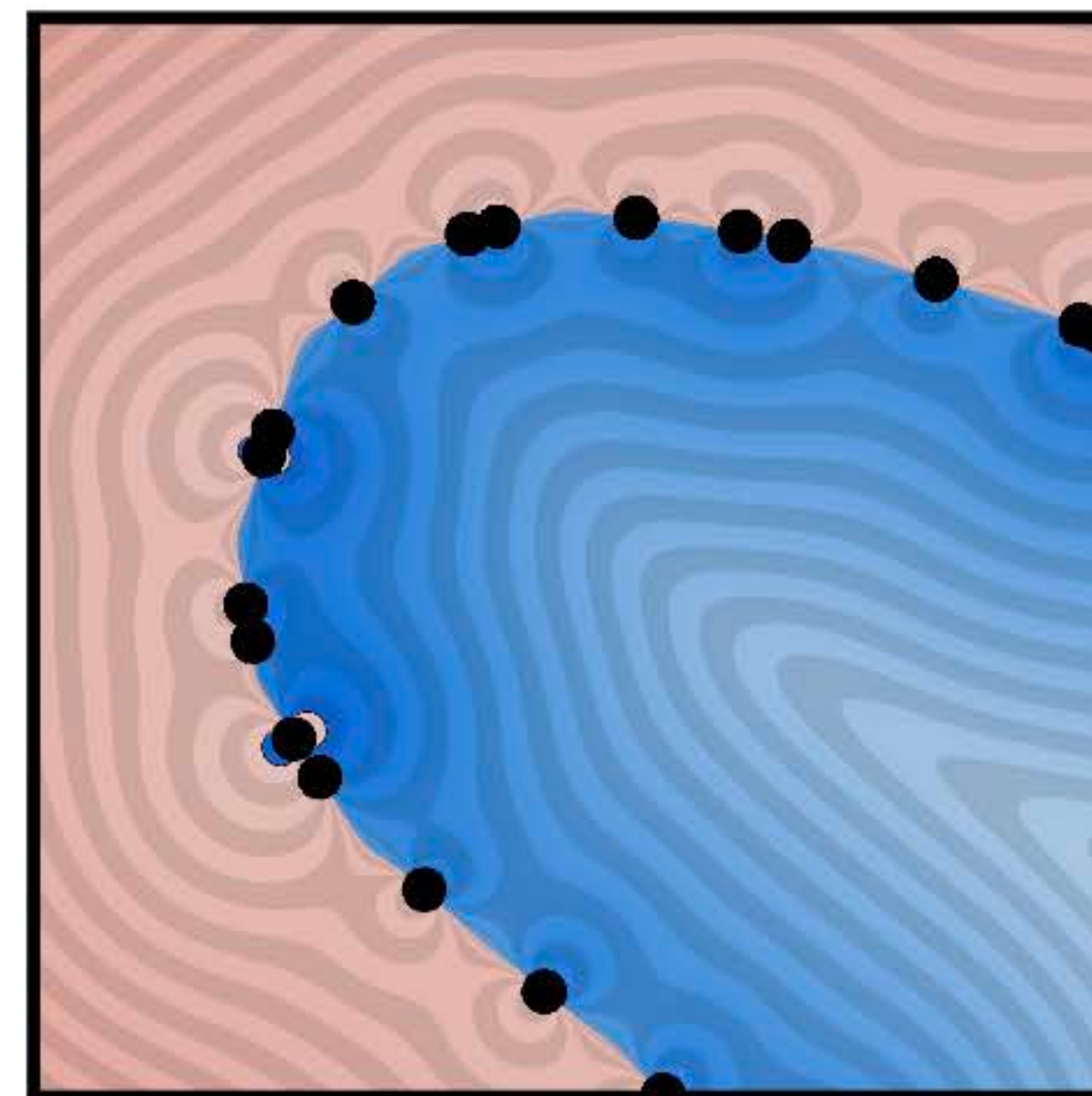
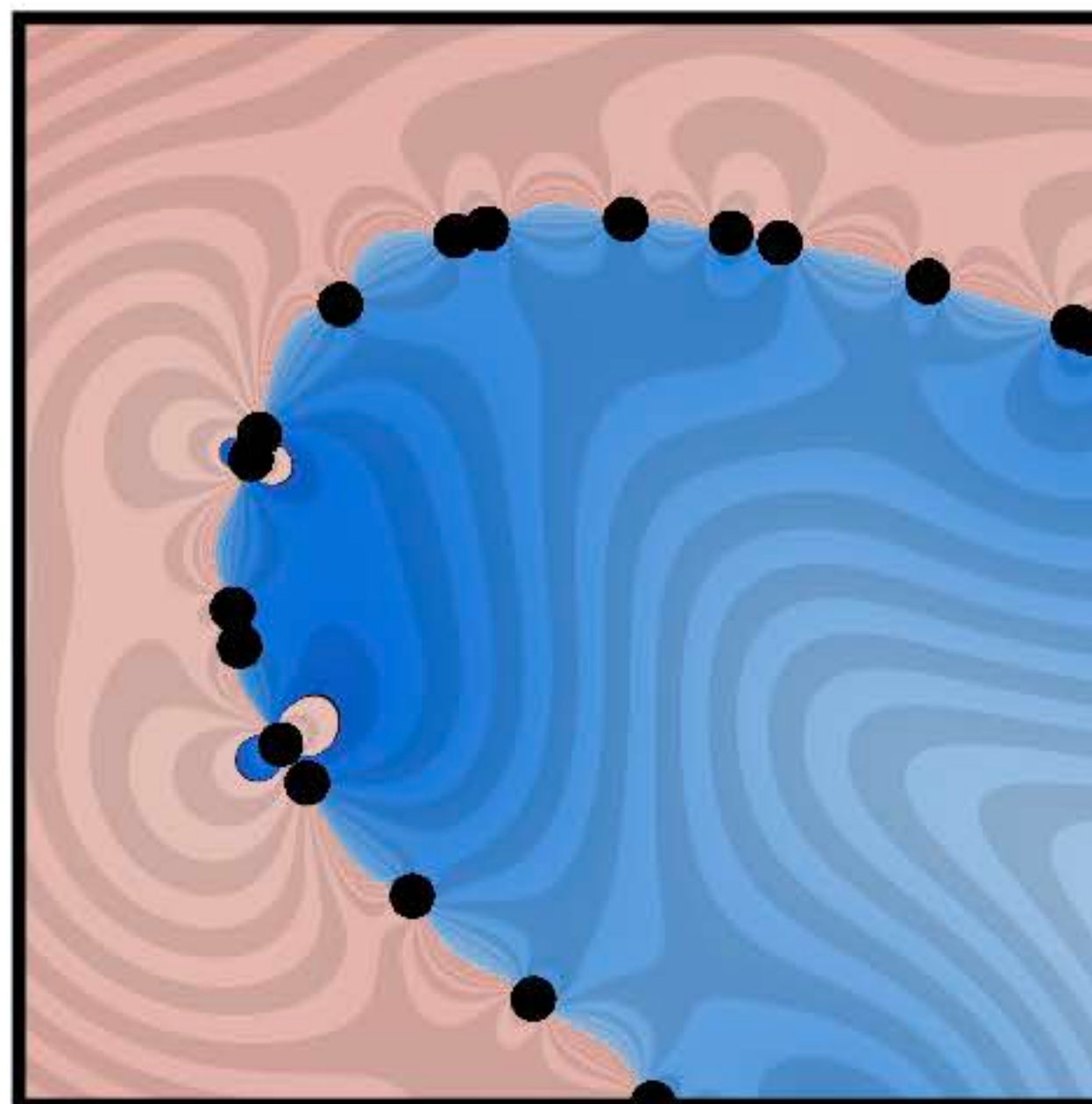
smaller λ



larger λ



worse eikonality
underfit level sets



better eikonality
overfit level sets

Better than naive convolutional distance

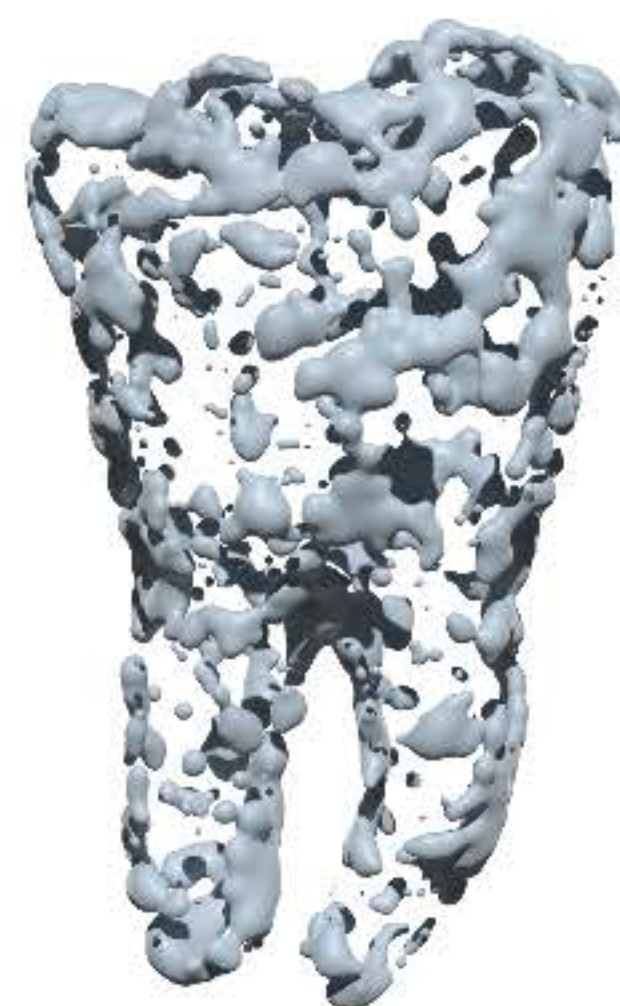
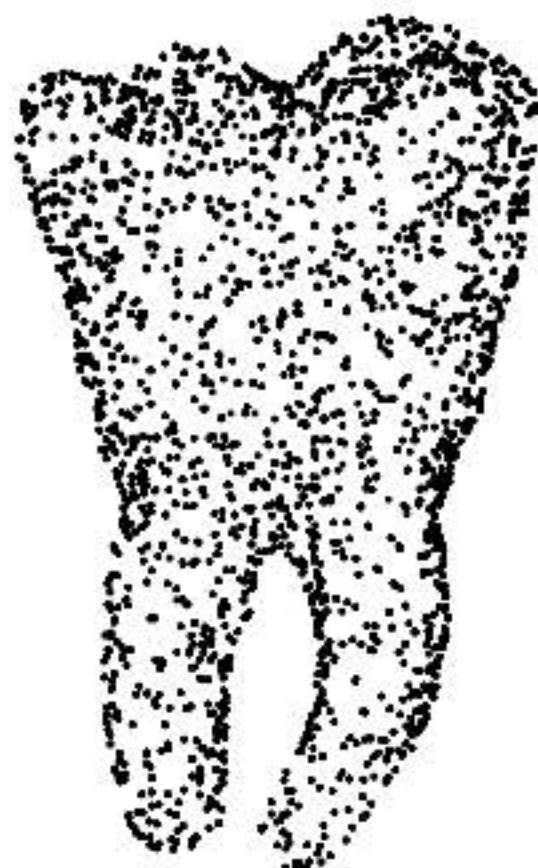
naive “LogSumExp” distances

[Madan & Levin 2022]

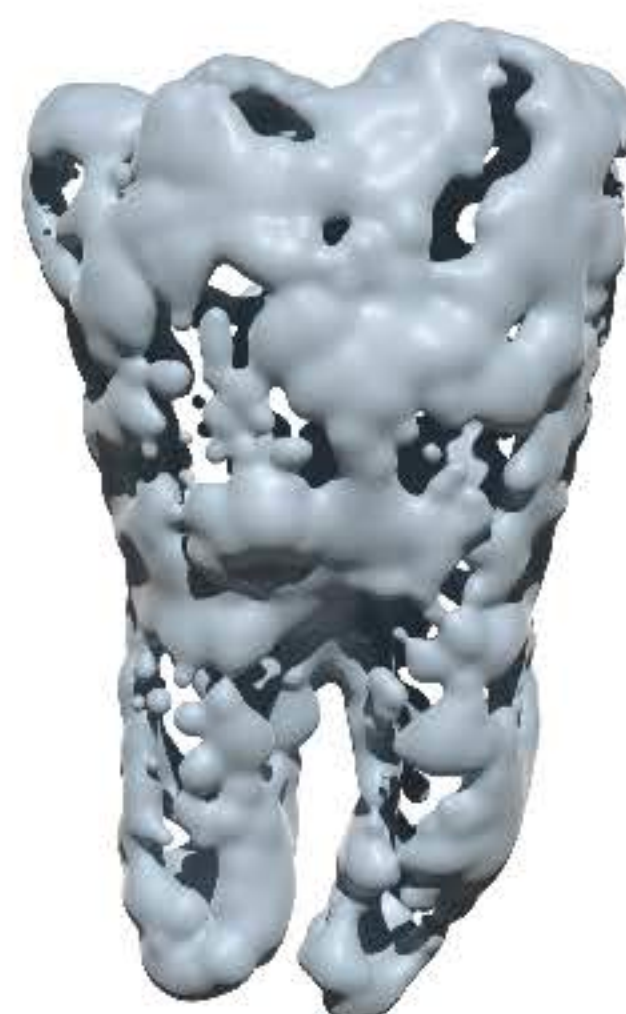
ground
truth



input



$\lambda = 50$



$\lambda = 40$



$\lambda = 35$

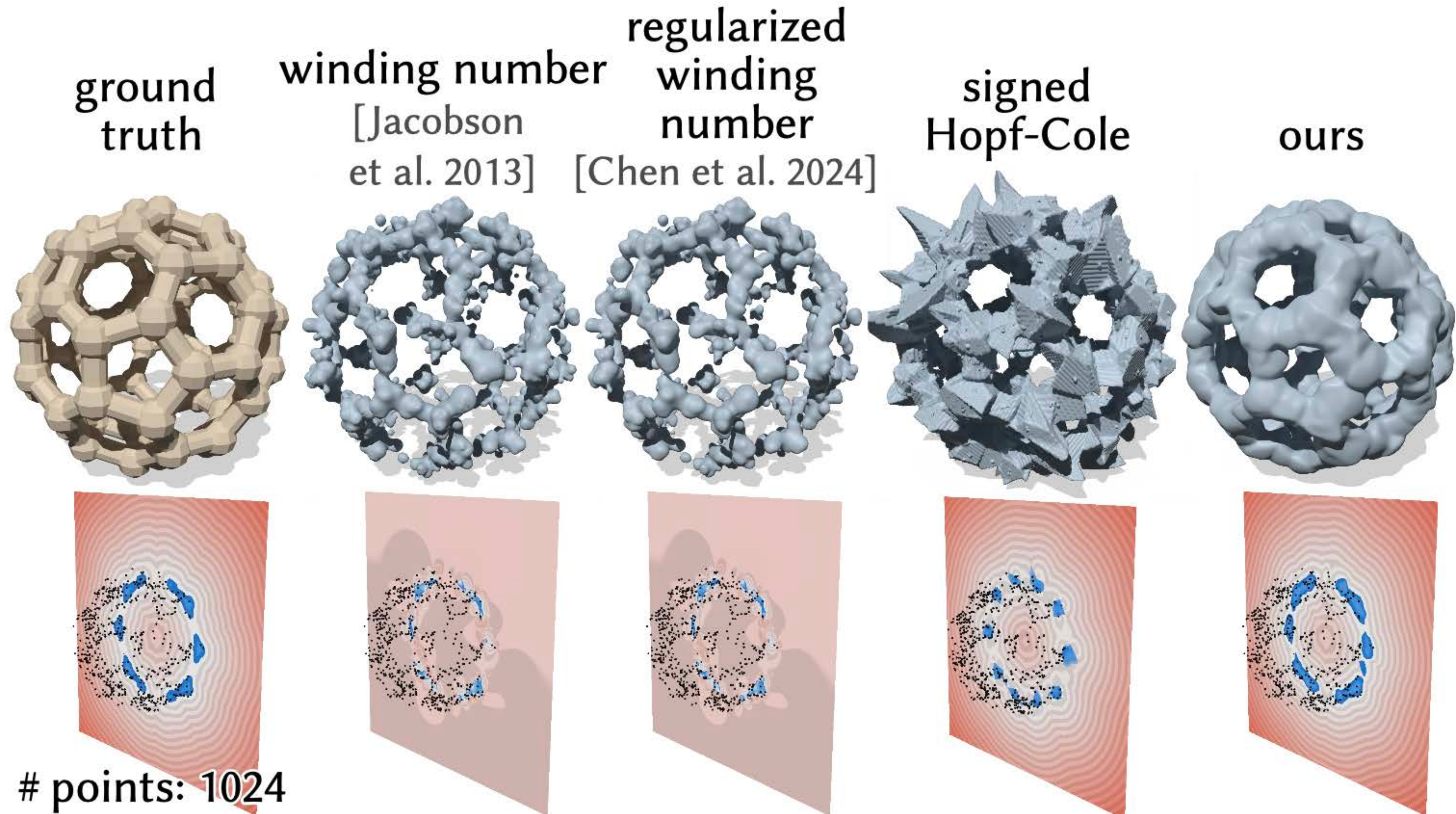


$\lambda = 30$

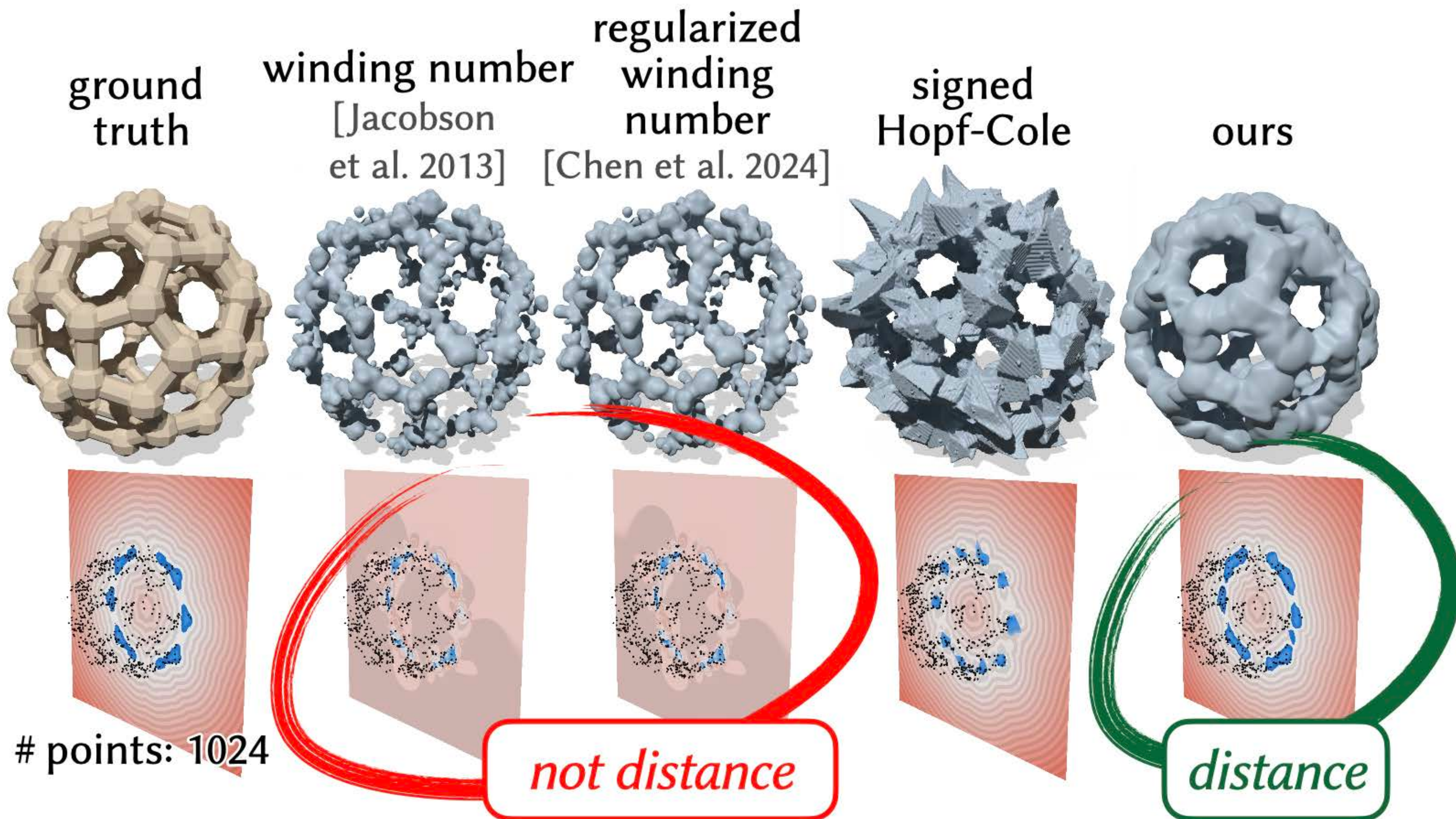
ours,
no tuning



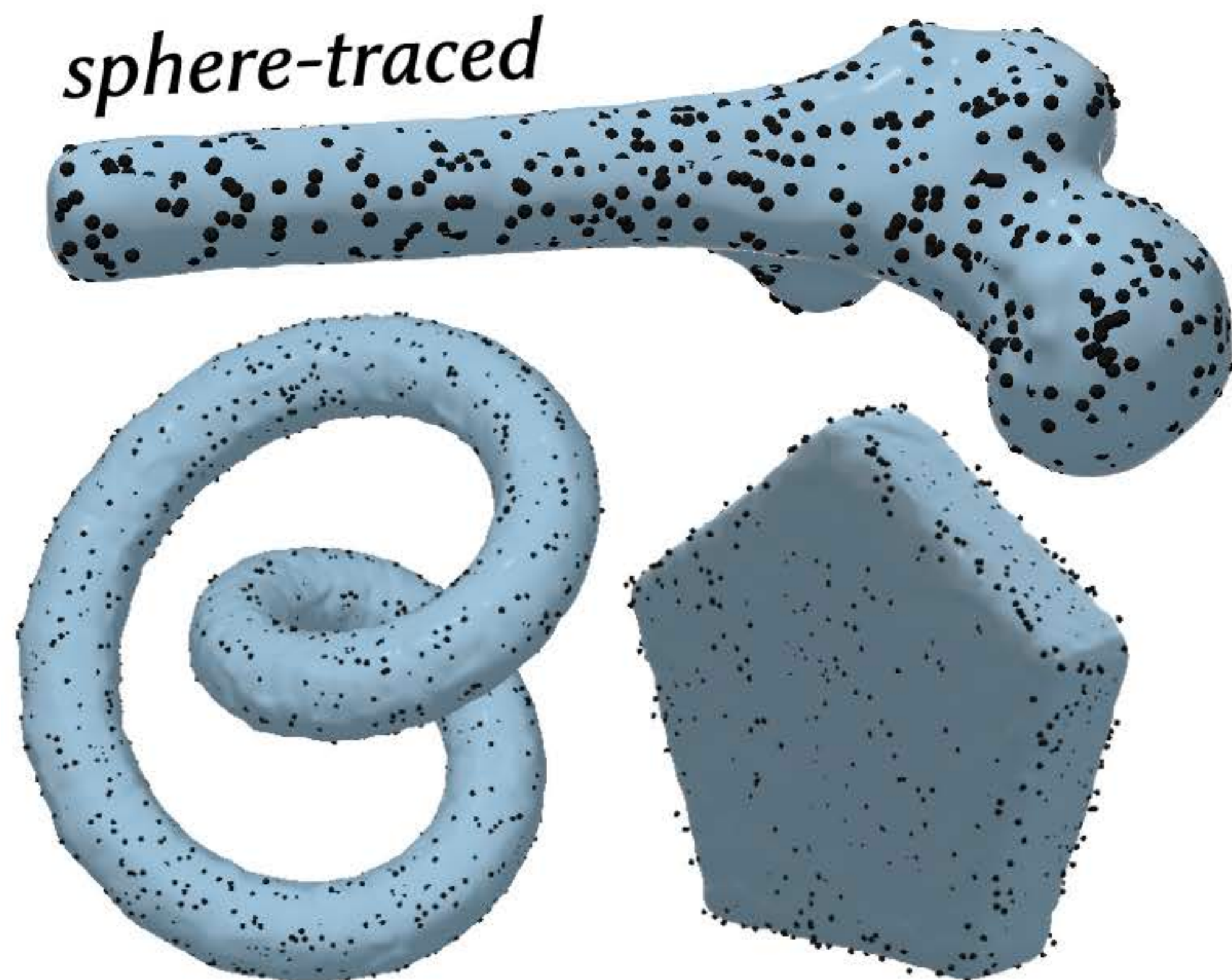
Better than winding number, and almost as fast



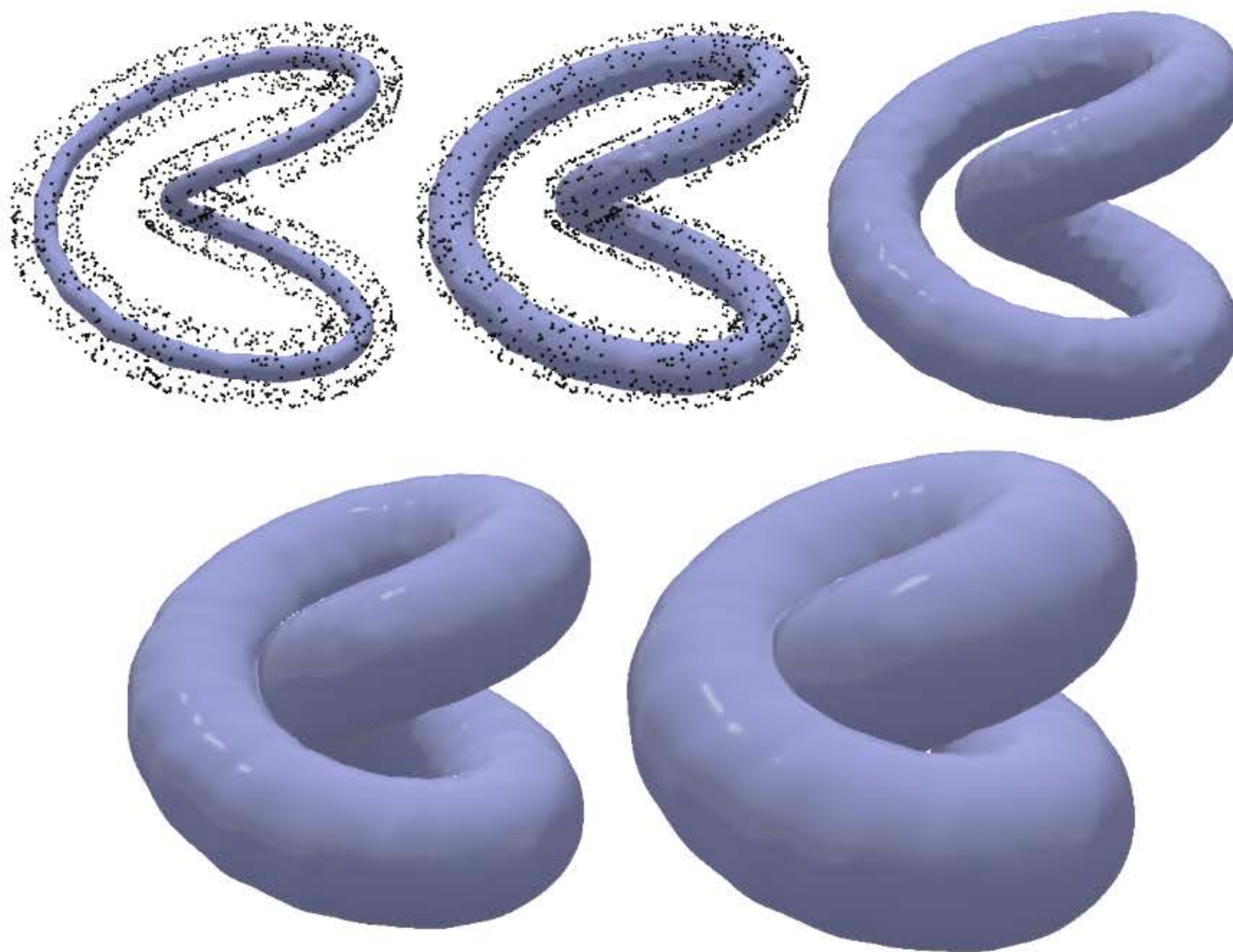
Better than winding number, and almost as fast



Pointwise evaluation

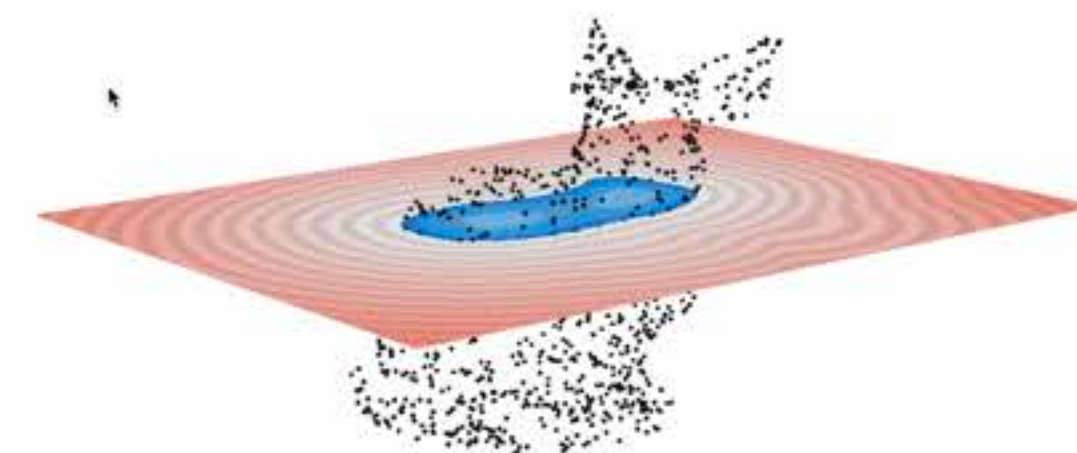
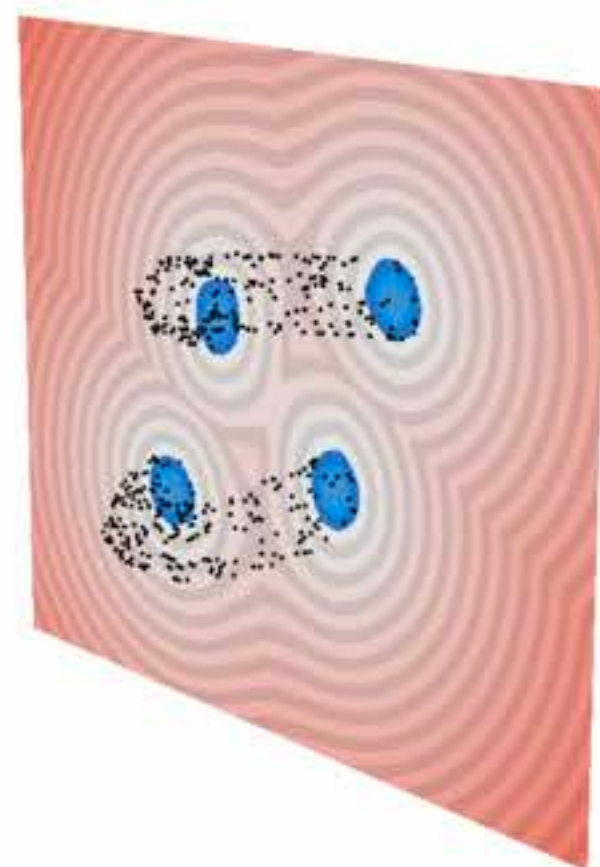
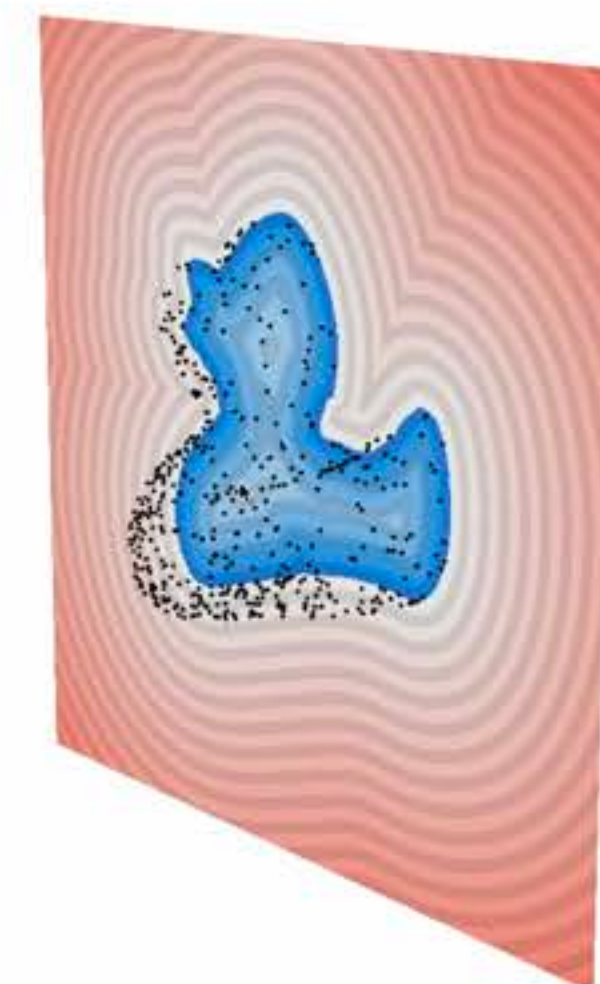
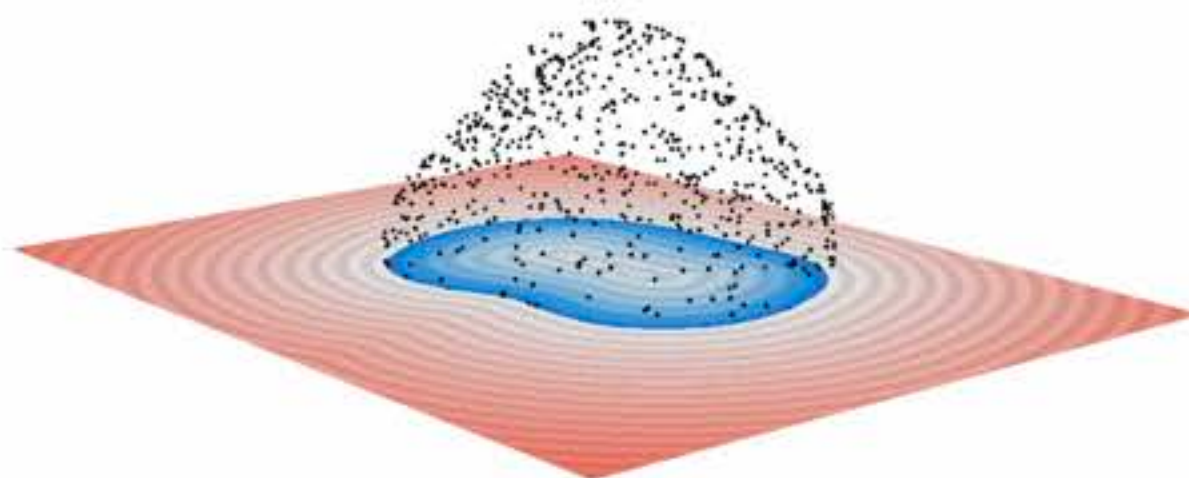
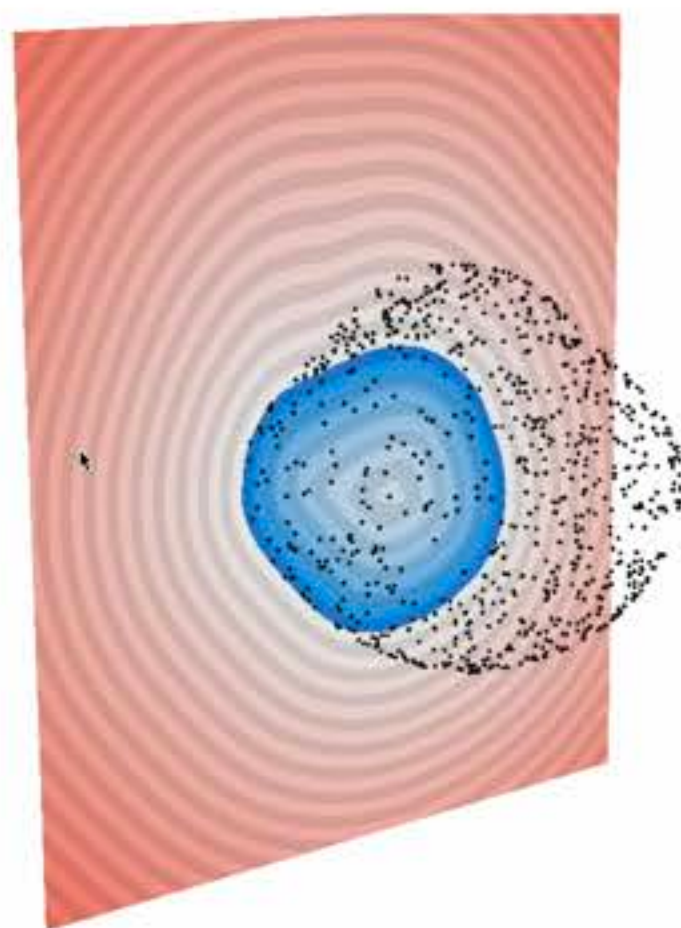


zero level set

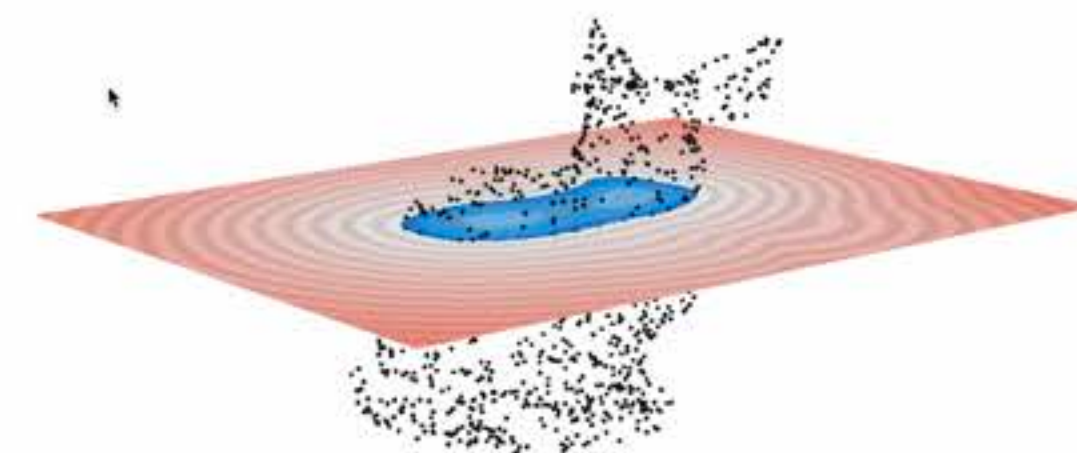
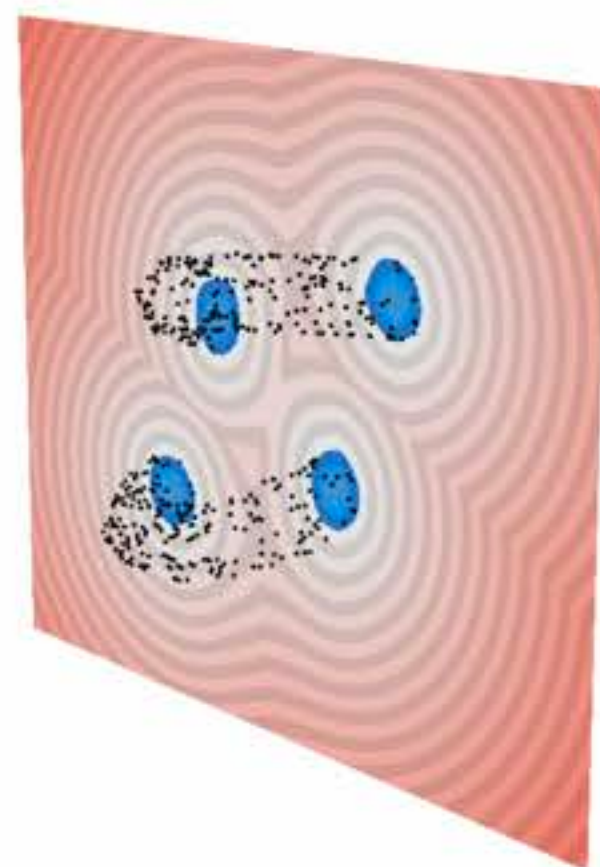
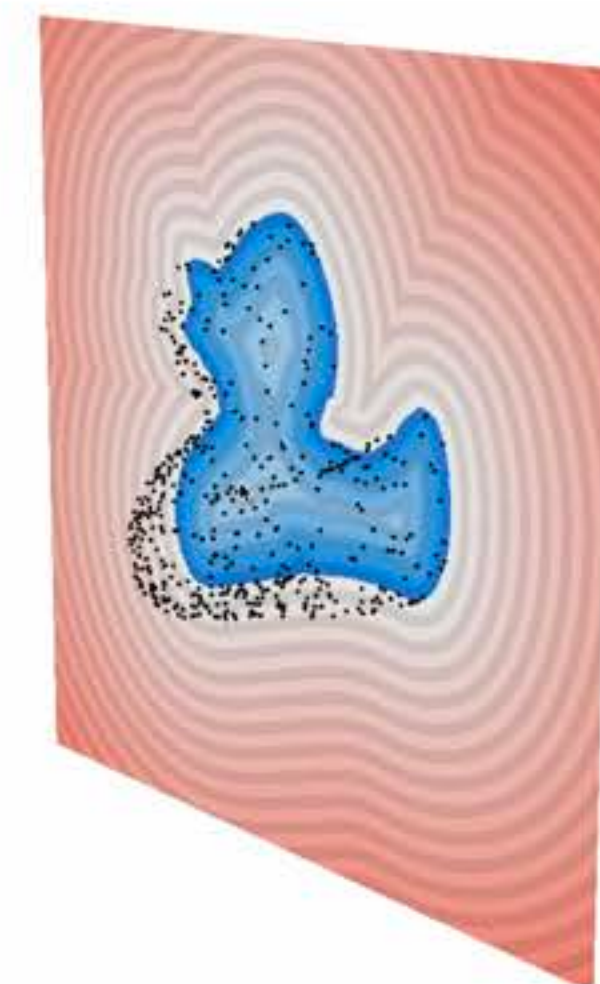
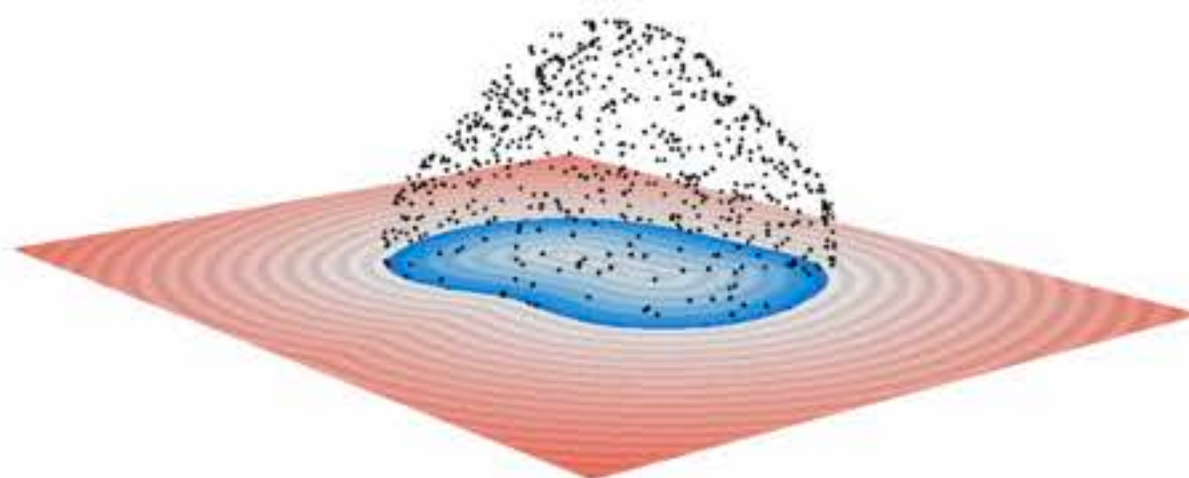
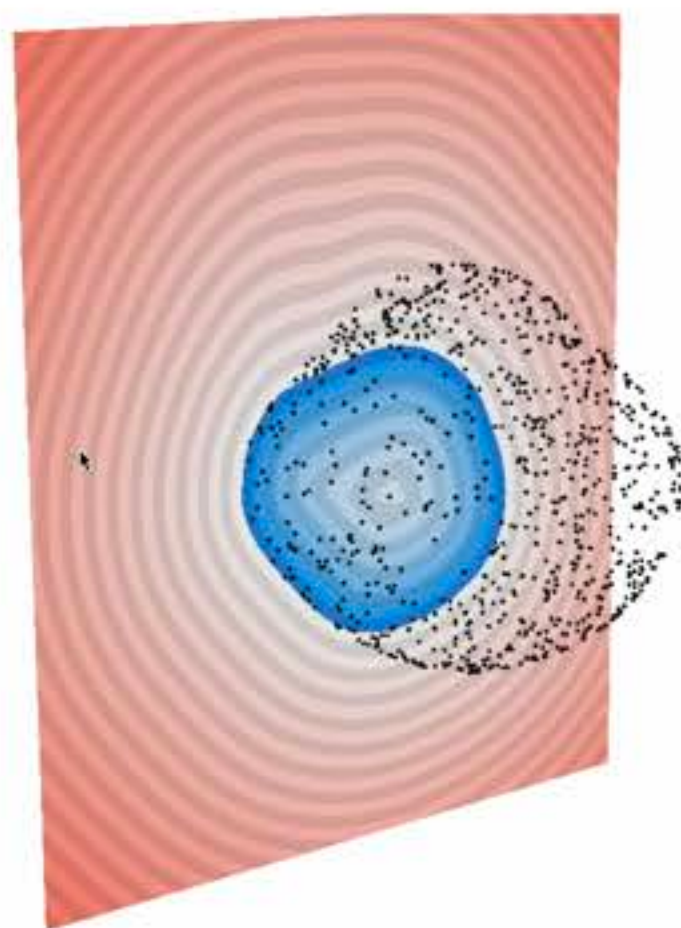


sphere-traced offset surfaces

Pointwise evaluation



Pointwise evaluation

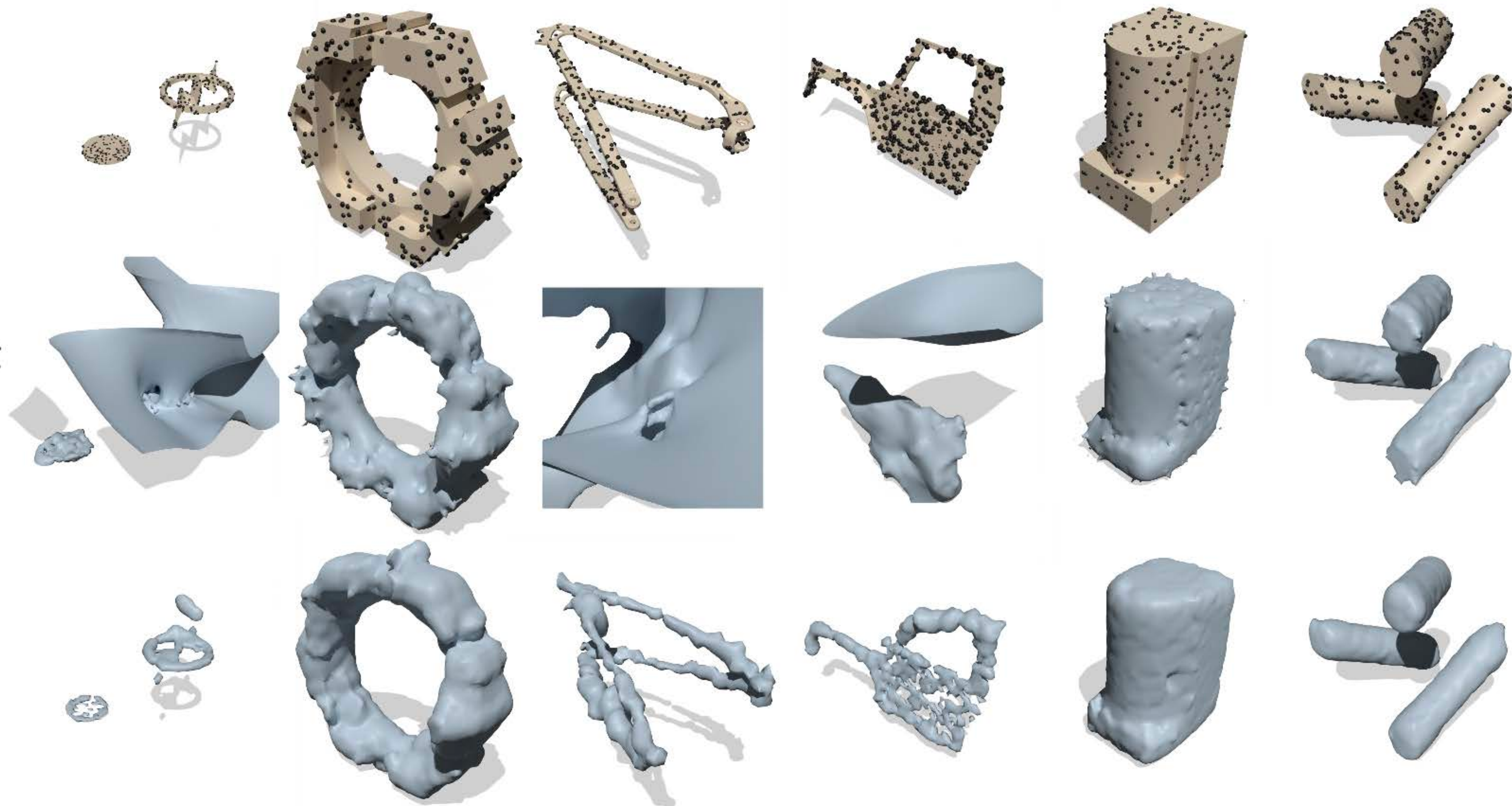


Reconstruction quality on sparse point clouds

true
(512 points)

**signed heat
method**
Feng & Crane
2024

ours



Reconstruction quality on sparse point clouds

true
(512 points)

**signed heat
method**
Feng & Crane
2024

ours



Reconstruction quality

true
(512 points)

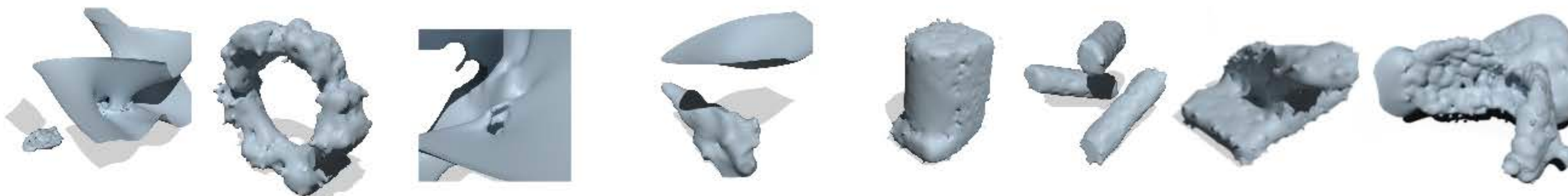


ours



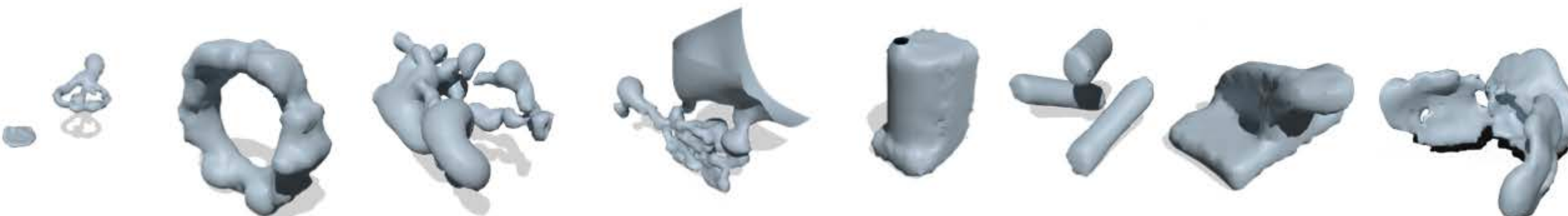
**signed heat
method**

Feng & Crane 2024



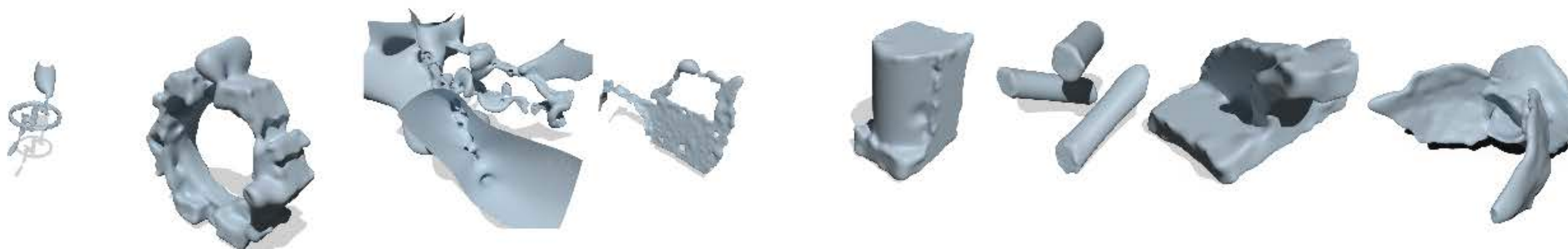
SPSR

Kazhdan & Hoppe 2013



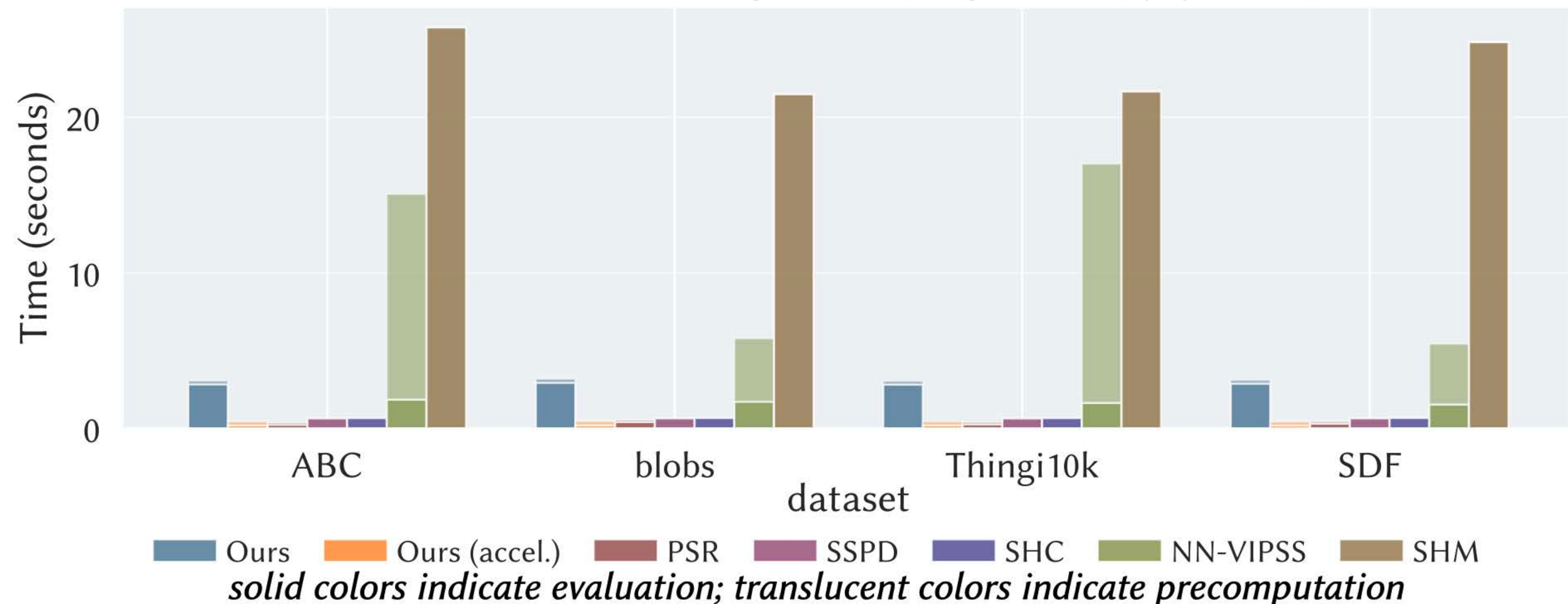
**NN-VIPSS
w/ normals**

Xia & Ju 2025



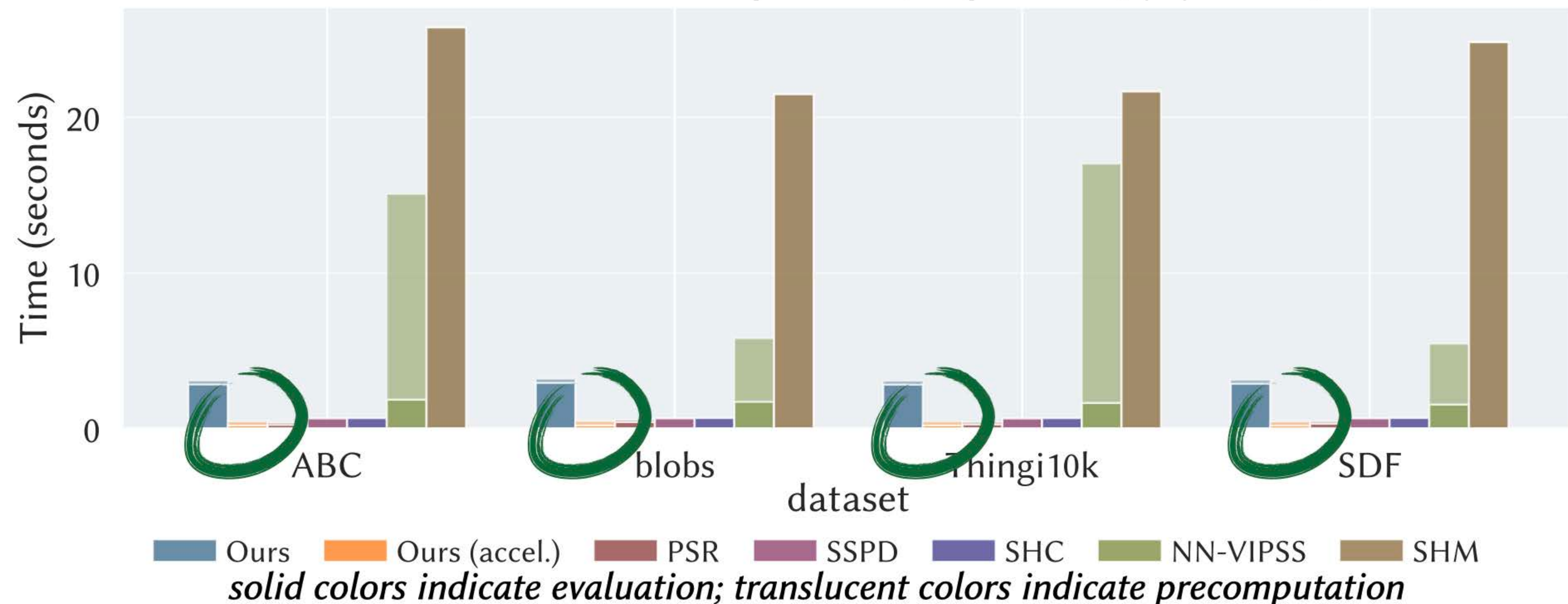
Performance

Total time for 64^3 queries (sequential, $|P| = 512$)



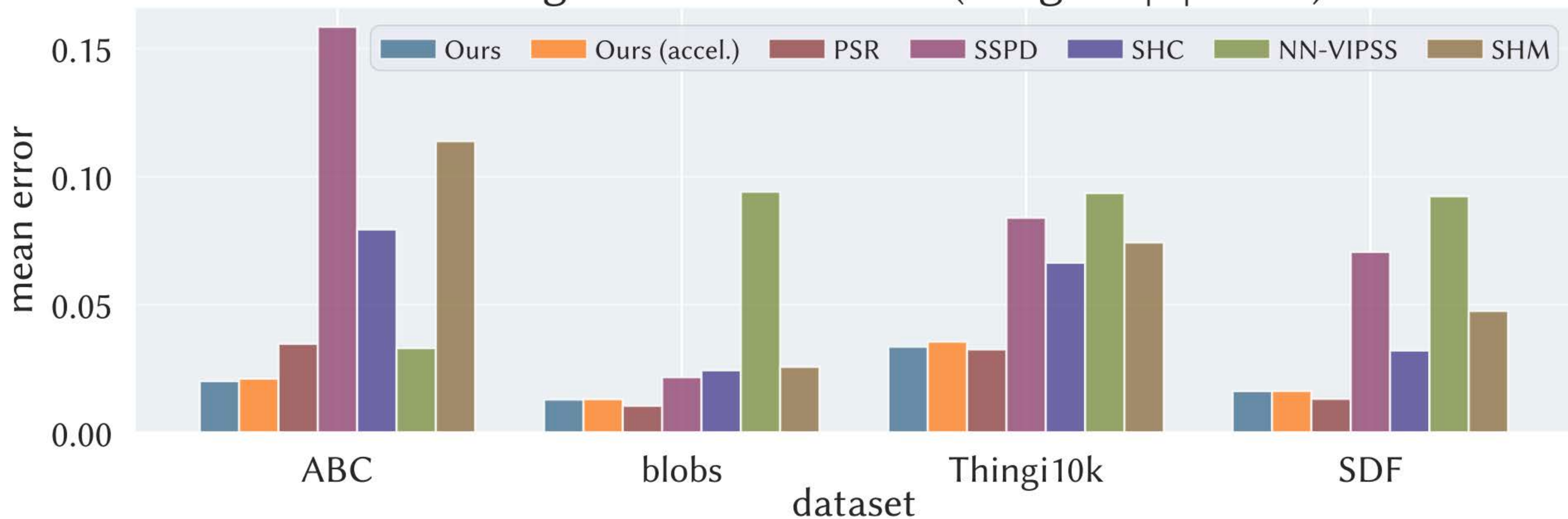
Performance

Total time for 64^3 queries (sequential, $|P| = 512$)



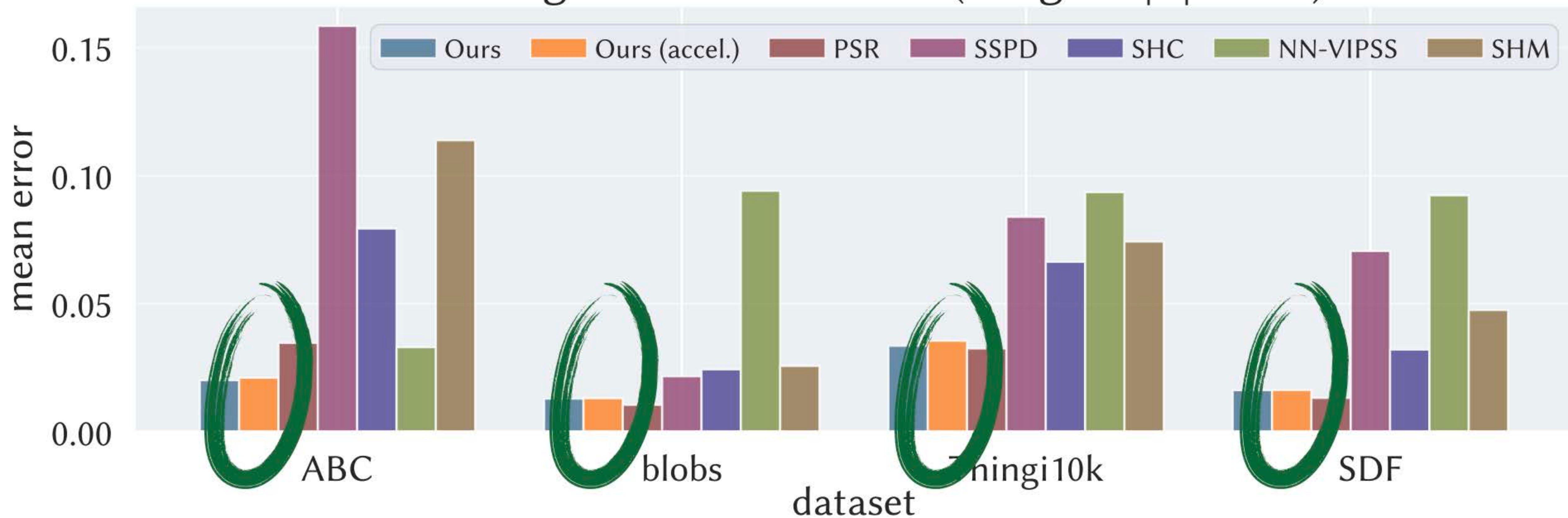
Accuracy

Mean signed distance error (64^3 grid $|P| = 512$)



Accuracy

Mean signed distance error (64^3 grid $|P| = 512$)



We intentionally do **not learn end-to-end**

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1, & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1, & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

**many local
minima**

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1, & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

**many local
minima**

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian

[Wang et al. 2023]

1-Lipschitz Neural Distance Fields

[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition

[Wang et al. 2025]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

etc. ...

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1, & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

**many local
minima**

**re-train for
every shape**

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian

[Wang et al. 2023]

1-Lipschitz Neural Distance Fields

[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition

[Wang et al. 2025]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

etc. ...

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1 & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

many local minima

**re-train for
every shape**

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

simple, analytical formula

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian

[Wang et al. 2023]

1-Lipschitz Neural Distance Fields

[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition

[Wang et al. 2025]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

etc. ...

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1 & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

many local minima

**re-train for
every shape**

local
surface

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

simple, analytical formula

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian

[Wang et al. 2023]

1-Lipschitz Neural Distance Fields

[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition

[Wang et al. 2025]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

etc. ...

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1 & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

many local minima

**re-train for
every shape**

re-usable,
pre-trained network

local
surface

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

simple, analytical formula

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural-Singular-Hessian: Implicit Neural Representation of Unoriented Point Clouds by Enforcing Singular Hessian

[Wang et al. 2023]

1-Lipschitz Neural Distance Fields

[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition

[Wang et al. 2025]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

etc. ...

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1 & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

many local minima

**re-train for
every shape**

re-usable,
pre-trained network

local
surface

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

simple, analytical formula

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural Representation of Signed Distance Fields by Enforcing Singular Hessian

[

1-Lipschitz Neural Distance Fields

[Coiffier & Béthune 2024]

HotSpot: Signed Distance Function Optimization with an Asymptotically Sufficient Condition

[Wang et al. 2025]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

etc. ...

less training

We intentionally do **not** learn end-to-end

eikonal equation

$$\begin{aligned}\|\nabla u(x)\|^2 &= 1 & x \notin \Omega \\ u(x) &= 0 & x \in \Omega \\ \frac{\partial u}{\partial n}(x) &= 1 & x \in \Omega\end{aligned}$$

nonlinear

many local minima

**re-train for
every shape**

re-usable,
pre-trained network

local
surface

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

simple, analytical formula

Implicit Geometric Regularization for Learning Shapes

[Gropp et al. 2020]

Neural Unsigned Distance Fields for Implicit Function Learning

[Chibane et al. 2020]

Neural-Pull: Learning Signed Distance Functions from Point Clouds by Learning to Pull Space onto Surfaces

[Ma et al. 2021]

Phase Transitions, Distance Functions, and Implicit Neural Representations

[Lipman 2019]

Critical Regularizations for Neural Surface Reconstruction in the Wild

[Zhang et al. 2022]

p-Poisson Surface Reconstruction in Curl-Free Flow From Point Clouds

[Park et al. 2023]

StEik: Stabilizing the Optimization of Neural Signed Distance Functions and Finer Shape Representation

[Yang et al. 2023]

Neural Representation of Signed Distance Fields by Enforcing Singular Hessian

[Zhang et al. 2024]

1-Lipschitz Neural Distance Fields

[Zhang et al. 2024]

Optimization with an

[Zhang et al. 2024]

[Zhang et al. 2024]

efunc: An Efficient Function Representation without Neural Networks

[Zhang & Wonka 2025]

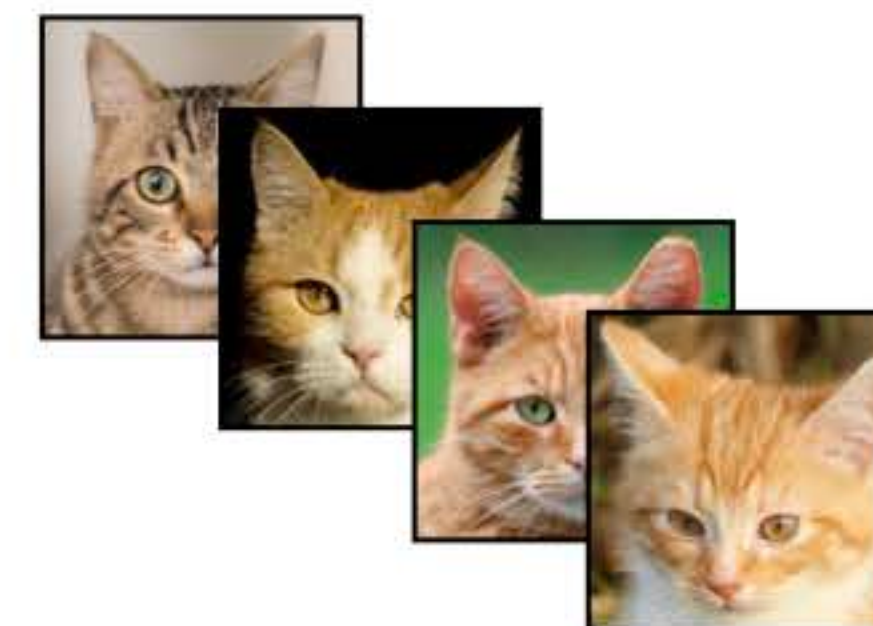
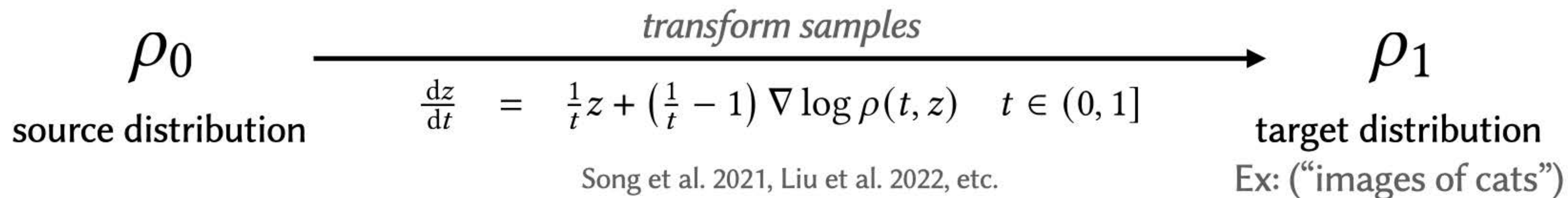
etc. ...

Higher-dimensional manifold learning

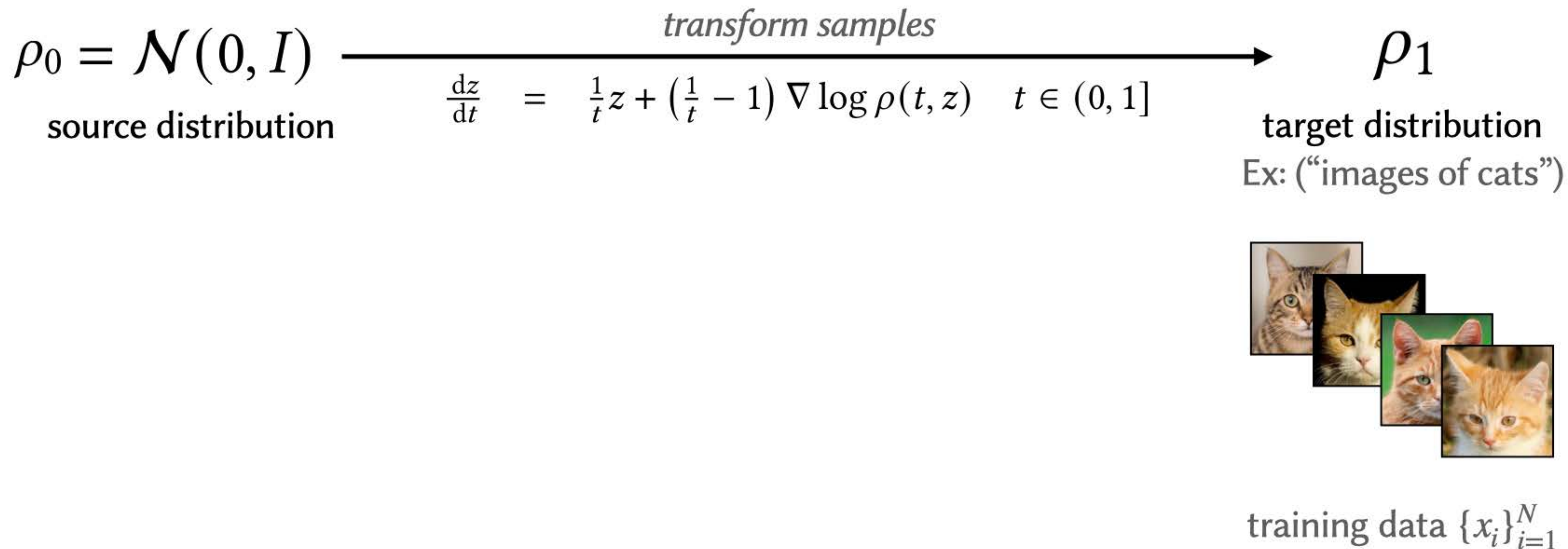
Higher-dimensional manifold learning

Goal: Sample from a target distribution.

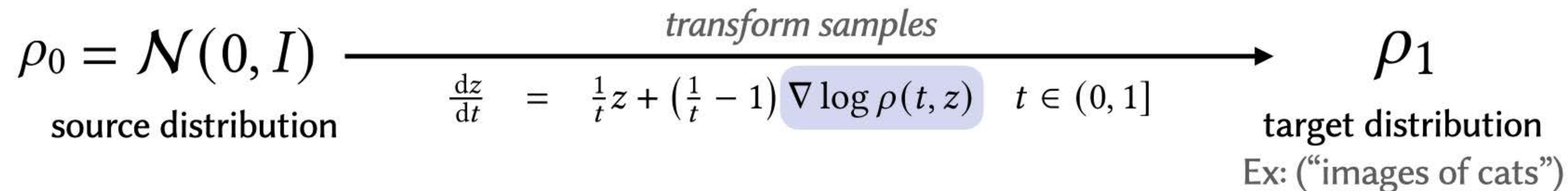
One approach: Sample from an easy-to-sample *source distribution*, then apply a transformation.



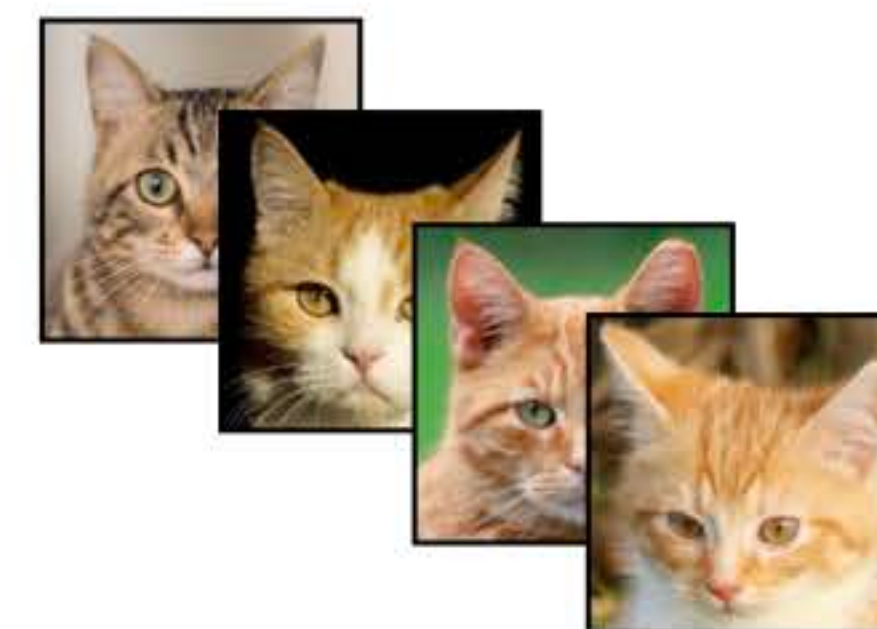
Connection to generative models



Connection to generative models



$$\nabla \log \rho(t, z) \propto \frac{\sum_{i=1}^N t x_i \exp \left(-\frac{1}{2} \frac{\|z - t x_i\|^2}{(1-t)^2} \right)}{\sum_{j=1}^N \exp \left(-\frac{1}{2} \frac{\|z - t x_j\|^2}{(1-t)^2} \right)} - z$$



training data $\{x_i\}_{i=1}^N$

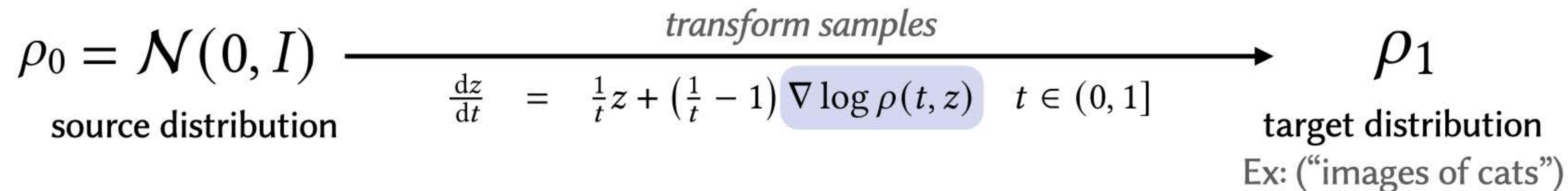
Flow Straight and Fast

[Liu et al. 2021]

Closed-form Diffusion Models

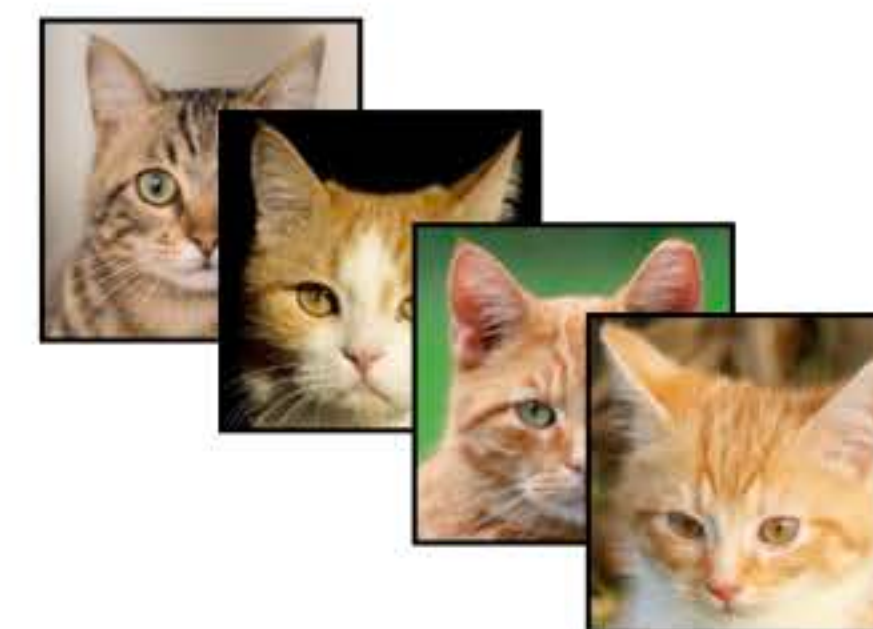
[Scarvelis et al. 2023]

Connection to generative models



$$\nabla \log \rho(t, z) \propto \frac{\sum_{i=1}^N t x_i \exp \left(-\frac{1}{2} \frac{\|z - t x_i\|^2}{(1-t)^2} \right)}{\sum_{j=1}^N \exp \left(-\frac{1}{2} \frac{\|z - t x_j\|^2}{(1-t)^2} \right)} - z$$

current sample



training data $\{x_i\}_{i=1}^N$

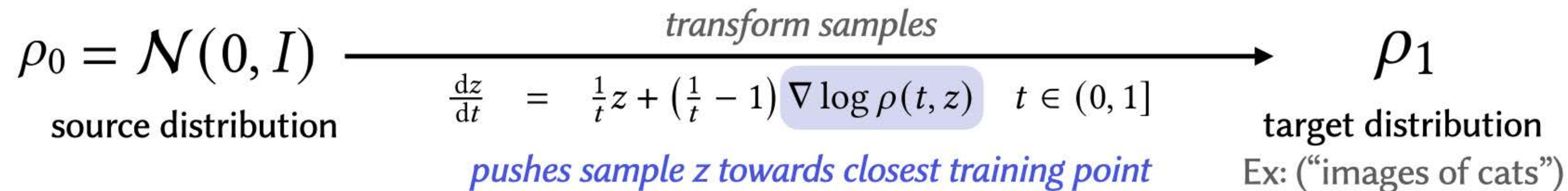
Flow Straight and Fast

[Liu et al. 2021]

Closed-form Diffusion Models

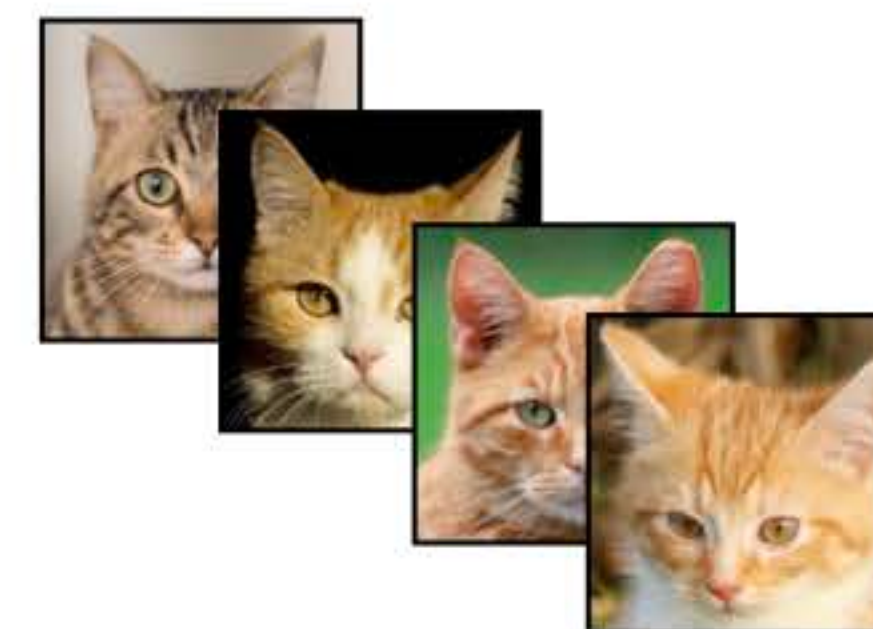
[Scarvelis et al. 2023]

Connection to generative models



$$\nabla \log \rho(t, z) \propto \frac{\sum_{i=1}^N tx_i \exp\left(-\frac{1}{2} \frac{\|z - tx_i\|^2}{(1-t)^2}\right)}{\sum_{j=1}^N \exp\left(-\frac{1}{2} \frac{\|z - tx_j\|^2}{(1-t)^2}\right)} - z$$

current sample *approximation of $\operatorname{argmin}_{tx_i} \|z - tx_i\|^2$*

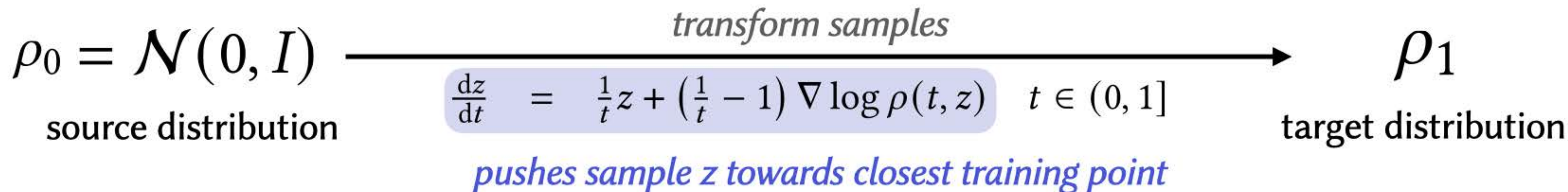


training data $\{x_i\}_{i=1}^N$

Flow Straight and Fast
[Liu et al. 2021]

Closed-form Diffusion Models
[Scarvelis et al. 2023]

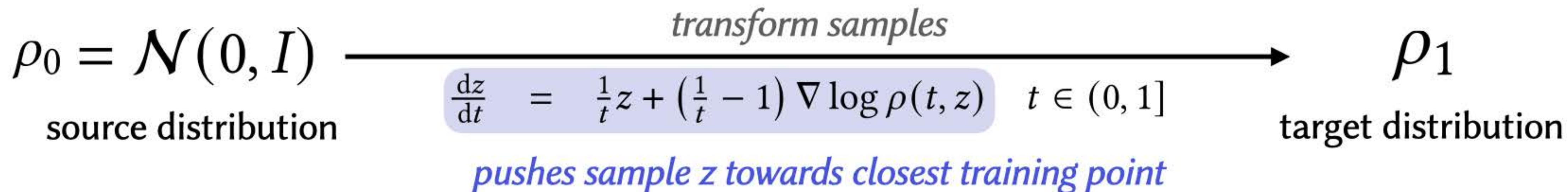
Connection to generative models



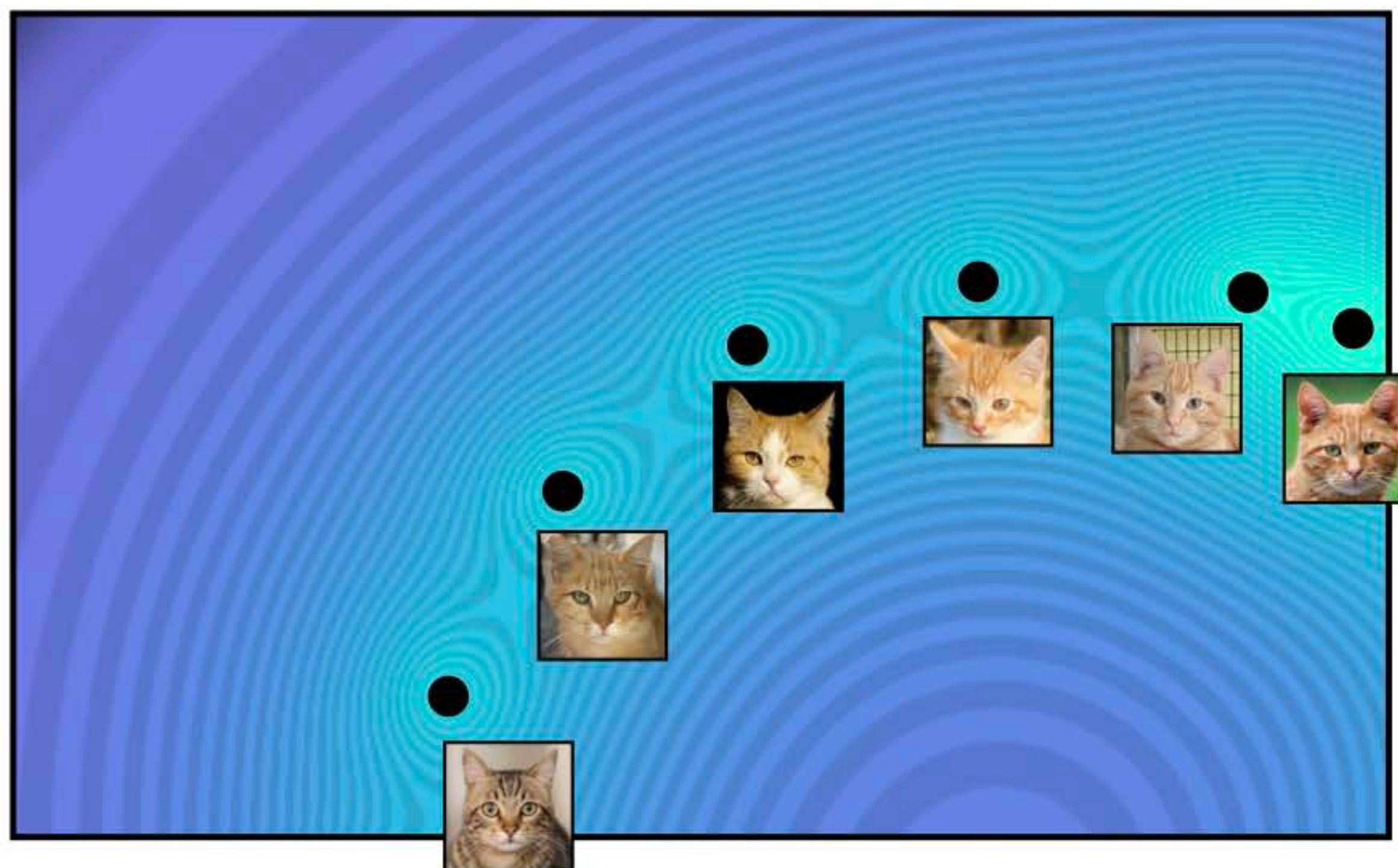
Flow-based models can only “memorize” their training data!

[Liu et al. 2021, Scarvelis et al. 2023, Gu et al. 2023, Gao & Li 2024, Carlini et al. 2023, Jain et al. 2024, Biroli et al. 2024, Gu et al. 2025, Somepalli et al. 2022; Pidstrigach 2022, Yoon et al. 2023, Kamb & Ganguli 2025]

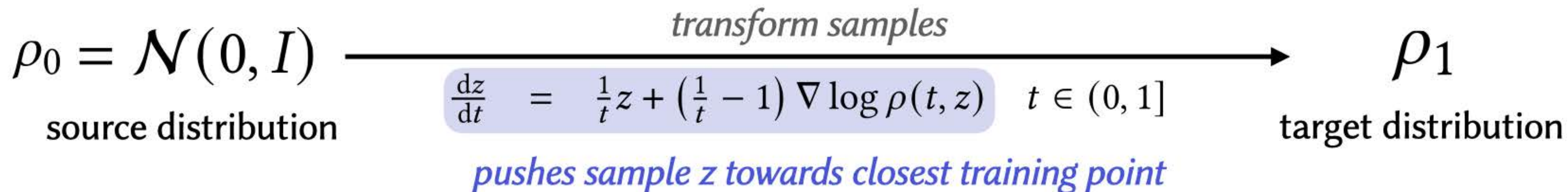
Connection to generative models



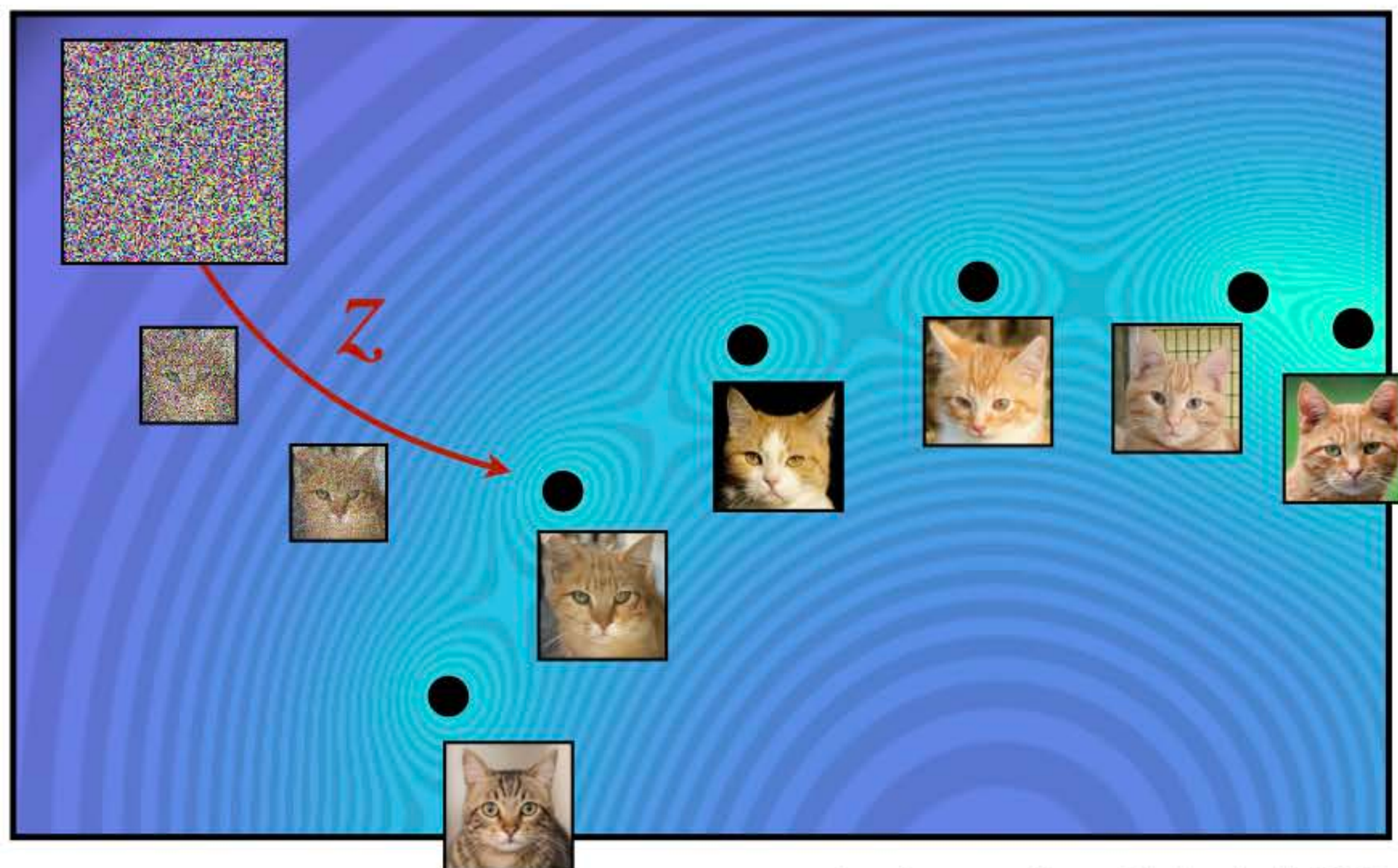
Flow-based models can only “memorize” their training data!



Connection to generative models



Flow-based models can only “memorize” their training data!



Connection to generative models

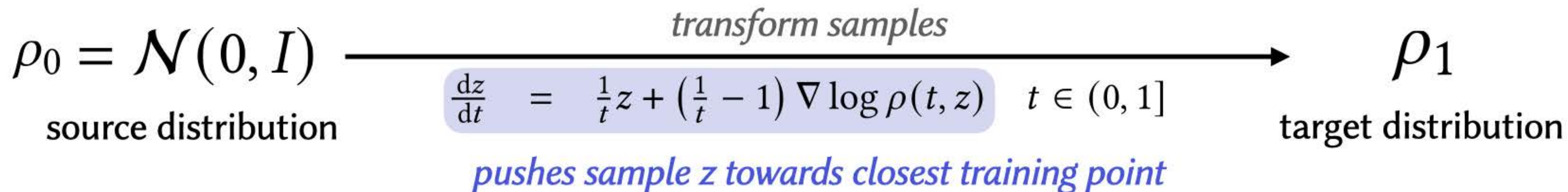
$$\rho_0 = \mathcal{N}(0, I) \xrightarrow[\substack{\frac{dz}{dt} = \frac{1}{t}z + \left(\frac{1}{t} - 1\right) \nabla \log \rho(t, z) \quad t \in (0, 1] \\ \text{pushes sample } z \text{ towards closest training point}}]{\text{transform samples}} \rho_1$$

source distribution target distribution

Flow-based models can only “memorize” their training data!



Connection to generative models



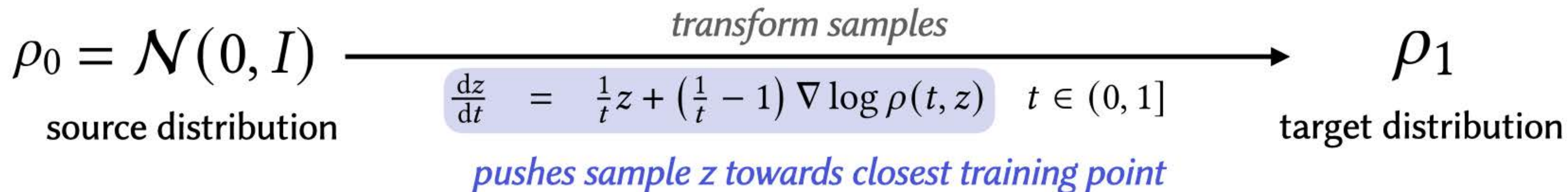
Flow-based models can only “memorize” their training data!



Many authors have observed the “memorization” phenomenon...

[Liu et al. 2021, Scarvelis et al. 2023, Gu et al. 2023, Gao & Li 2024, Carlini et al. 2023, Jain et al. 2024, Biroli et al. 2024, Gu et al. 2025, Somepalli et al. 2022; Pidstrigach 2022, Yoon et al. 2023, Kamb & Ganguli 2025]

Connection to generative models



Flow-based models can only “memorize” their training data!

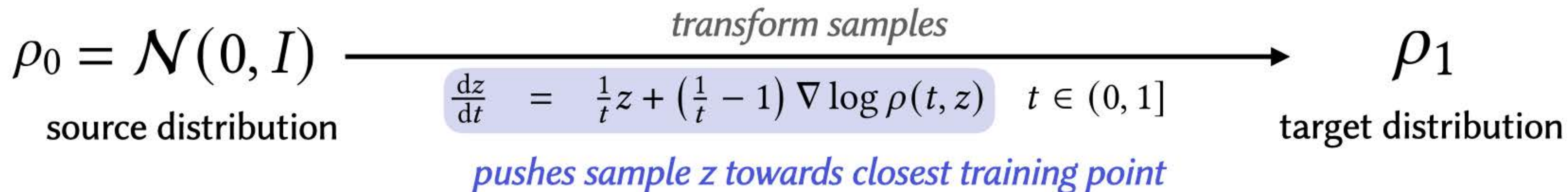


Many authors have observed the “memorization” phenomenon...

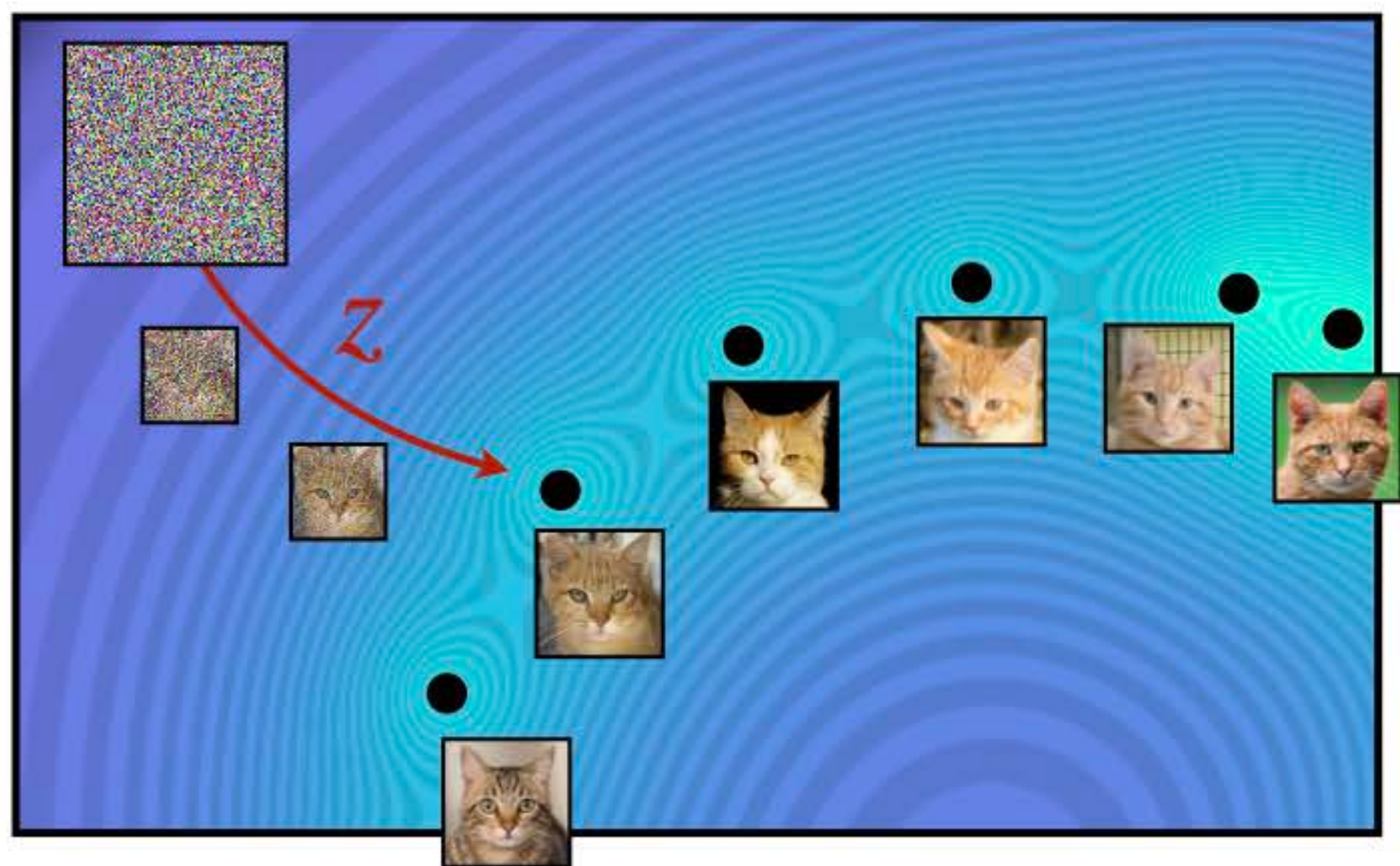
[Liu et al. 2021, Scarvelis et al. 2023, Gu et al. 2023, Gao & Li 2024, Carlini et al. 2023, Jain et al. 2024, Biroli et al. 2024, Gu et al. 2025, Somepalli et al. 2022; Pidstrigach 2022, Yoon et al. 2023, Kamb & Ganguli 2025]

Current state-of-the-art: train a (really expensive!) neural network that implicitly regularizes the problem.

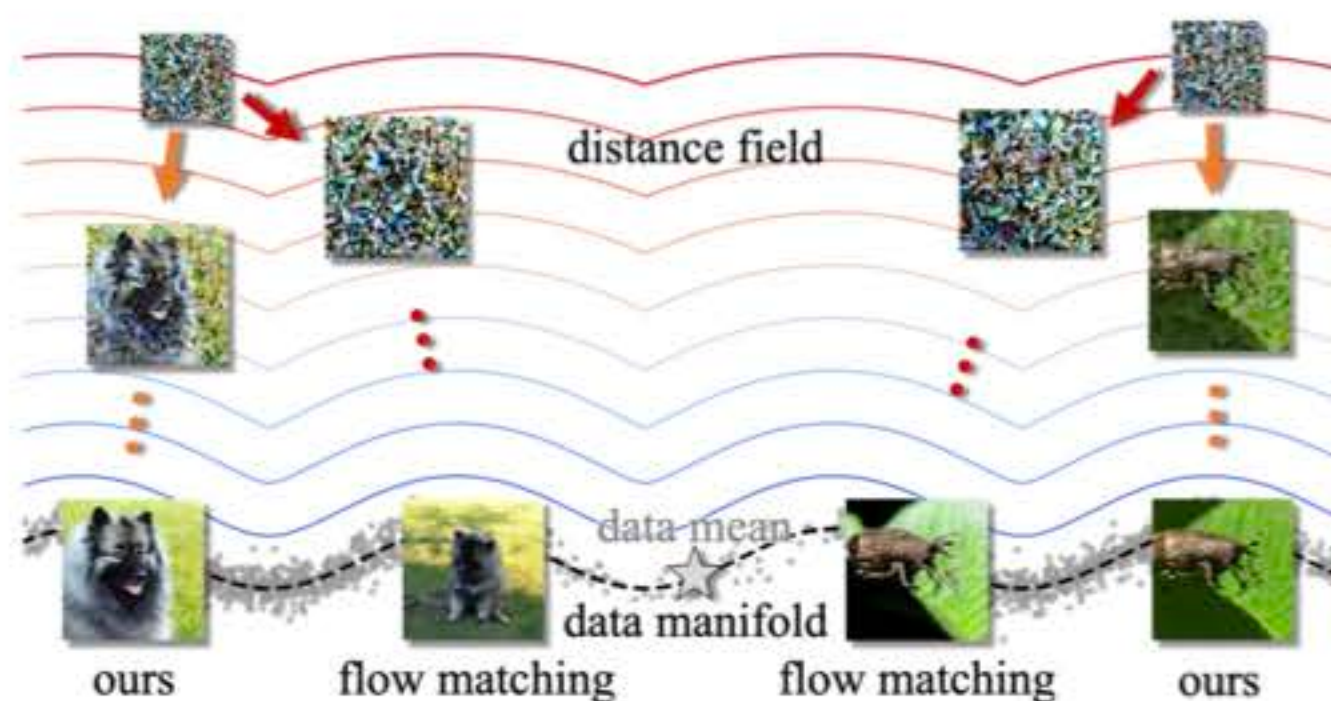
Connection to generative models



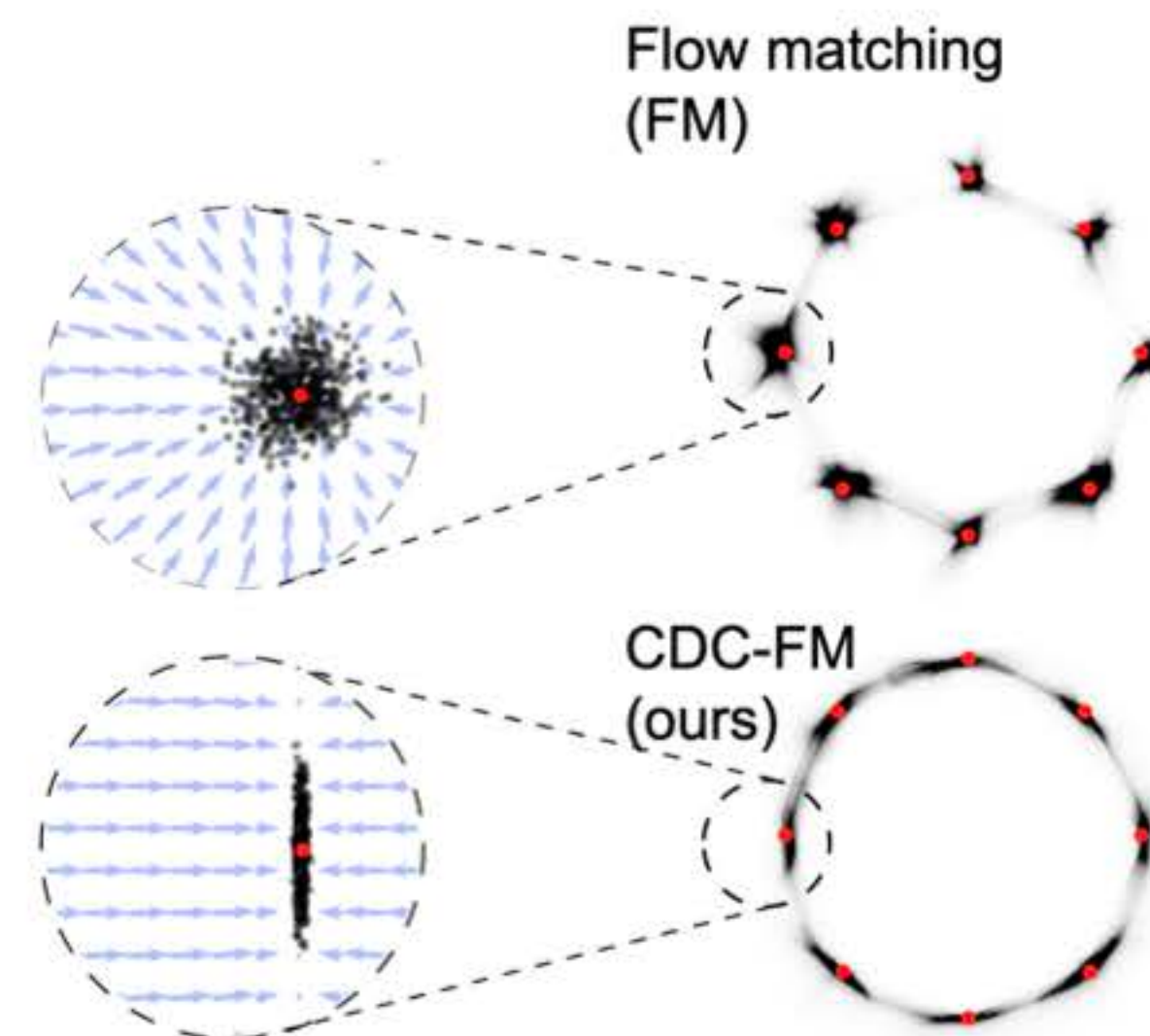
Flow-based models can only “memorize” their training data!



Distance Marching for Generative Modeling
[Wang et al. 2026]



Carré du champ Flow Matching: Better Quality-Generalisation Tradeoff in Generative Models
[Bamberger et al. 2026]

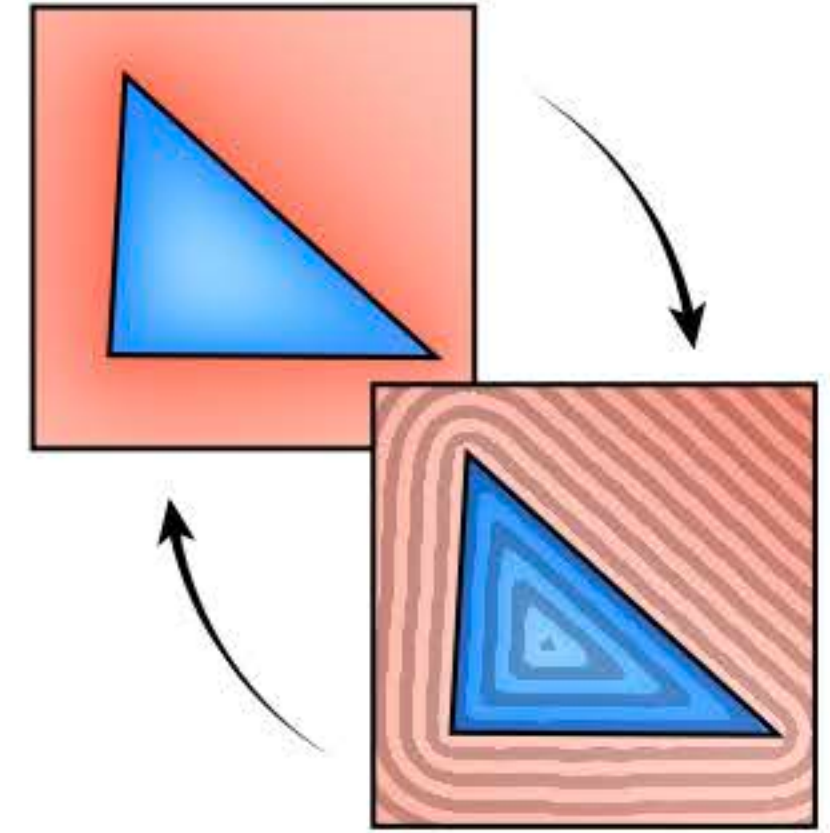


CONCLUSION

Takeaways

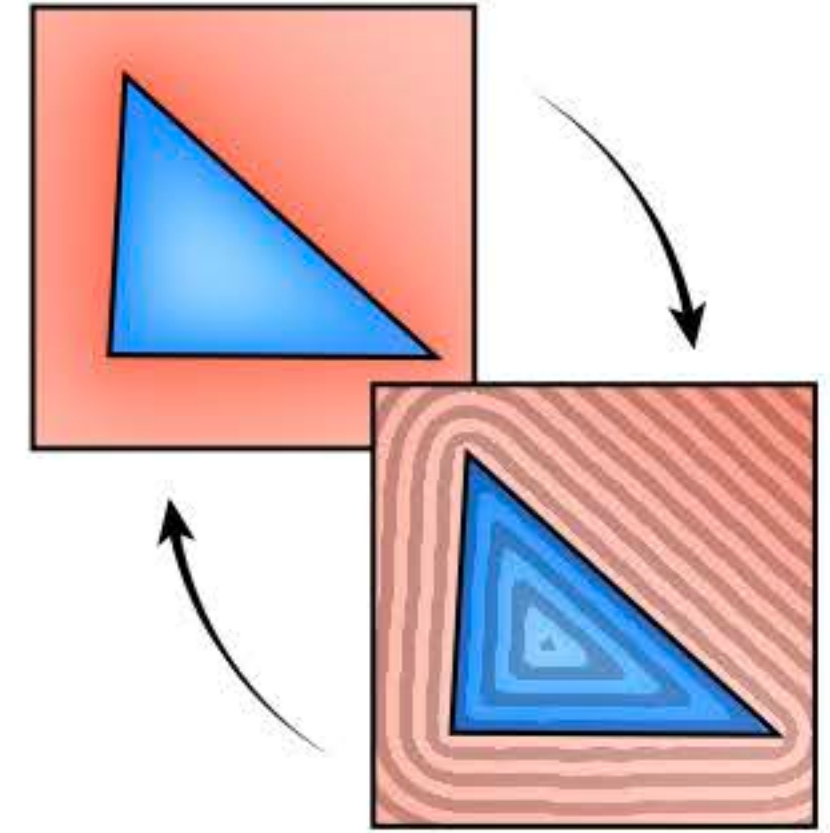
Takeaways

- Signed distance, winding numbers, and Poisson Surface Reconstruction are, in some sense, *equivalent*



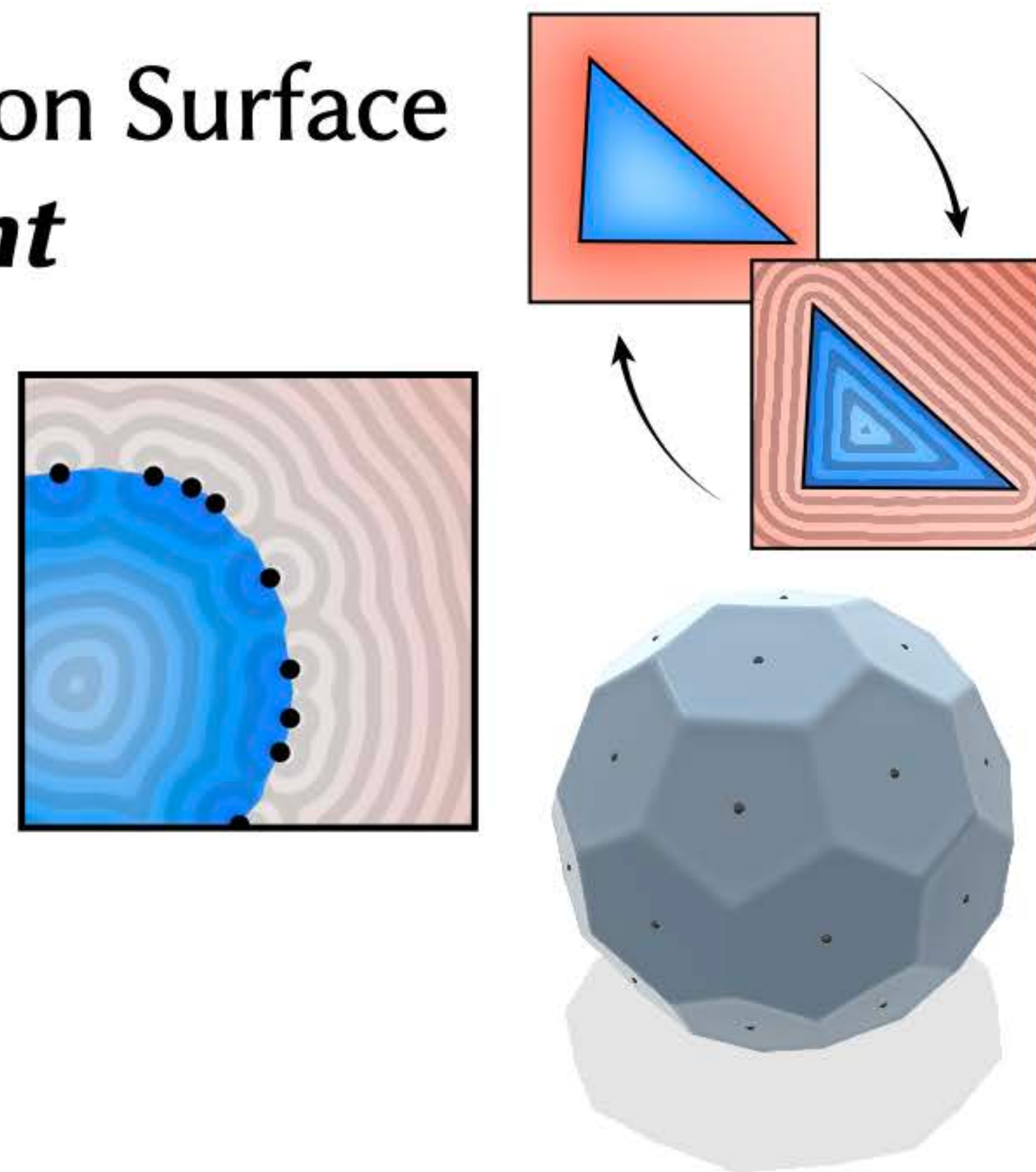
Takeaways

- Signed distance, winding numbers, and Poisson Surface Reconstruction are, in some sense, *equivalent*
- Decades of distance approximations can be distilled into two formulas



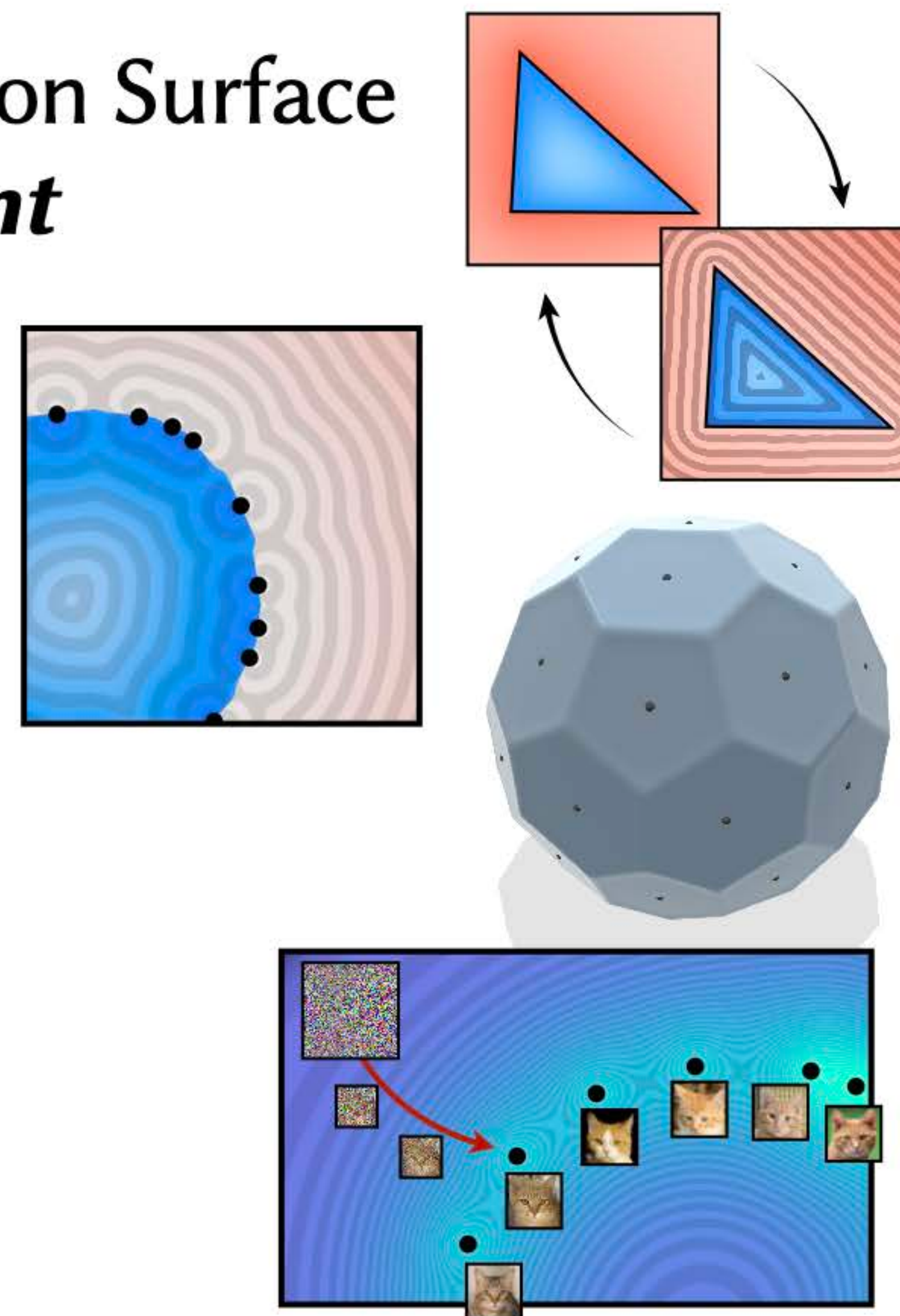
Takeaways

- Signed distance, winding numbers, and Poisson Surface Reconstruction are, in some sense, *equivalent*
- Decades of distance approximations can be distilled into two formulas — but **none of them work for point clouds**



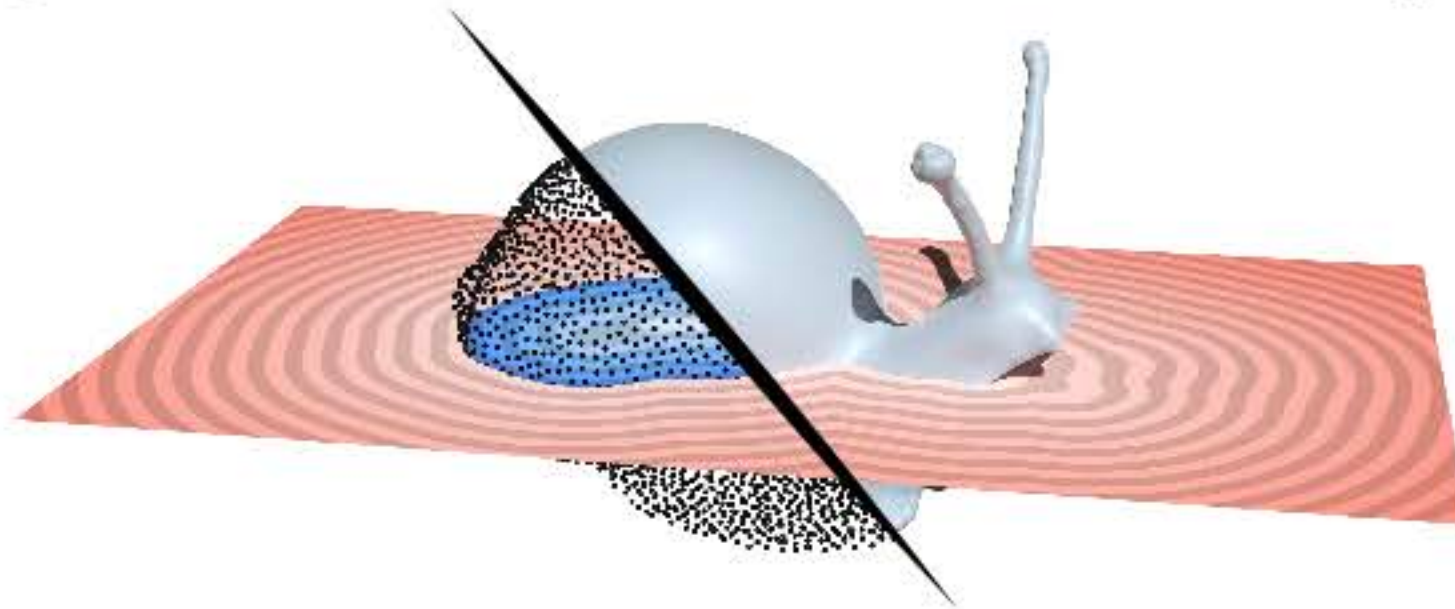
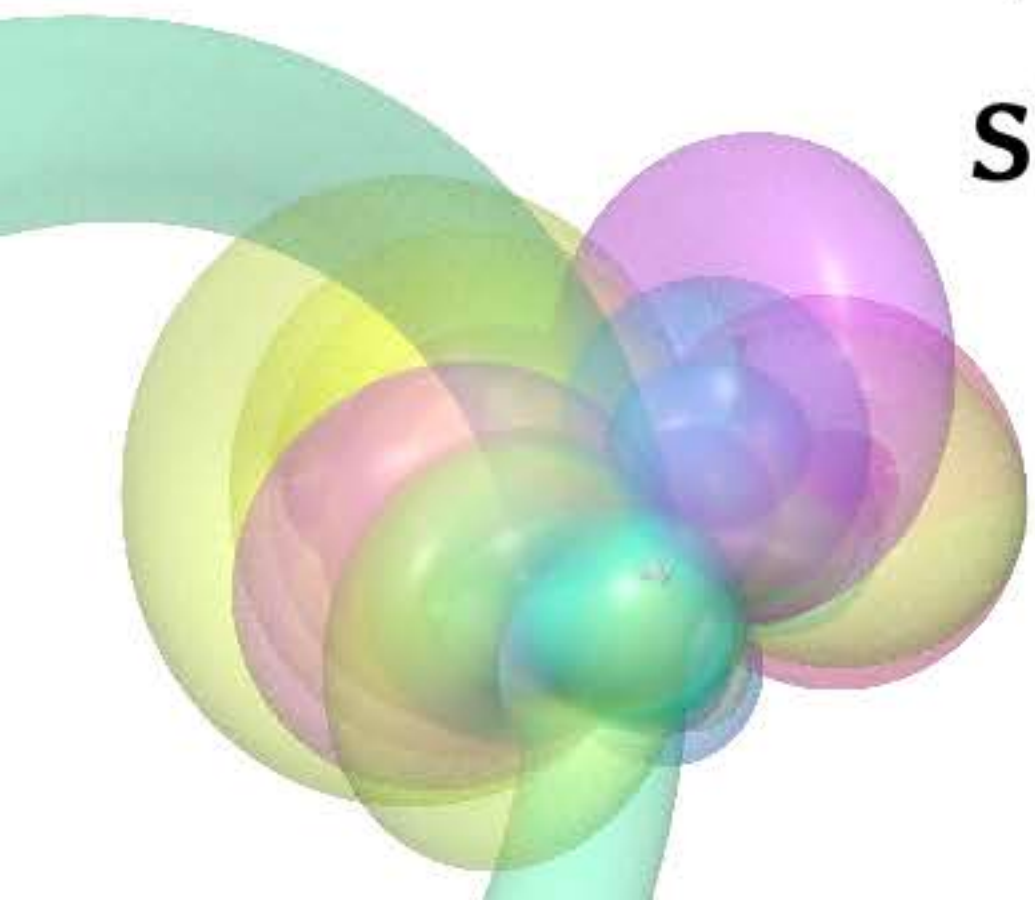
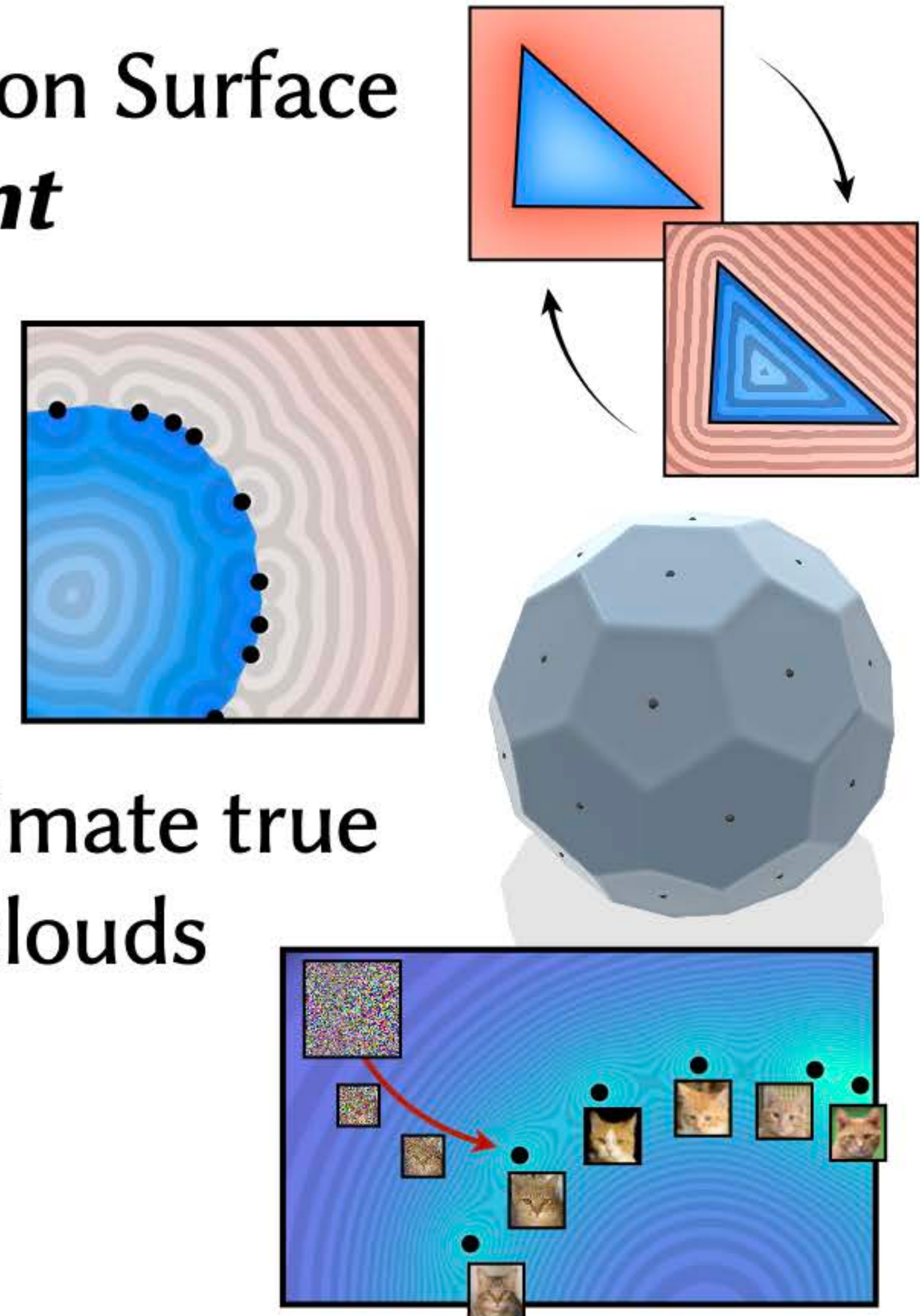
Takeaways

- Signed distance, winding numbers, and Poisson Surface Reconstruction are, in some sense, *equivalent*
- Decades of distance approximations can be distilled into two formulas — but **none of them work for point clouds**



Takeaways

- Signed distance, winding numbers, and Poisson Surface Reconstruction are, in some sense, *equivalent*
- Decades of distance approximations can be distilled into two formulas — but **none of them work for point clouds**
 - We treat *points as tori*, and approximate true signed distance directly from point clouds



Future directions

- Accelerating pre-computation (torus fitting)

Future directions

- Accelerating pre-computation (torus fitting)

Precompute times using one NVIDIA RTX 3090:

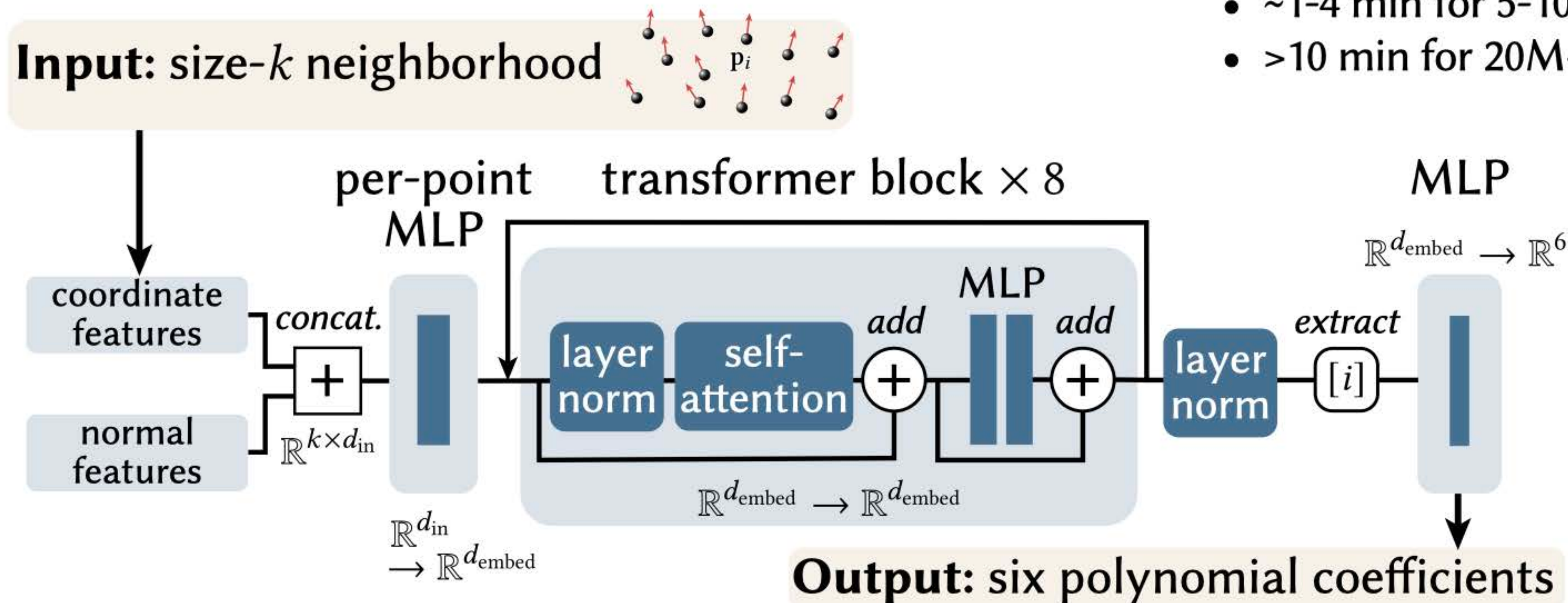
- <10s for 100k points
- <60s for 1M points
- ~1-4 min for 5-10M points
- >10 min for 20M+ points

Future directions

- Accelerating pre-computation (torus fitting)

Precompute times using one NVIDIA RTX 3090:

- <10s for 100k points
- <60s for 1M points
- ~1-4 min for 5-10M points
- >10 min for 20M+ points

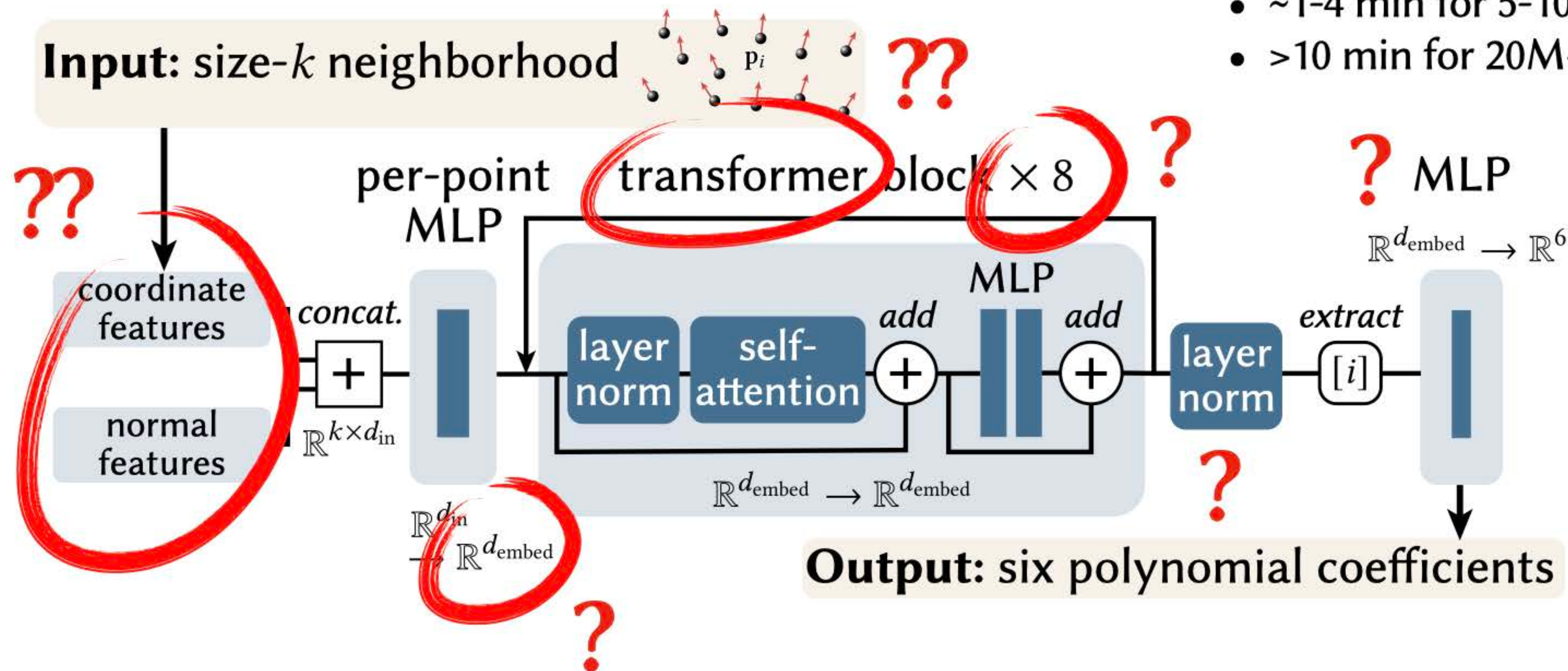


Future directions

- Accelerating pre-computation (torus fitting)

Precompute times using one NVIDIA RTX 3090:

- <10s for 100k points
- <60s for 1M points
- ~1-4 min for 5-10M points
- >10 min for 20M+ points

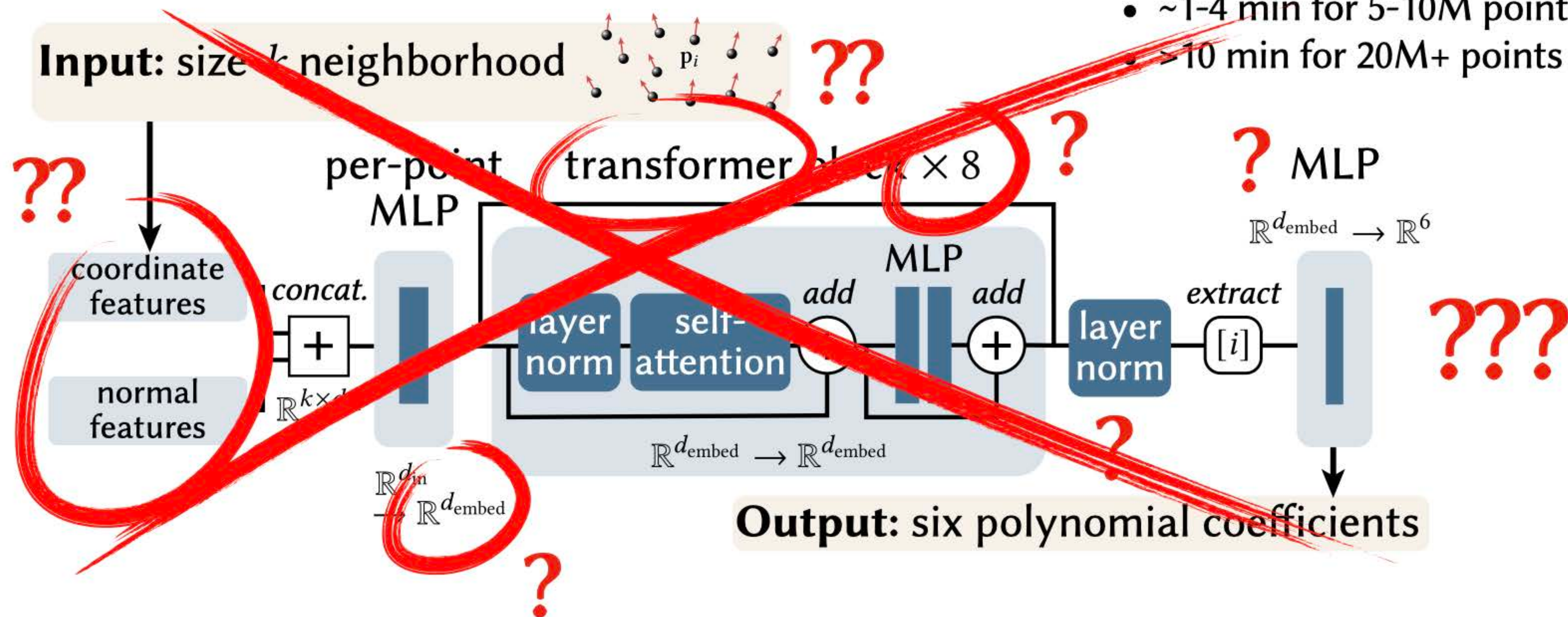


Future directions

- Accelerating pre-computation (torus fitting)

Precompute times using one NVIDIA RTX 3090:

- <10s for 100k points
- <60s for 1M points
- ~1-4 min for 5-10M points
- >10 min for 20M+ points

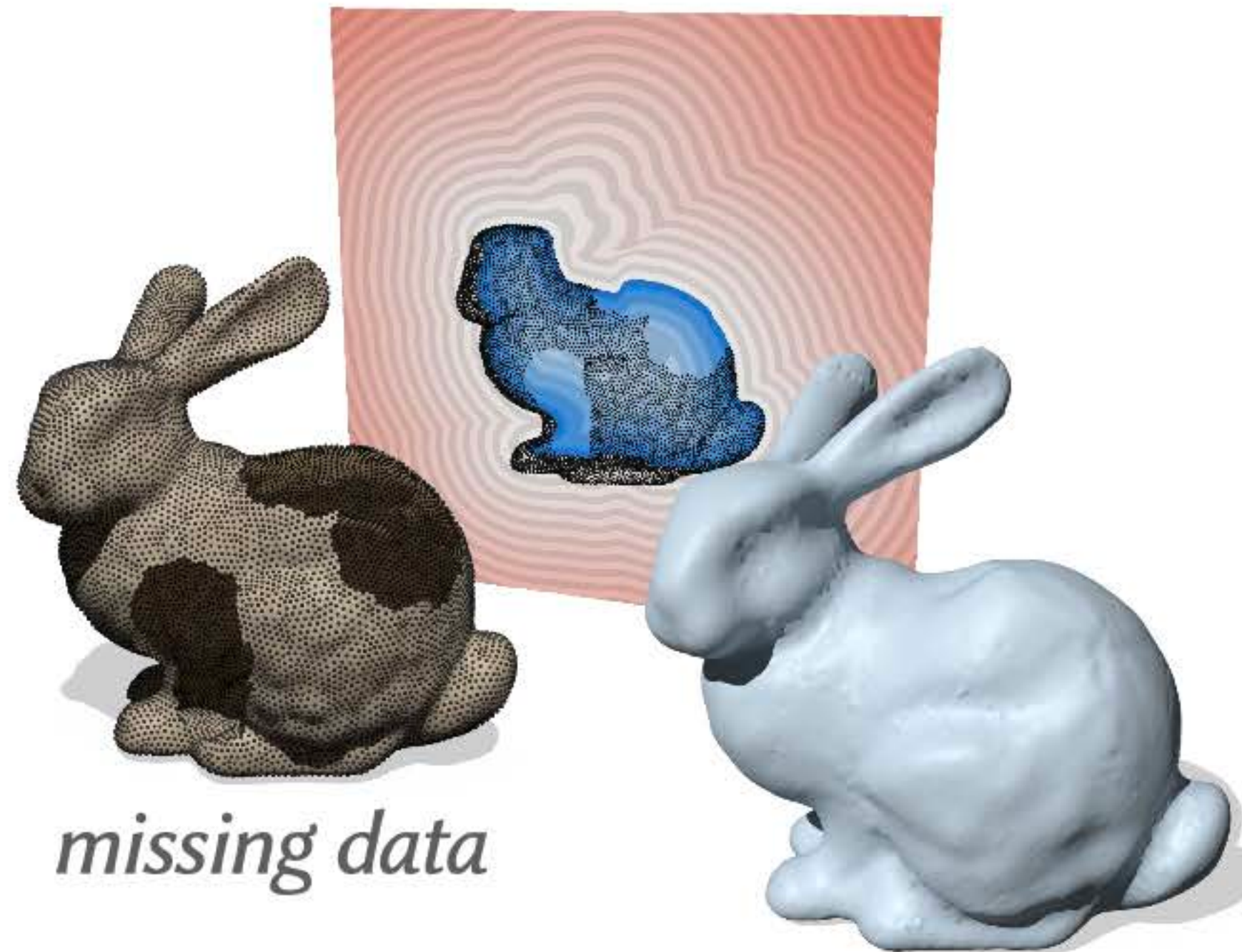
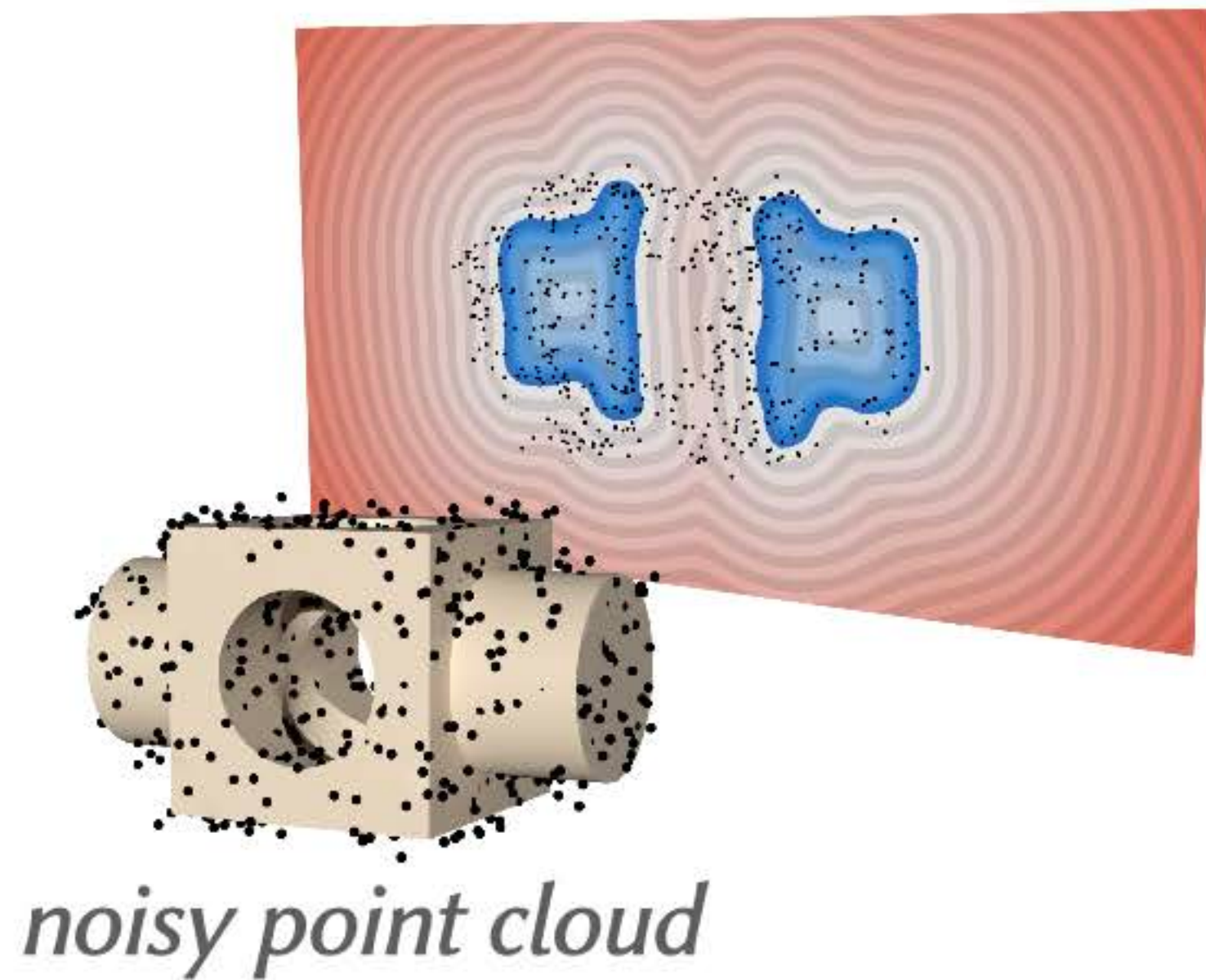


Future directions

- Accelerating pre-computation (torus fitting)
- Robustness to corruption (noise, missing data, outliers, etc.)

Future directions

- Accelerating pre-computation (torus fitting)
- Robustness to corruption (noise, missing data, outliers, etc.)



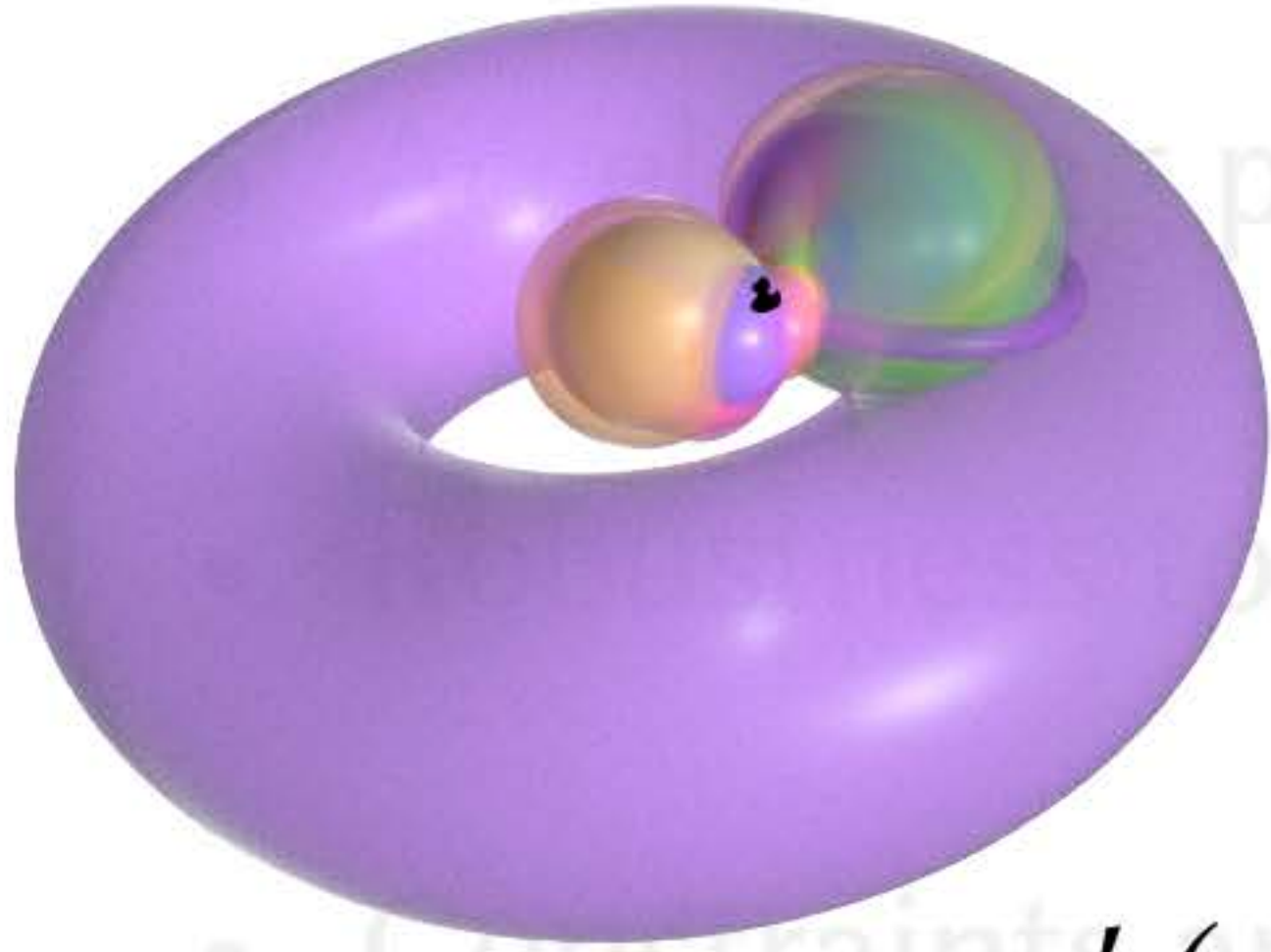
Future directions

- Accelerating pre-computation (torus fitting)
- Robustness to corruption (noise, missing data, outliers, etc.)

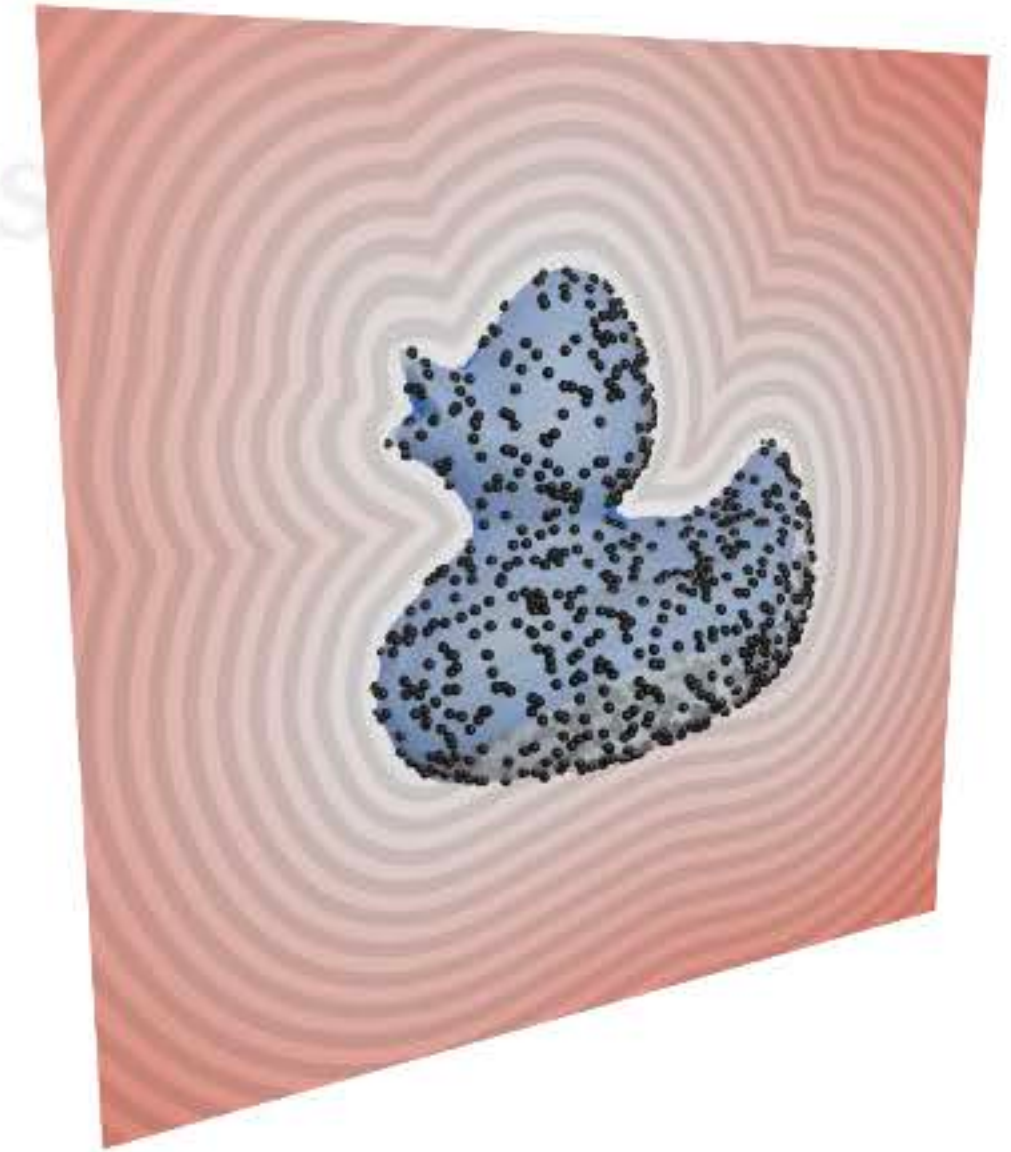
Future directions

- Accelerating pre-computation (torus fitting)
- Robustness to corruption (noise, missing data, outliers, etc.)
- Constraints and priors
- Compression/simplification
- ...

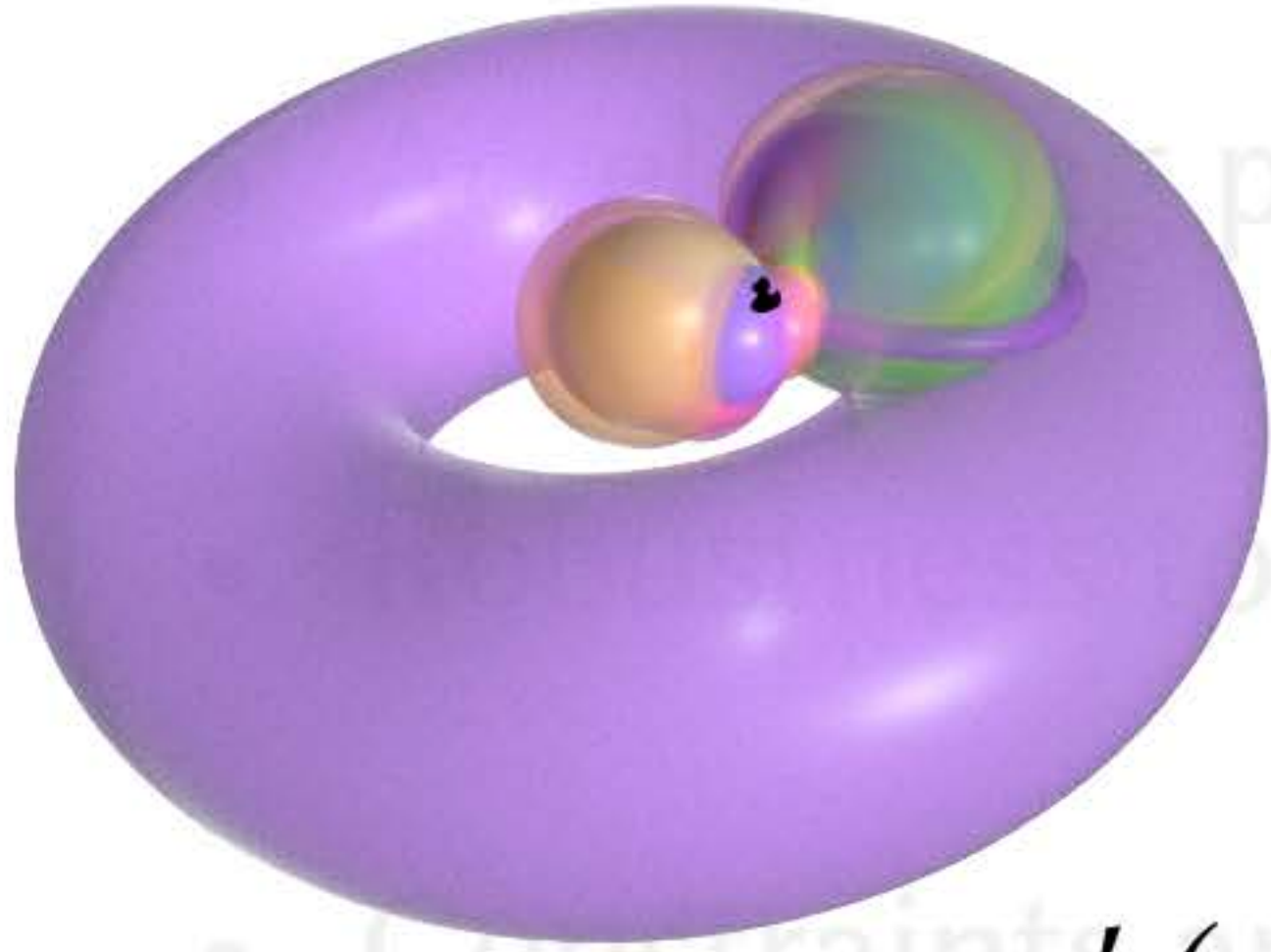
Future directions



$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$



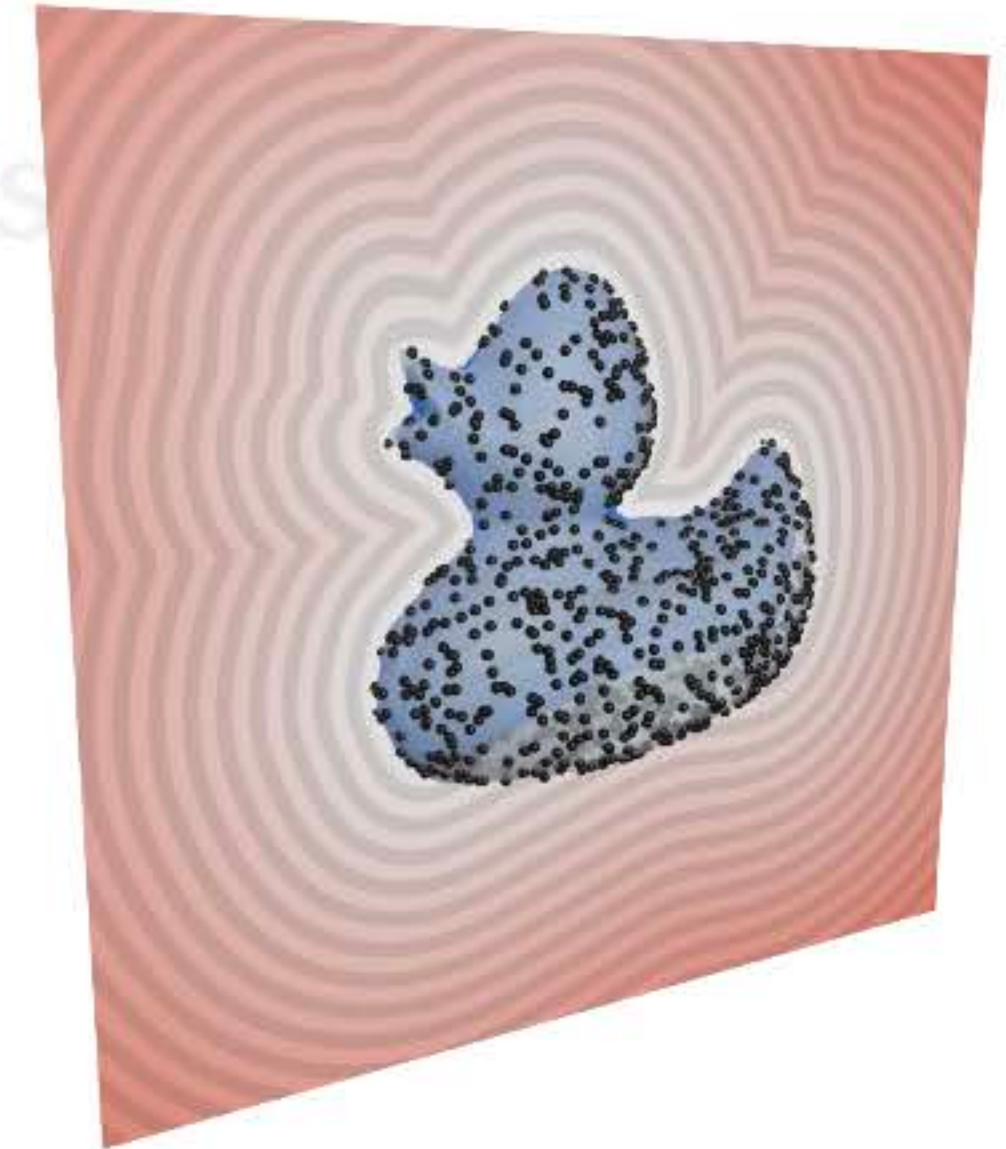
Future directions



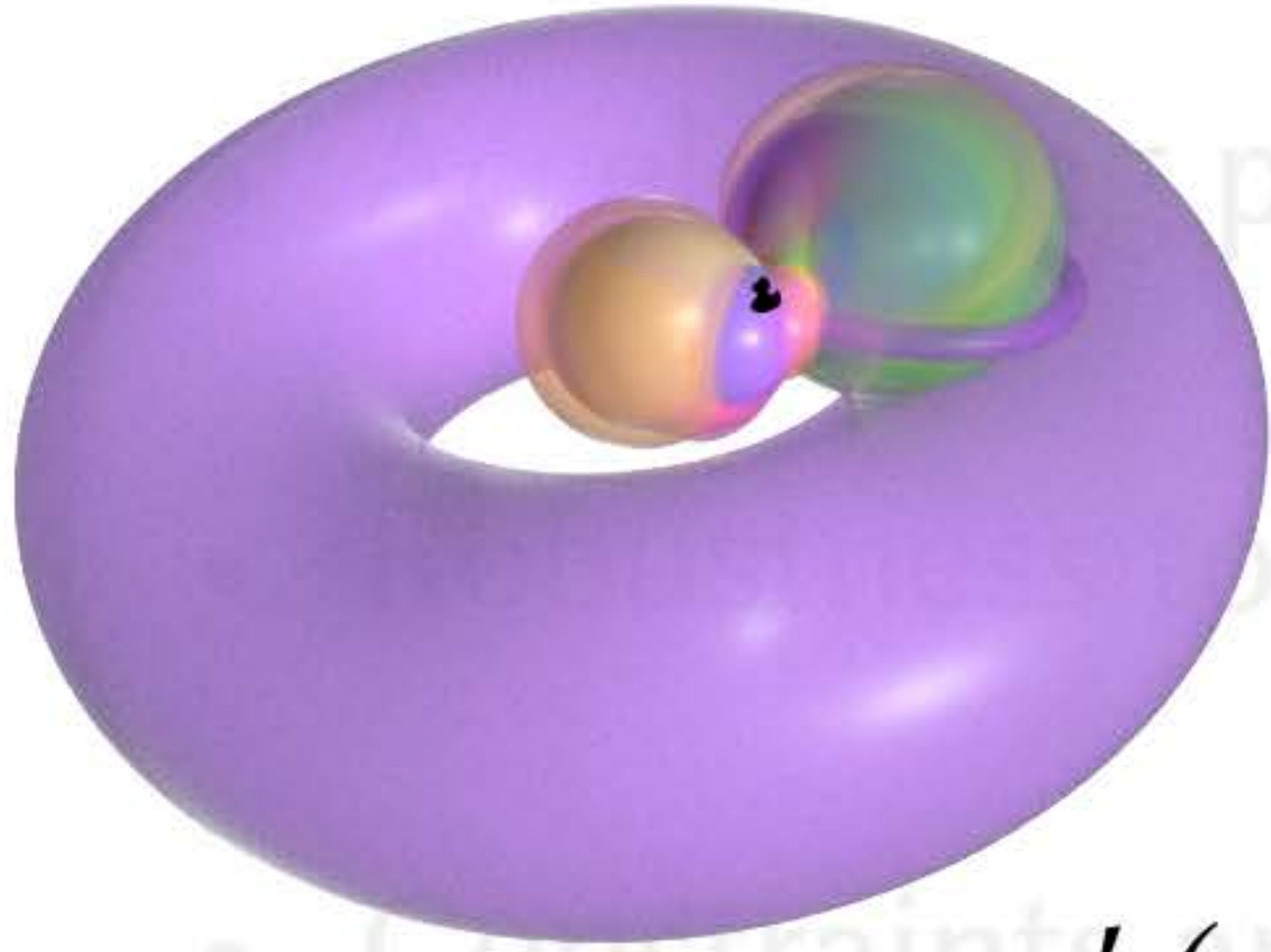
pre-computation (torus fitting)

differentiable

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$



Future directions

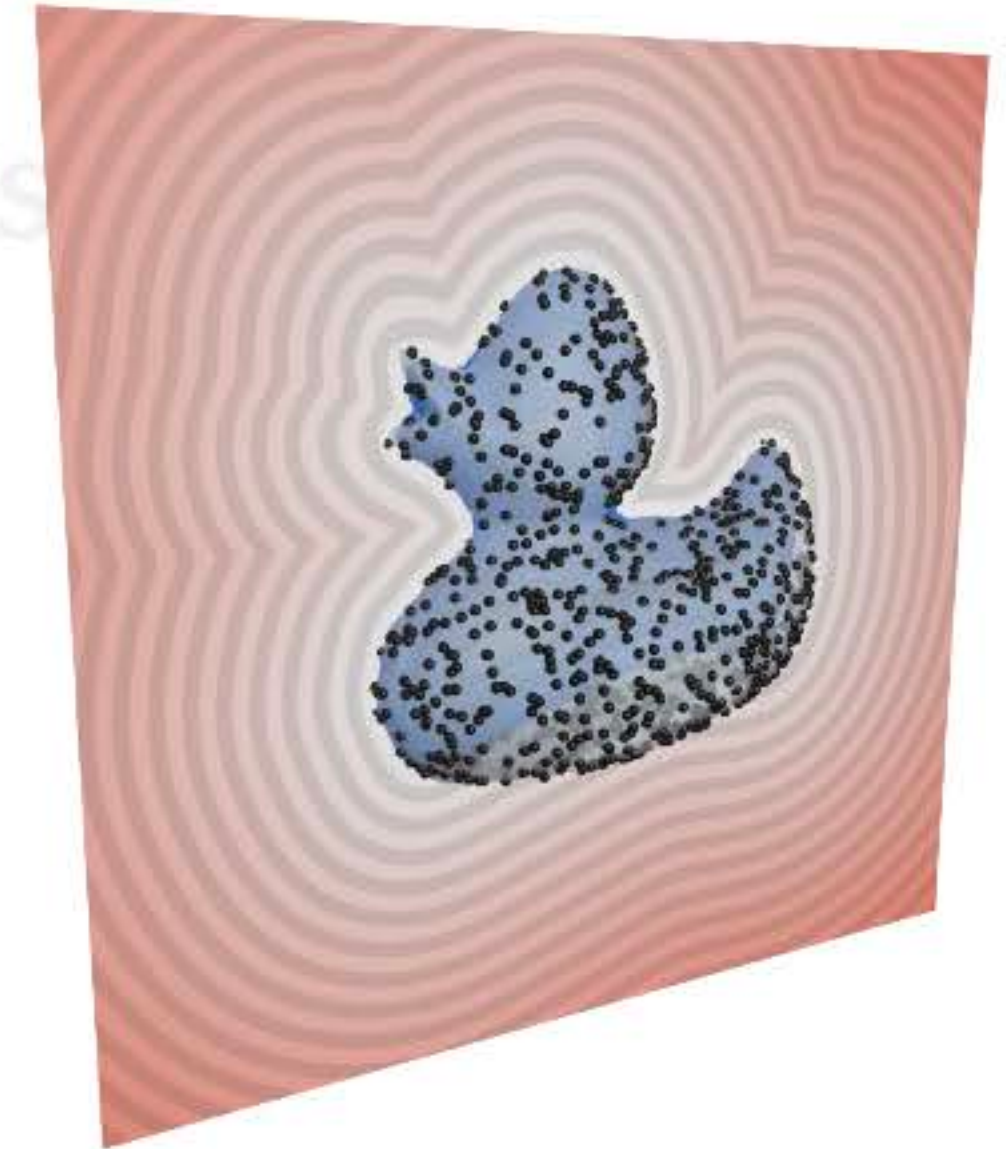


pre-computation (torus fitting)

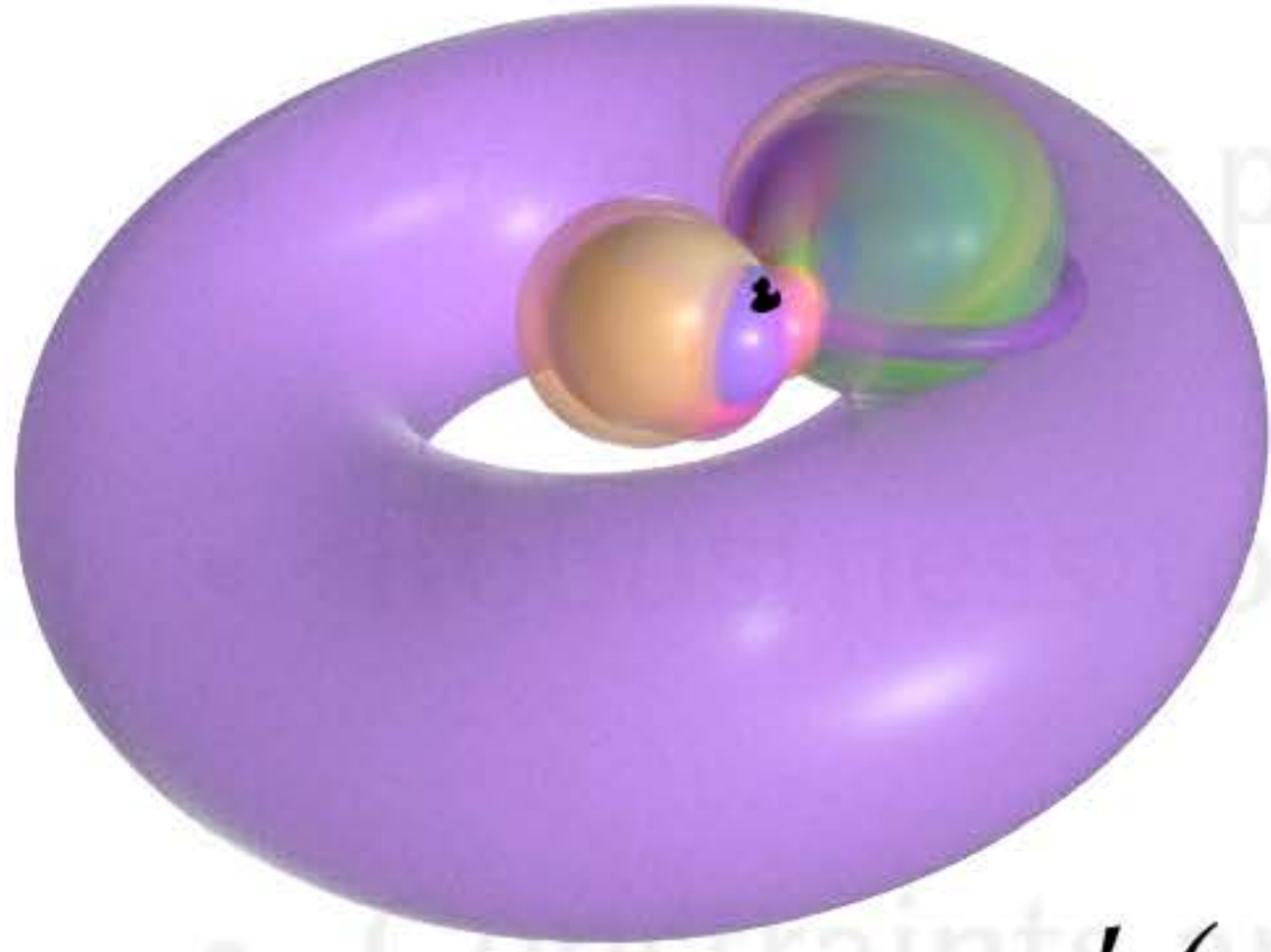
differentiable

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

point-based



Future directions

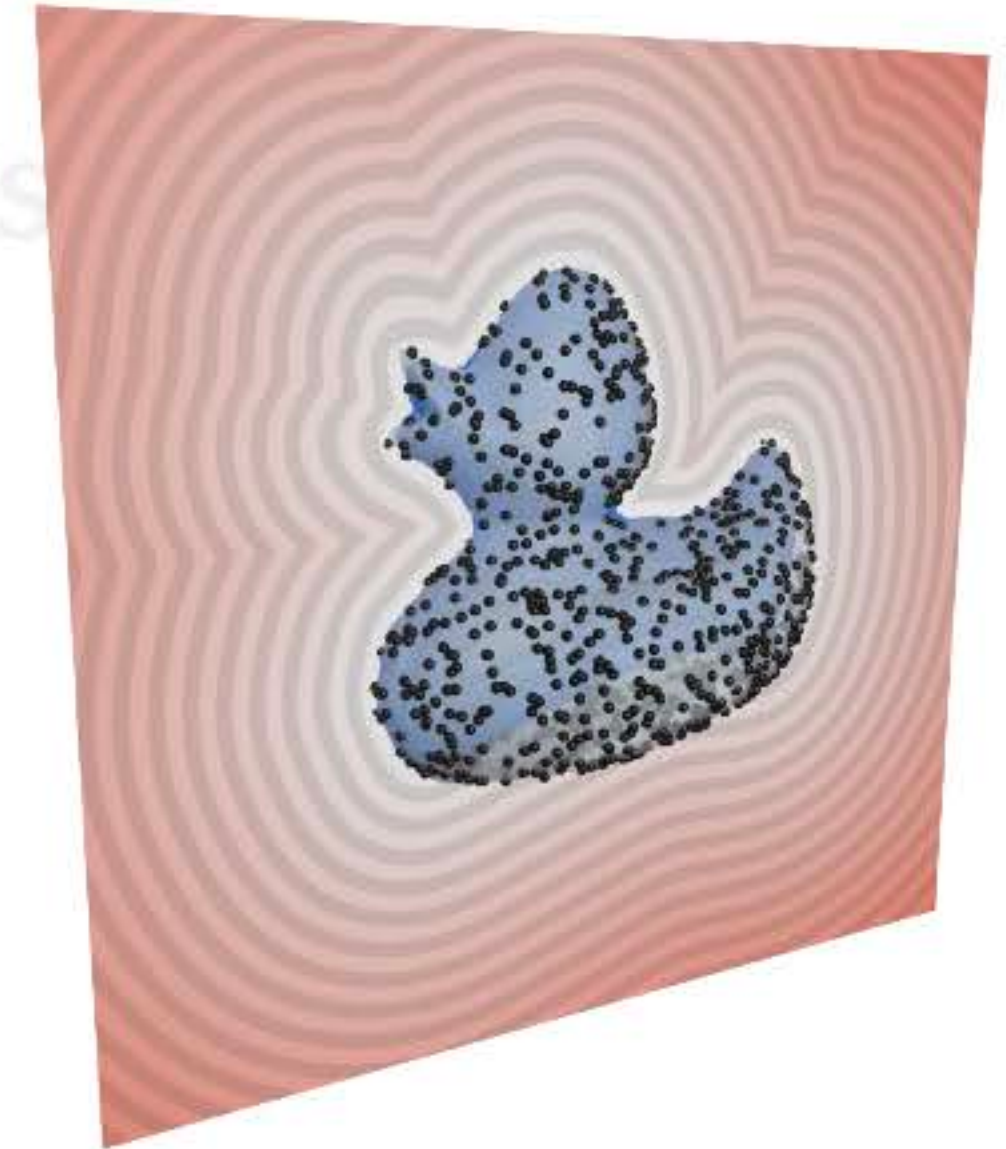


point-based

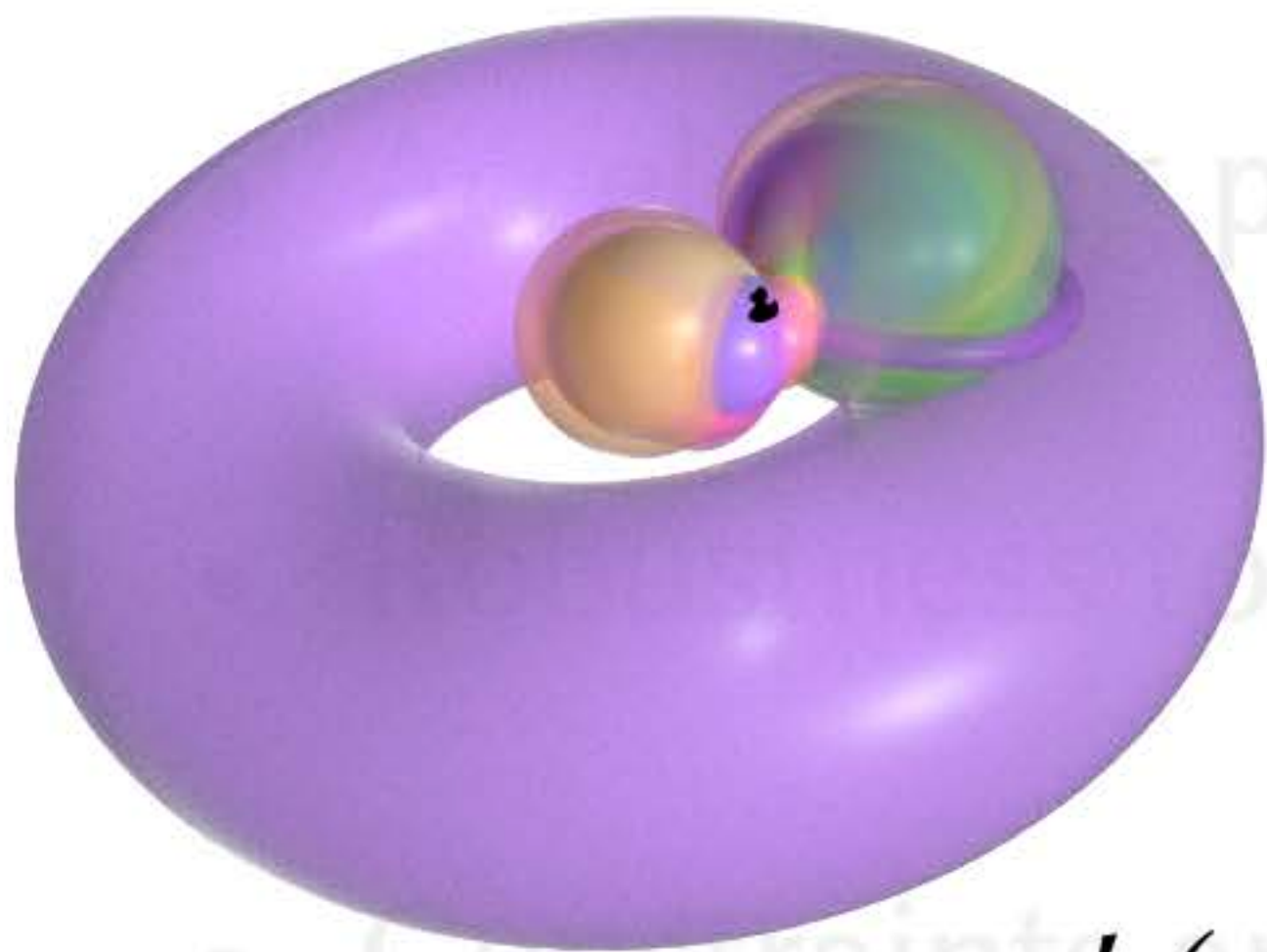
differentiable

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

output-sensitive



Future directions

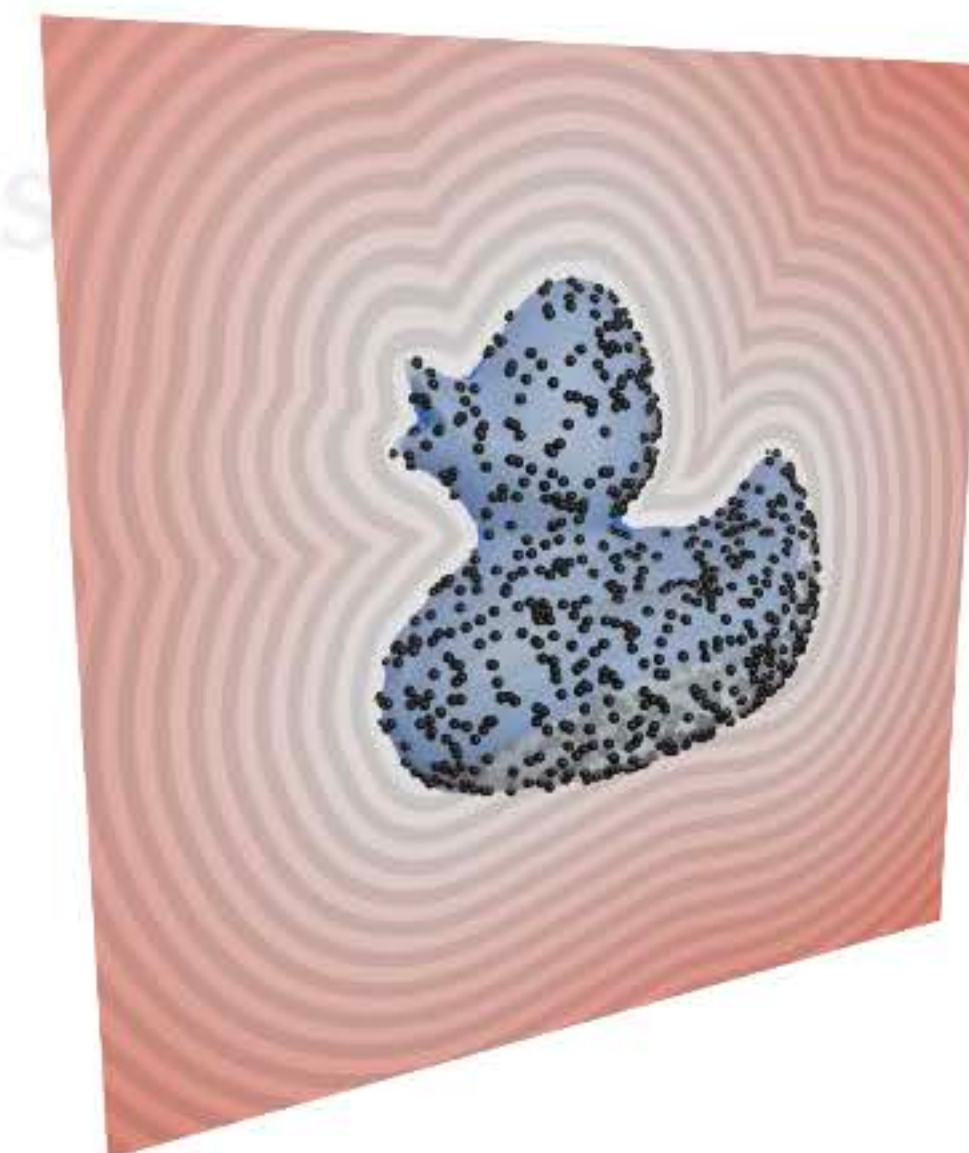


differentiable

$$\phi(x) = \frac{\sum_{i=1}^{|P|} g_i(x) \exp(-\lambda \|x - p_i\|)}{\sum_{i=1}^{|P|} \exp(-\lambda \|x - p_i\|)}$$

output-sensitive

point-based



- Torus signed distance as a differentiable shape representation

THANKS!