



Pacific Ballroom #137



Neural Inverse Knitting: From Images to Manufacturing Instruction

Alexandre Kaspar*, Tae-Hyun Oh*, Liane Makatura,
Petr Kellnhofer and Wojciech Matusik

MIT CSAIL

Industrial Knitting



SHIMA SEIKI WHOLEGARMENT™

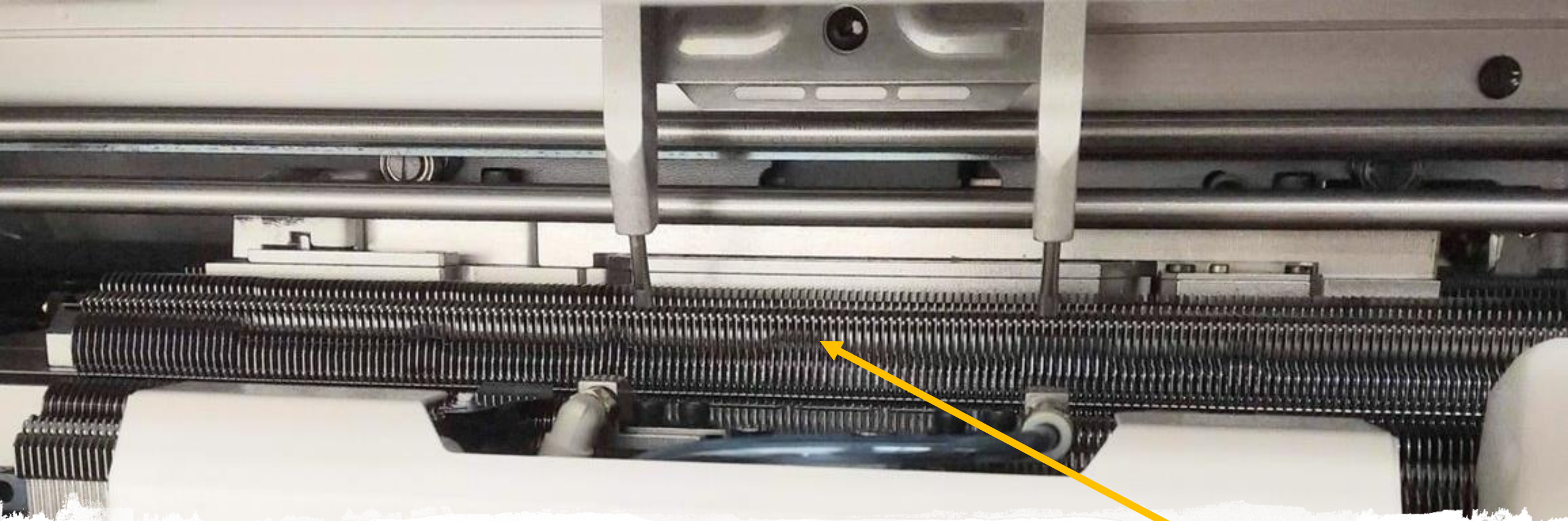
SWG 091 N₂

Pacific Ballroom #137, <http://deepknitting.csail.mit.edu>



Industrial Knitting

- Whole garments from scratch



Industrial Knitting

- Control of individual needles
- Whole garments from scratch

Knitted Garment & Patterns

Many garments are knitted:

- Beanies, scarves
- Gloves, socks and underwear
- Sweaters, sweatpants

Current machines can create those garments **seamlessly** (no sewing needed).



Knitted Garment & Patterns

Those garments have **various types** of surface **patterns** (knitting patterns).

These can be fully controlled by industrial knitting machine.

= **User customization!**

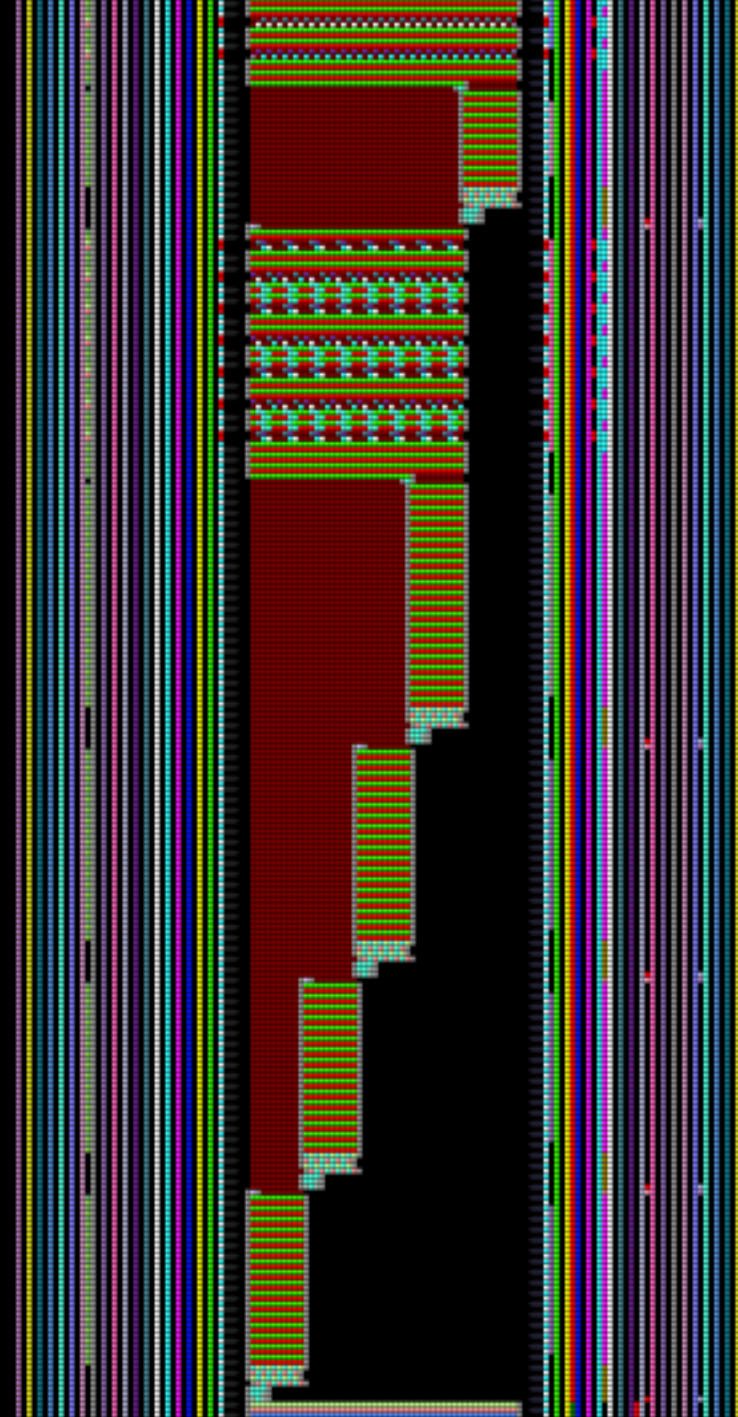


Machine Knitting Programming

Low-level machine code
requires skilled experts
= knitting masters

Good news

- Many hand knitting patterns available online and in books
- Online communities of knitting enthusiasts sharing patterns



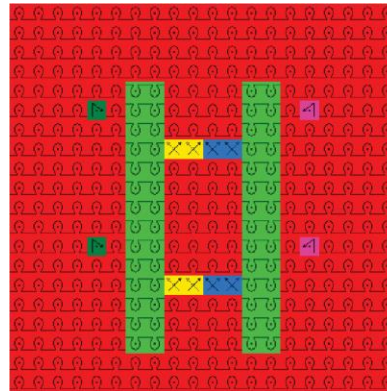
Scenario

1. User takes picture of knitting pattern



Scenario

1. User takes picture of knitting pattern
2. System creates knitting instructions

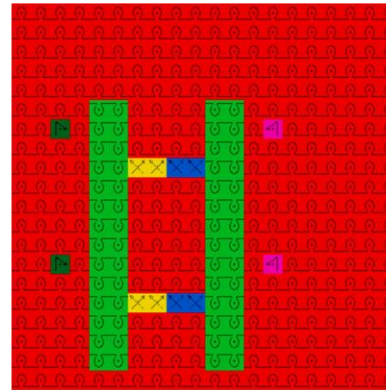


← Inverse Neural Knitting →

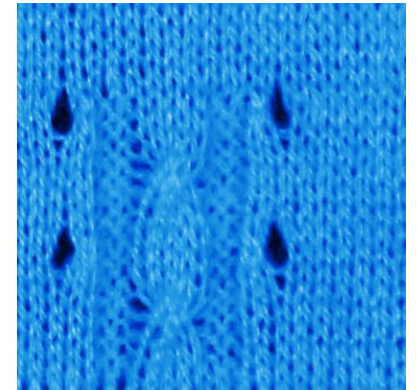


Scenario

1. User takes picture of knitting pattern
2. System creates knitting instructions
3. User reuses pattern for new garment

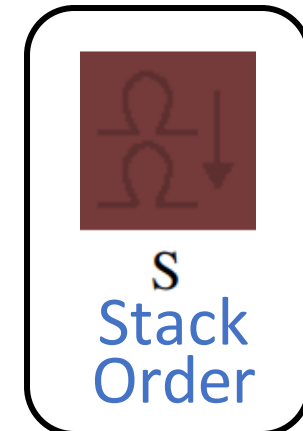
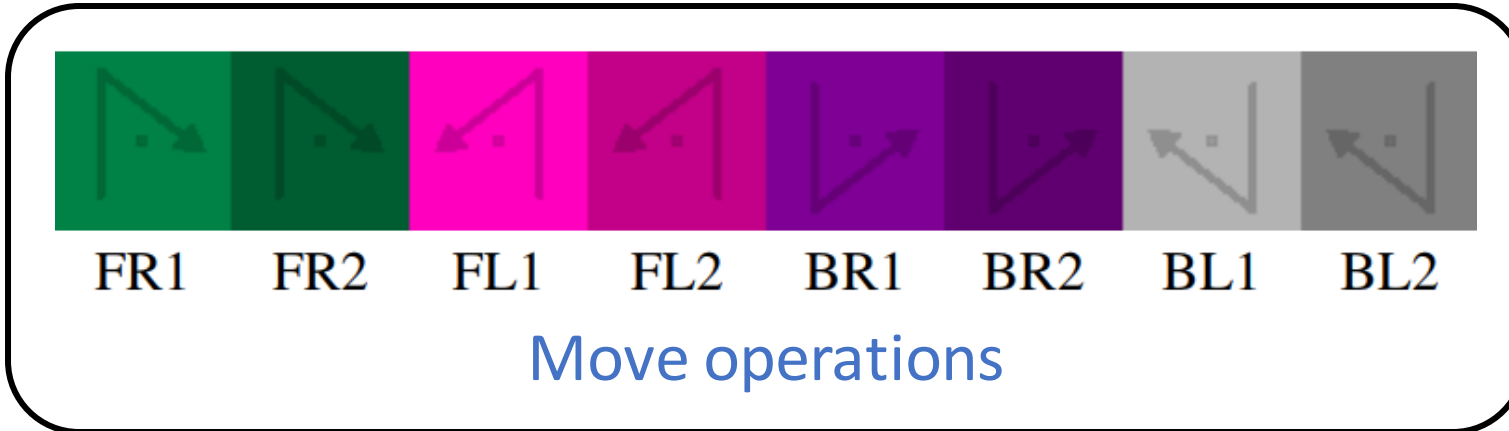
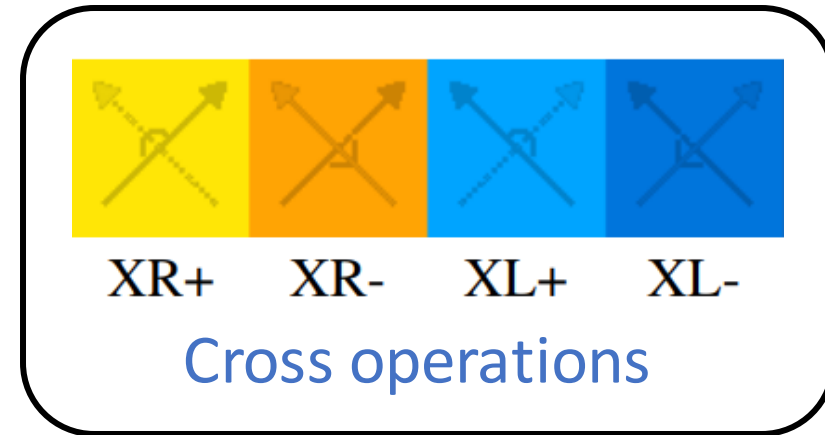
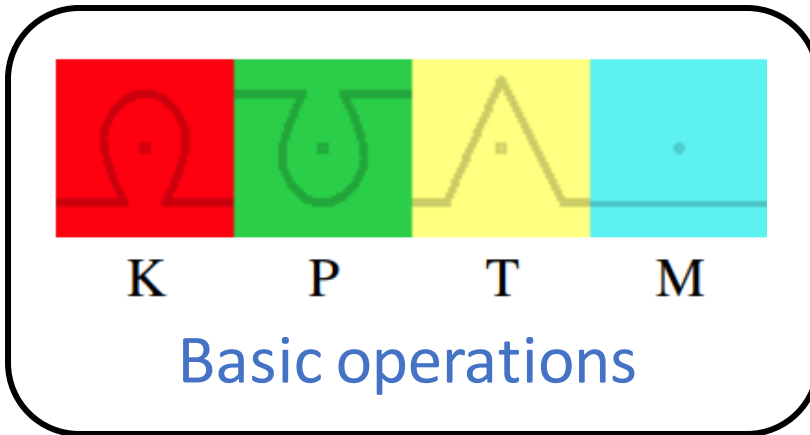


Machine
Knitting



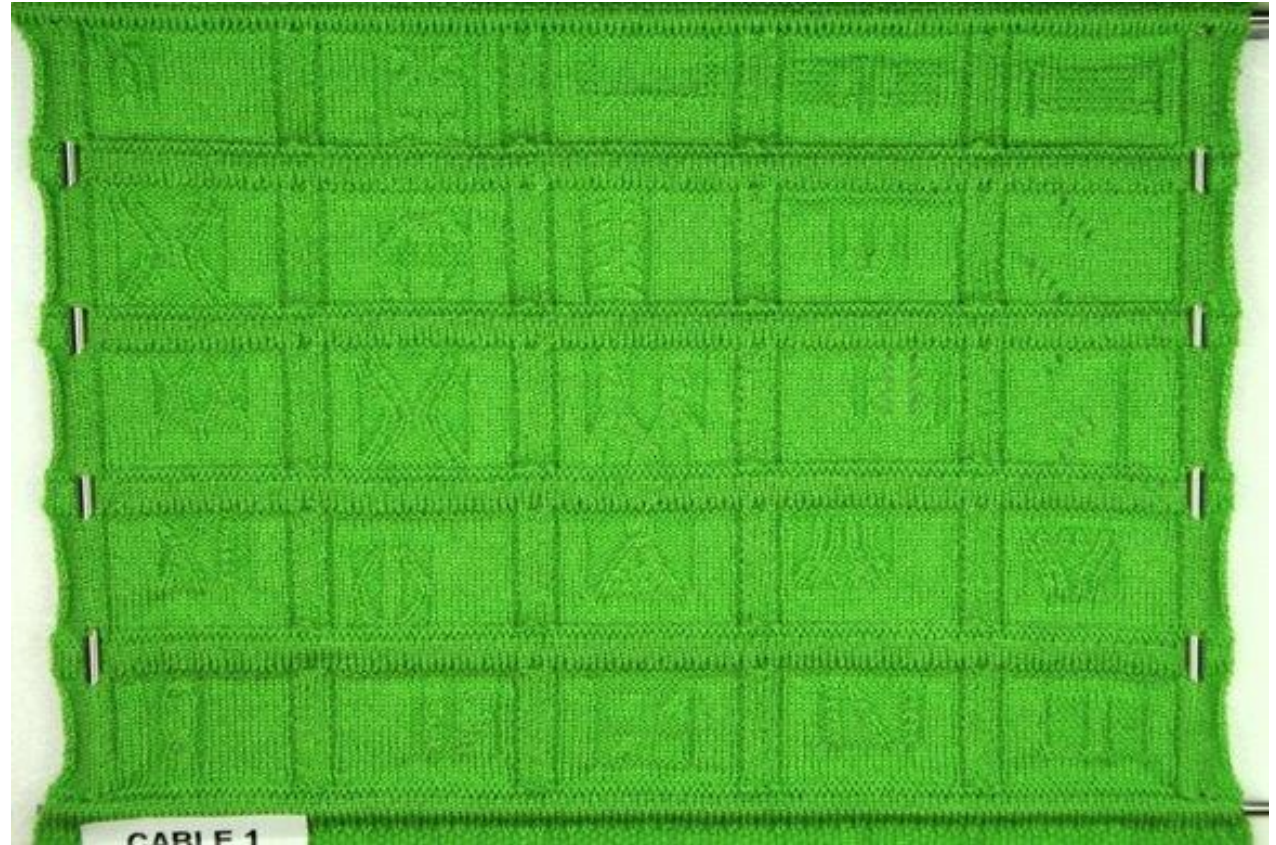
Dataset: DSL

Domain Specific Language (DSL) for regular knitting patterns



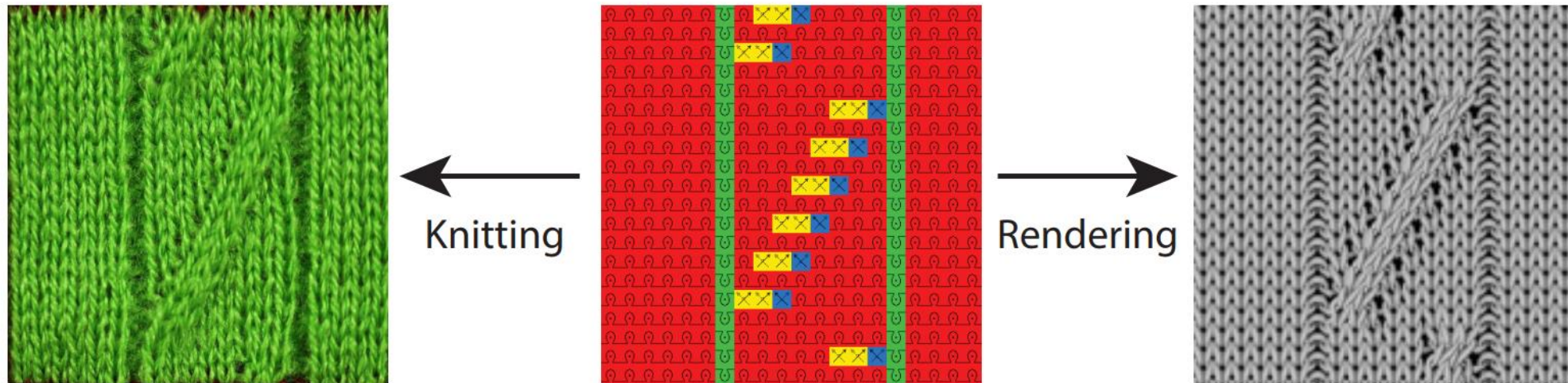
Dataset: Capture

Capture setup with steel rods to normalize tension



Dataset Content

- Paired instructions with real (2,088) and synthetic (14,440) images.
- Available on project page.



Learning Problem

Mapping **images** to discrete **instruction maps**

= CE loss minimization

Using two domains of input data
(one real, one synthetic)

= How to best combine both

Generalization Bound with Two Domains

With probability at least $1 - \delta$

$$\underbrace{\frac{1}{2} |\mathcal{L}_T(\hat{h}, y) - \mathcal{L}_T(h_T^*, y)|}_{\text{Generalization gap}} \leq \alpha (\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) + \lambda) + \epsilon$$

Ideal min.

Generalization Bound with Two Domains

With probability at least $1 - \delta$

$$\frac{1}{2} \underbrace{|\mathcal{L}_T(\hat{h}, y) - \mathcal{L}_T(h_T^*, y)|}_{\text{Generalization gap}} \leq \alpha (\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) + \lambda) + \epsilon$$

Ideal min.

Empirical min. $\arg \min_h \alpha \mathcal{L}_{\hat{S}}(h, y) + (1 - \alpha) \mathcal{L}_{\hat{T}}(h, y)$

Generalization Bound with Two Domains

With probability at least $1 - \delta$

$$\begin{aligned} \frac{1}{2} |\mathcal{L}_T(\hat{h}, y) - \mathcal{L}_T(h_T^*, y)| \\ \leq \alpha (\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) + \lambda) + \epsilon \end{aligned}$$

Generalization Bound with Two Domains

With probability at least $1 - \delta$

$$\frac{1}{2} |\mathcal{L}_T(\hat{h}, y) - \mathcal{L}_T(h_T^*, y)| \leq \alpha (\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) + \lambda) + \epsilon$$

$$\epsilon(m, \alpha, \beta, \delta) = \sqrt{\frac{1}{2m} \left(\frac{\alpha^2}{\beta} + \frac{(1-\alpha)^2}{1-\beta} \right) \log\left(\frac{2}{\delta}\right)},$$

Parameter dependent term

Generalization Bound with Two Domains

With probability at least $1 - \delta$

$$\frac{1}{2} |\mathcal{L}_T(\hat{h}, y) - \mathcal{L}_T(h_T^*, y)| \leq \alpha (\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) + \lambda) + \epsilon$$

$$\lambda = \min_{h \in \mathcal{H}} \mathcal{L}_S(h, y) + \mathcal{L}_T(h, y).$$

Ideal error of the combined losses

Generalization Bound with Two Domains

With probability at least $1 - \delta$

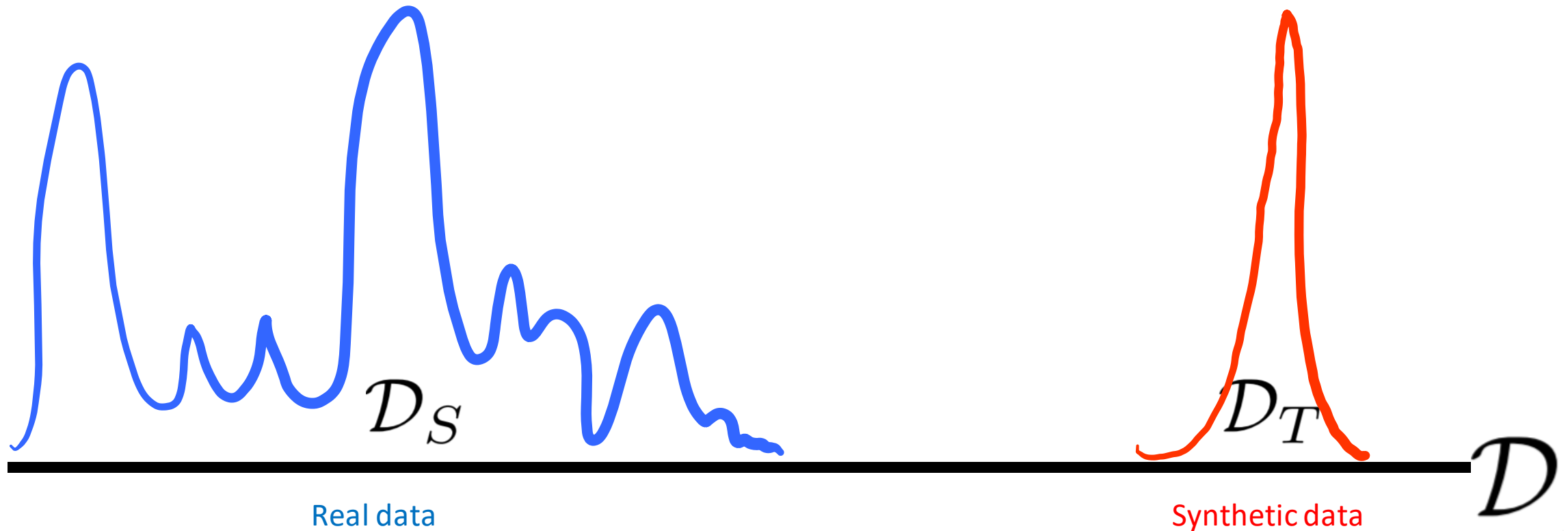
$$\frac{1}{2} |\mathcal{L}_T(\hat{h}, y) - \mathcal{L}_T(h_T^*, y)| \leq \alpha (\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) + \lambda) + \epsilon$$

Discrepancy between distributions

$$\text{disc}_{\mathcal{H}}(\mathcal{D}_S, \mathcal{D}_T) = \max_{h, h' \in \mathcal{H}} |\mathcal{L}_{\mathcal{D}_S}(h, h') - \mathcal{L}_{\mathcal{D}_T}(h, h')|$$

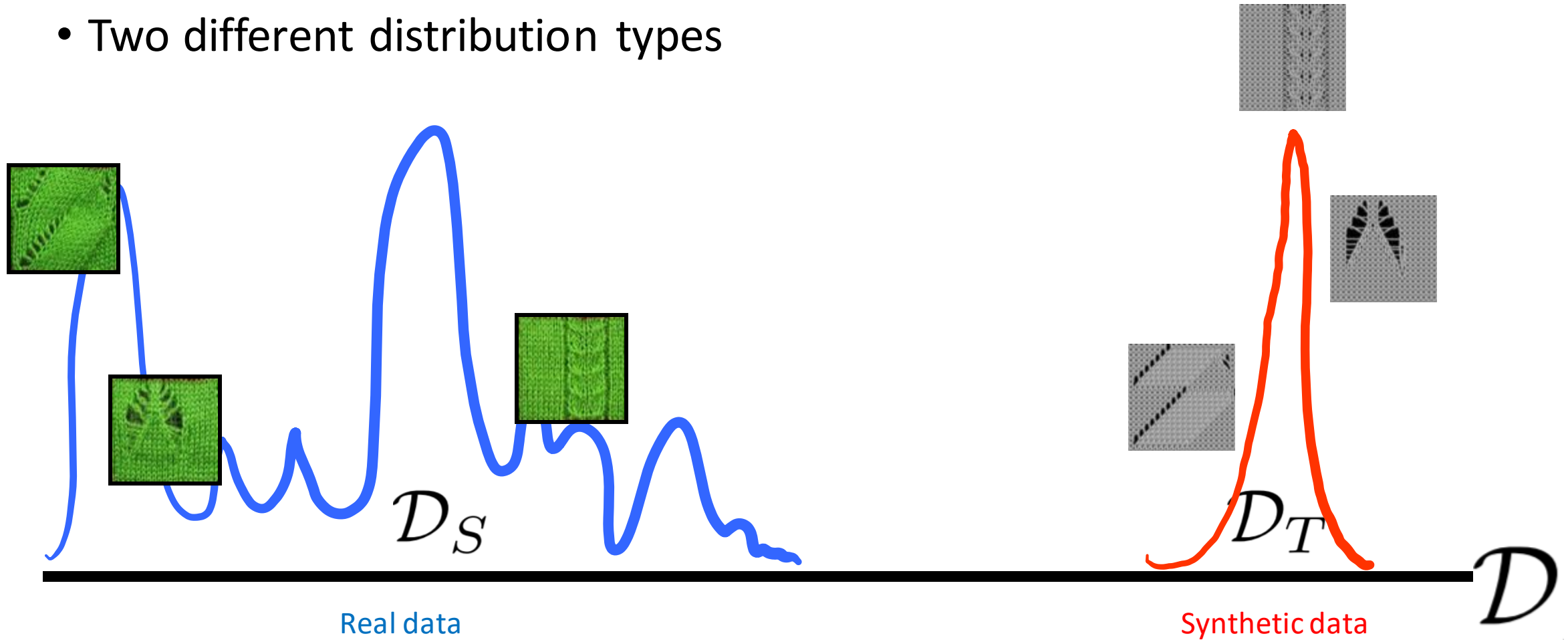
Data distributions

- Two different distribution types



Data distributions

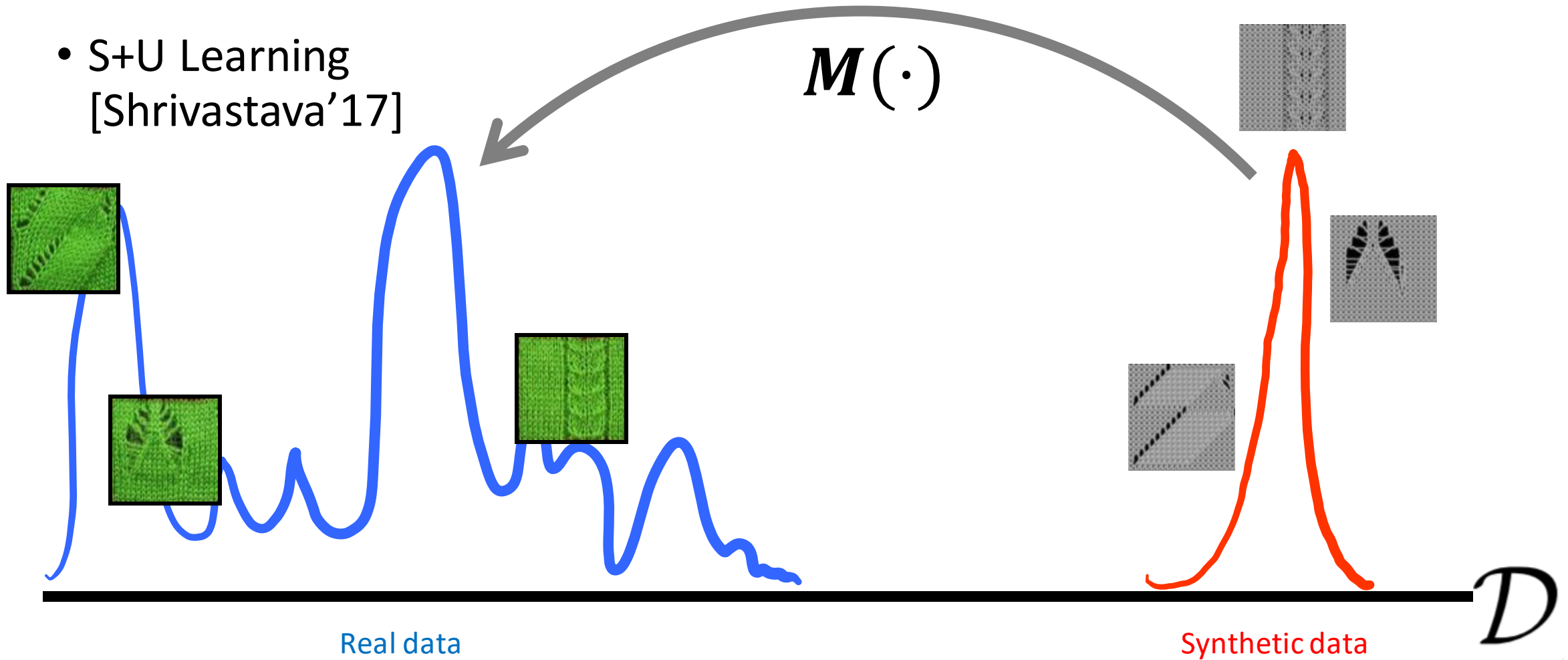
- Two different distribution types



From synthetic to real

$$\min_M \text{disc}(\mathcal{D}_S, M(\mathcal{D}_T))_{S \leftarrow T}$$

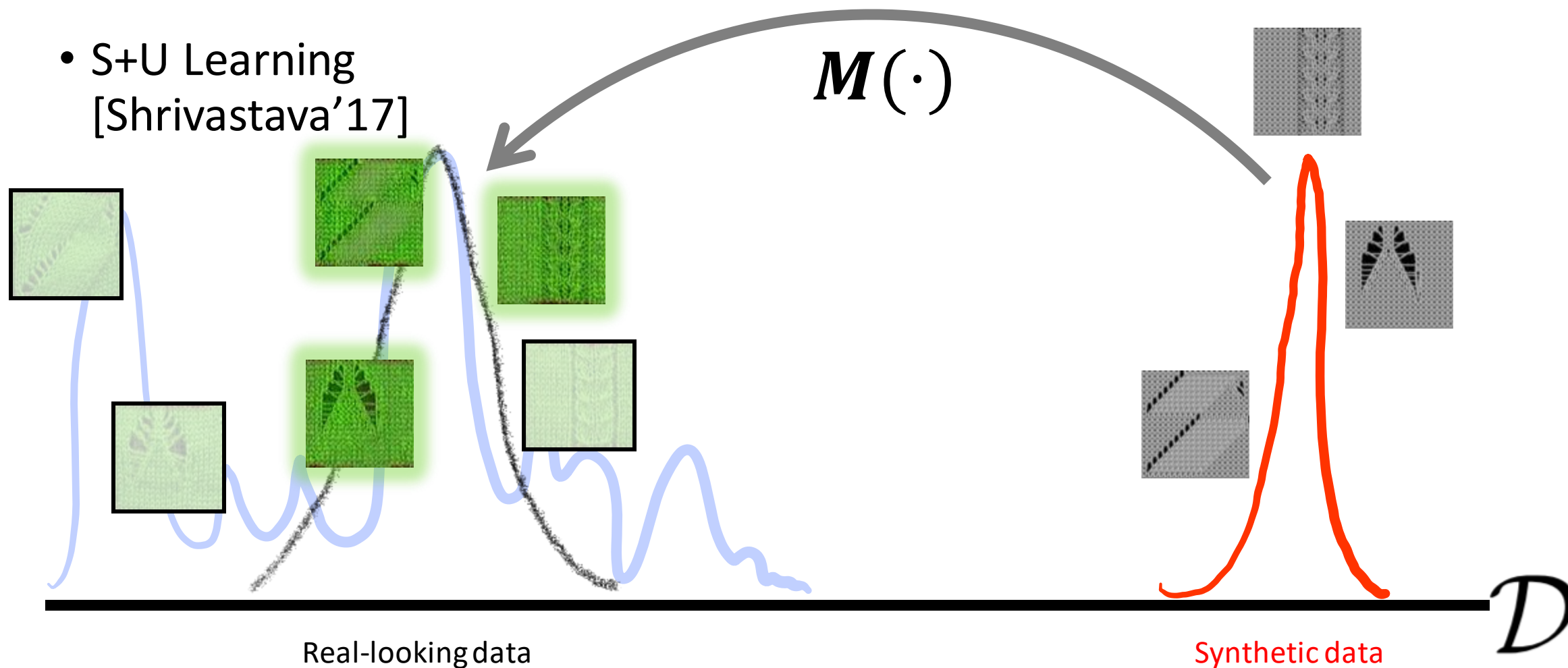
- S+U Learning [Shrivastava'17]



From synthetic to real

$$\min_M \text{disc}(\mathcal{D}_S, M(\mathcal{D}_T))_{S \leftarrow T}$$

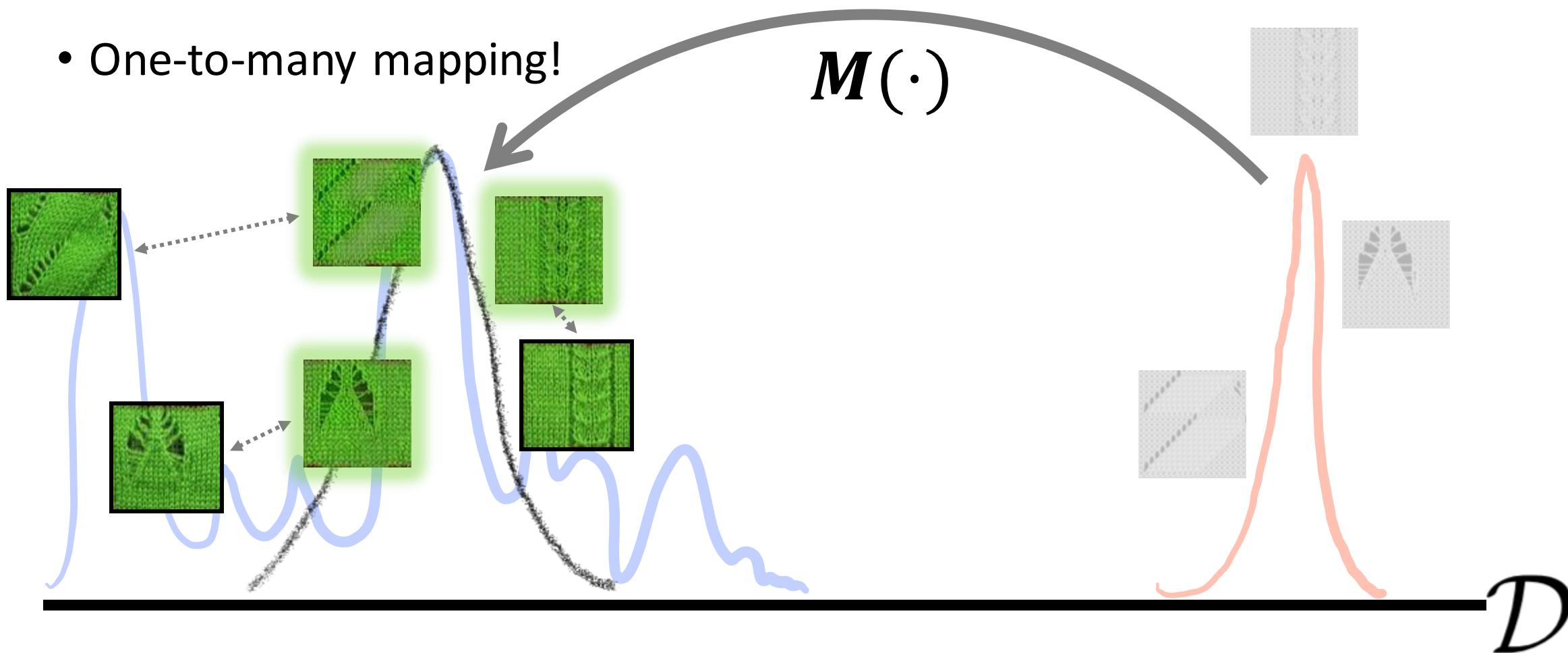
- S+U Learning [Shrivastava'17]



From synthetic to real

$$\min_M \text{disc}(\mathcal{D}_S, M(\mathcal{D}_T))_{S \leftarrow T}$$

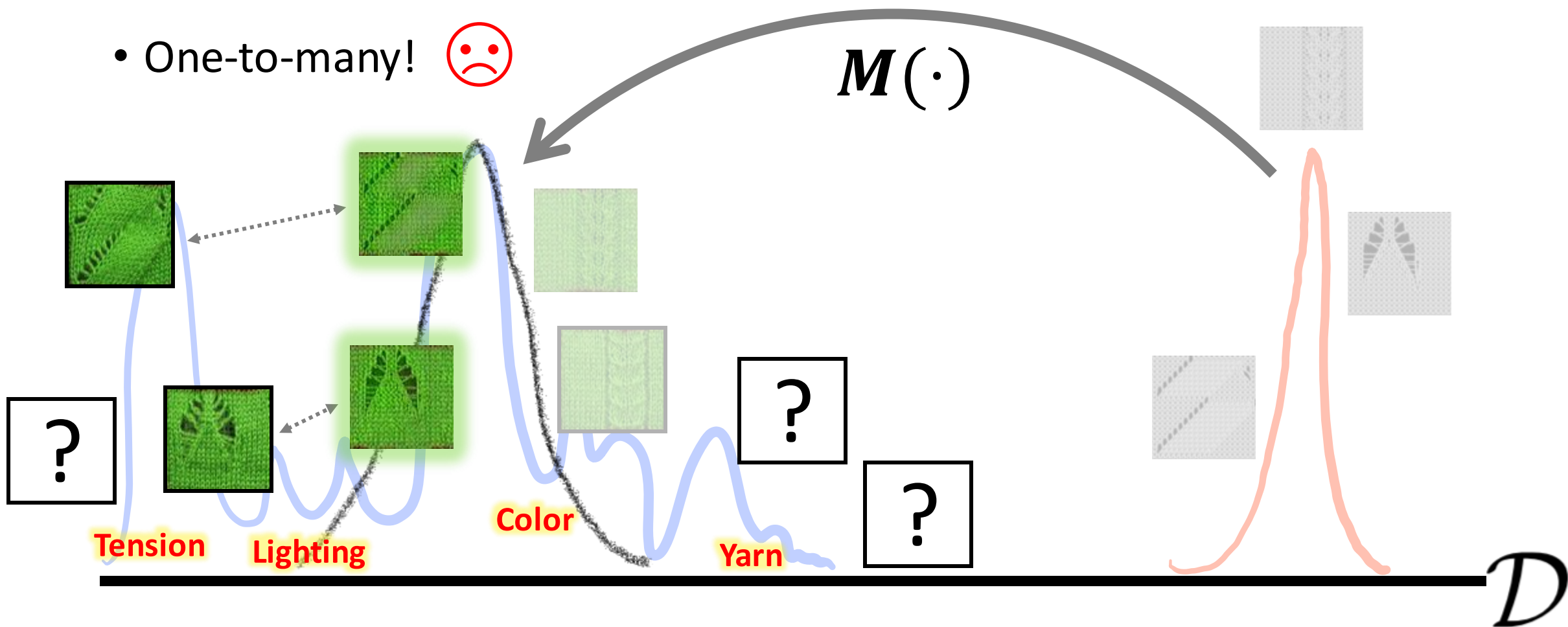
- One-to-many mapping!



From synthetic to real

$$\min_M \text{disc}(\mathcal{D}_S, M(\mathcal{D}_T))_{S \leftarrow T}$$

- One-to-many! 😞



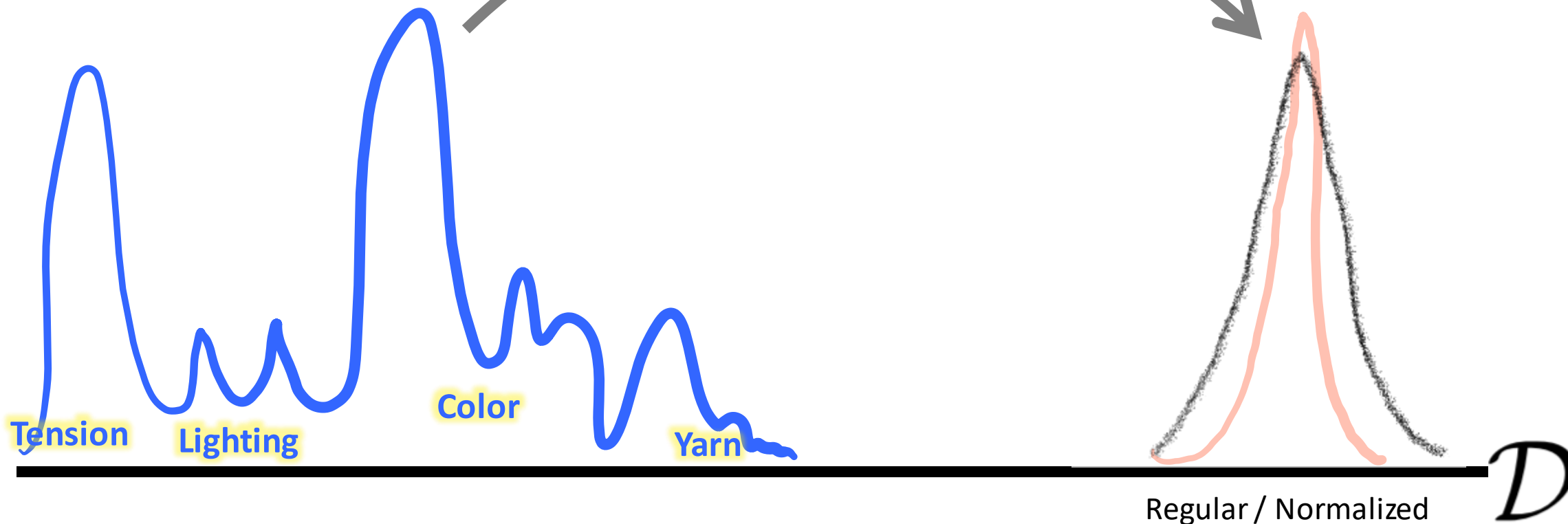
From real to synthetic

$$\min_{M} \text{disc}(M(\mathcal{D}_S), \mathcal{D}_T)_{S \Rightarrow T}$$

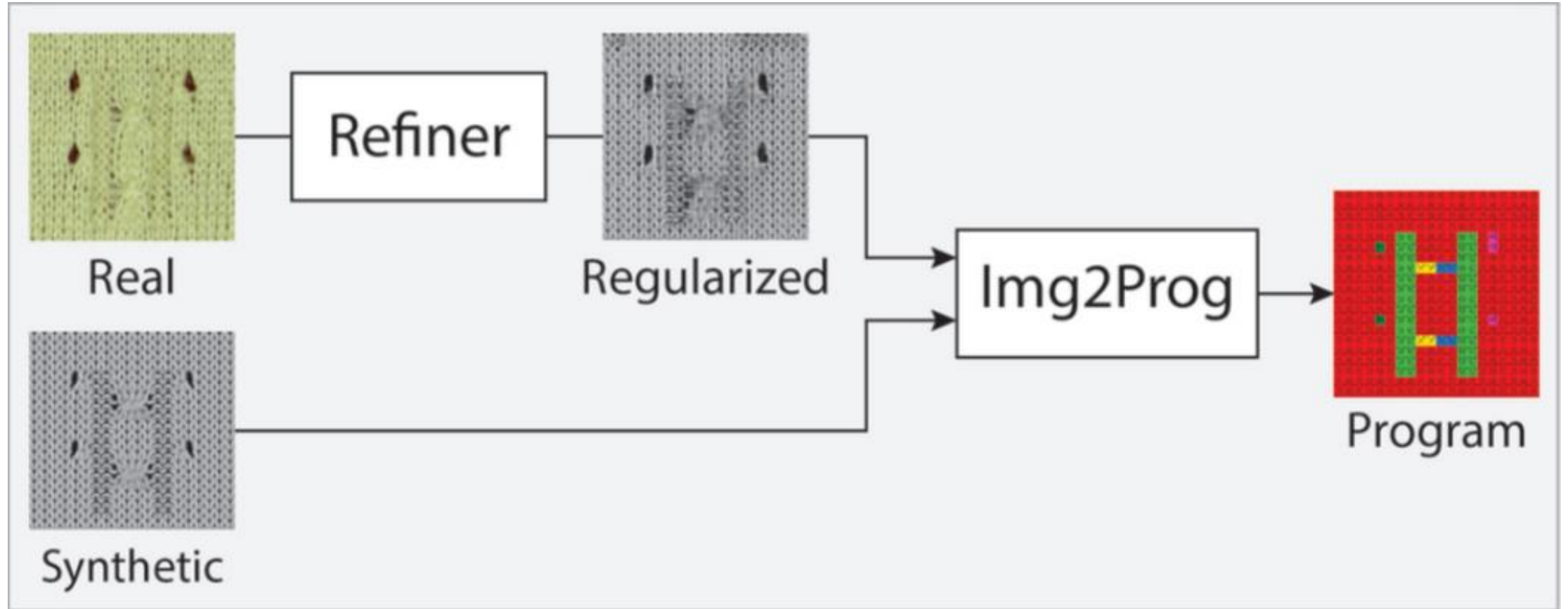
- Many-to-one!

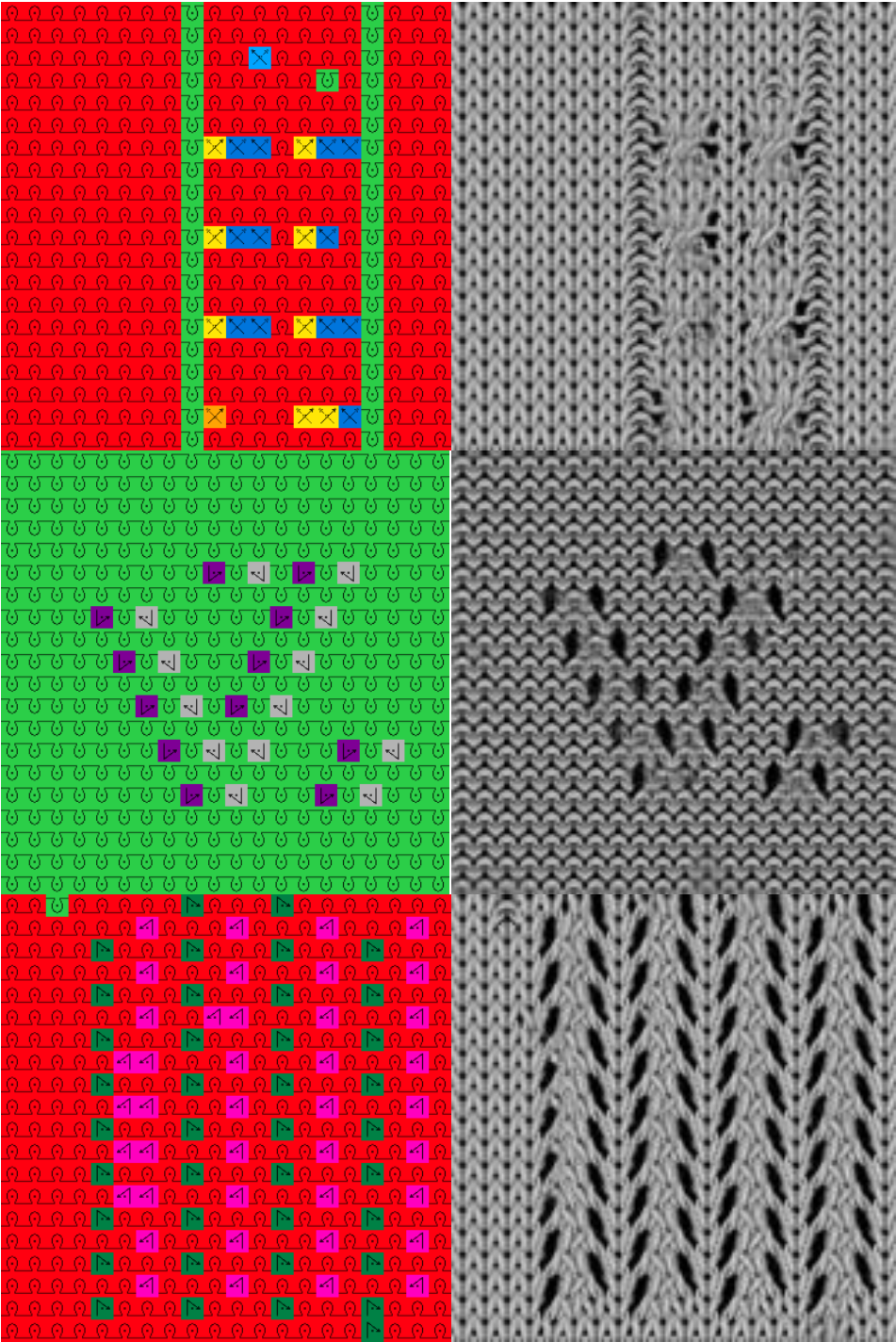


$M(\cdot)$

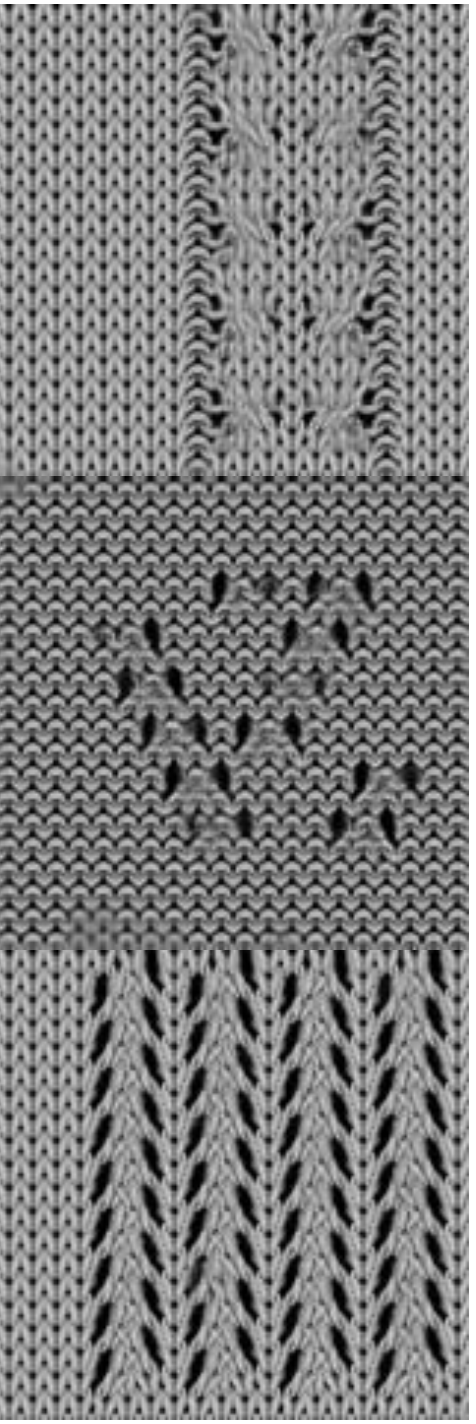


Network composition

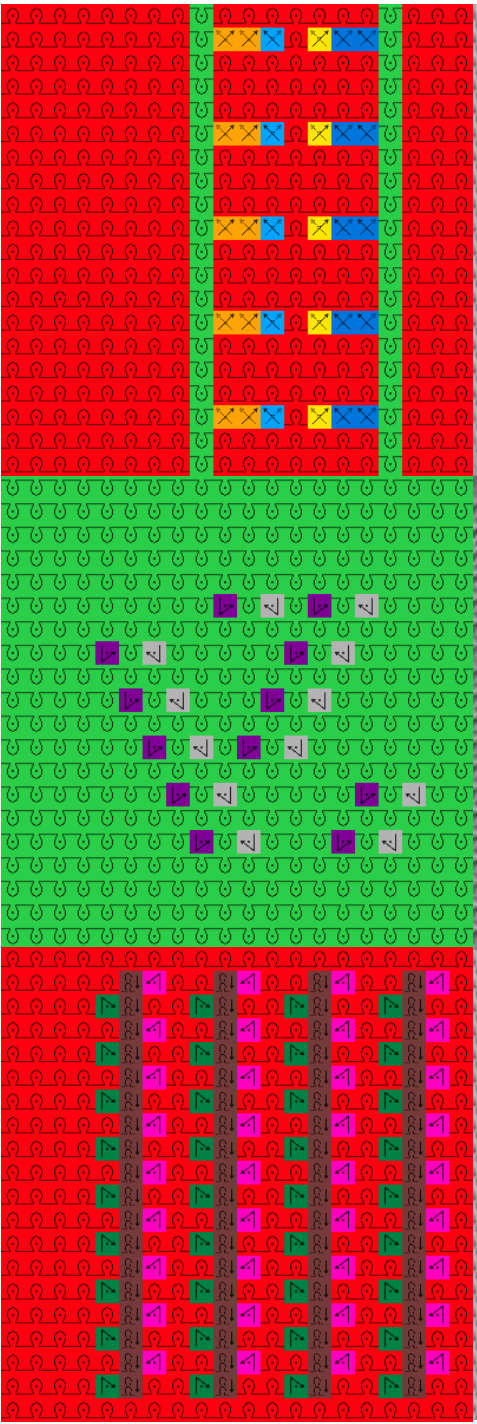




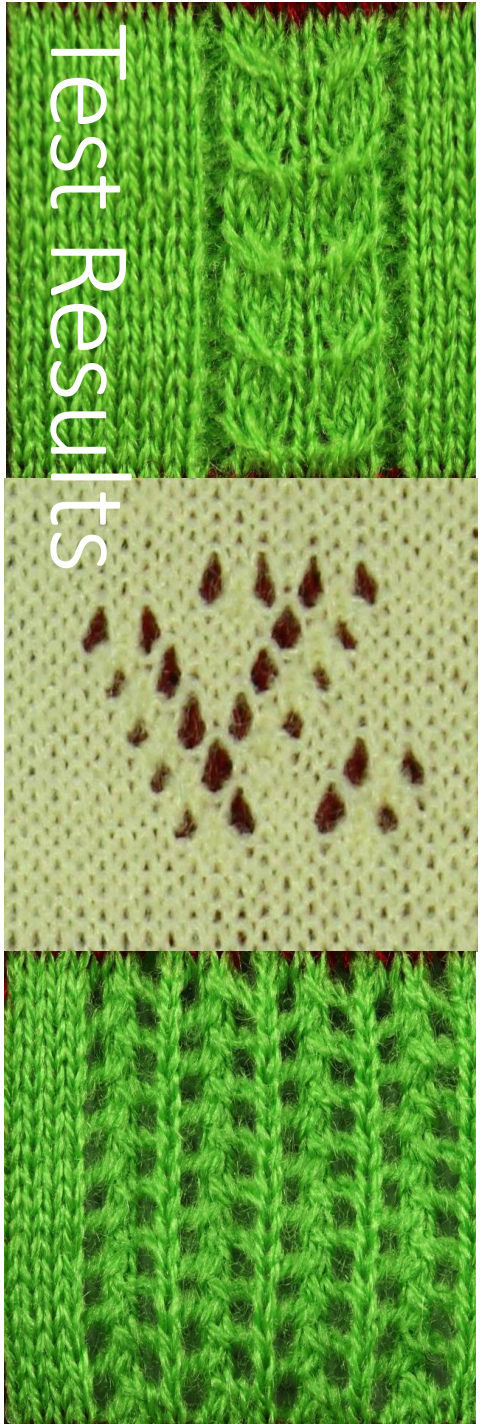
Our Result

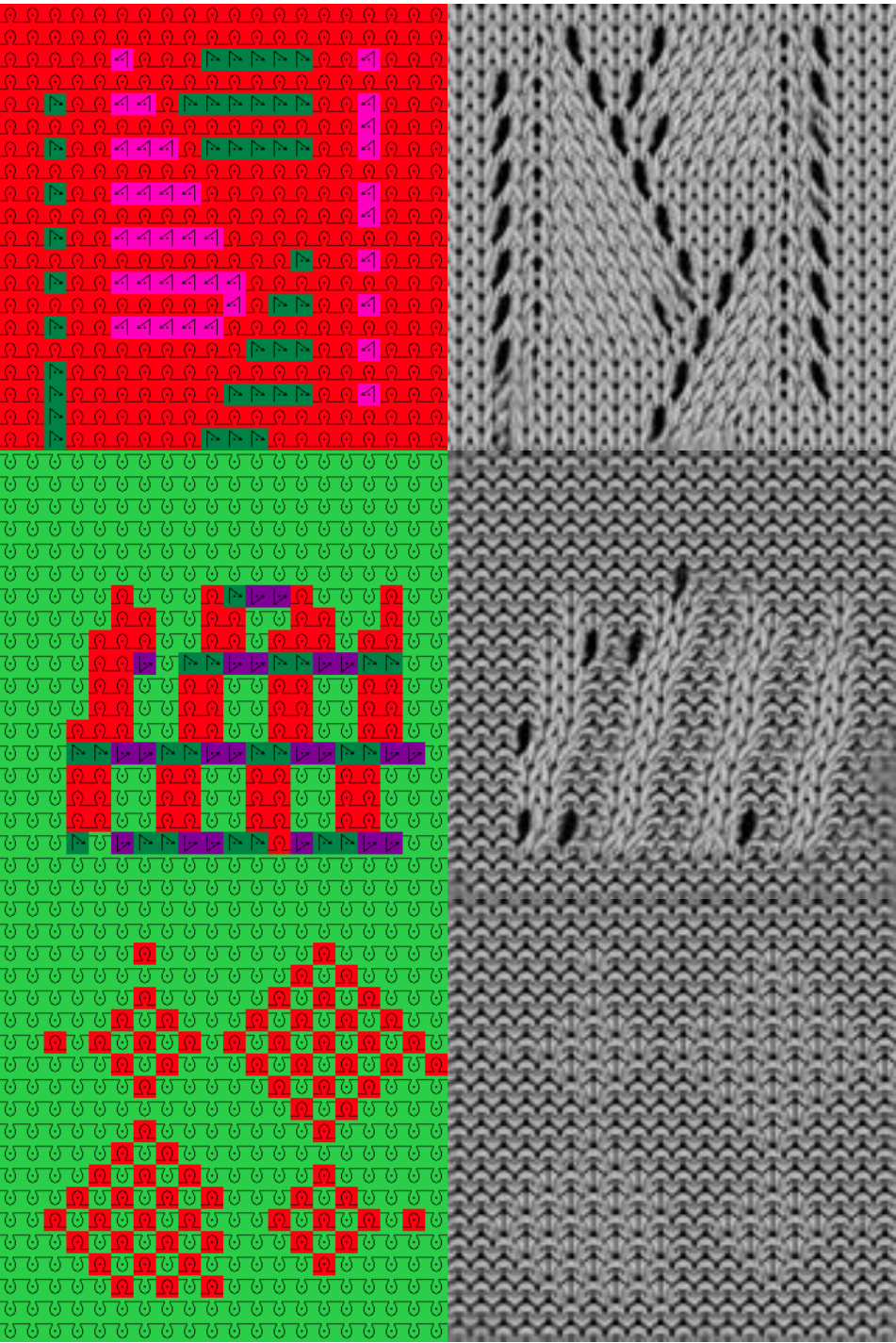


Ground Truth

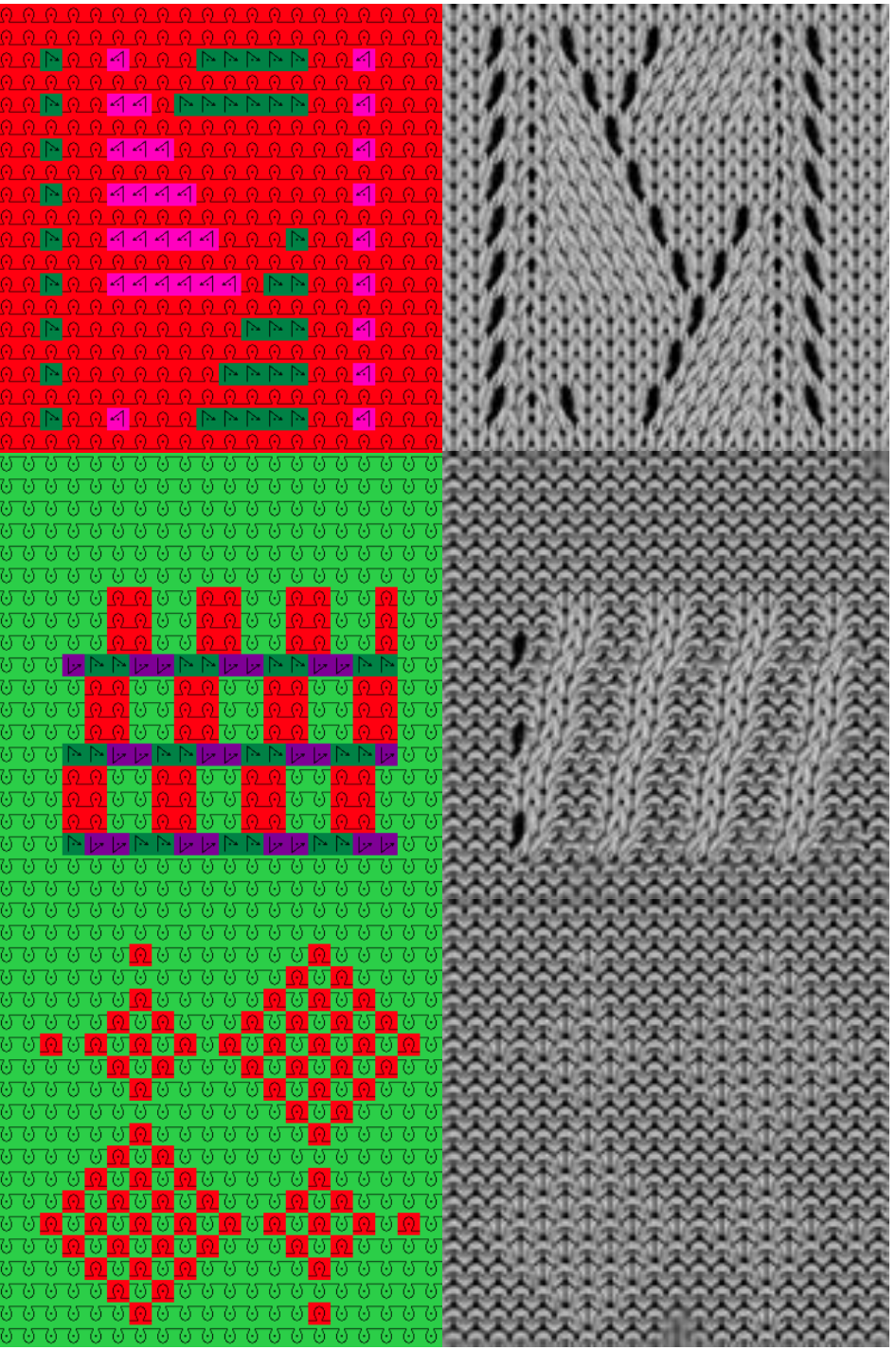


Test Results

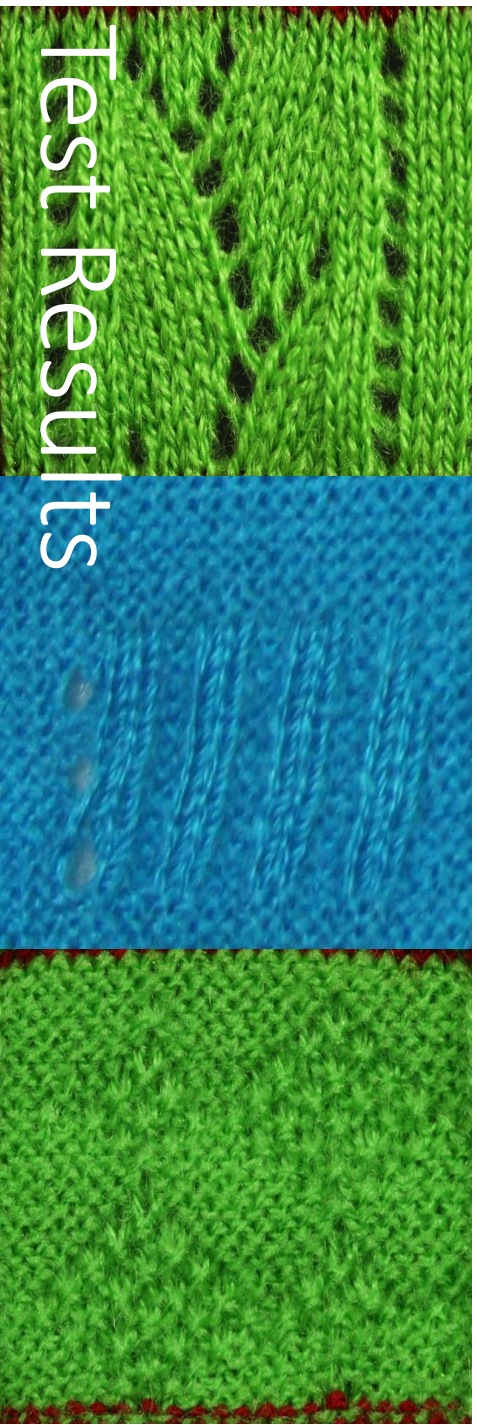




Our Result



Ground Truth



Test Results

Pacific Ballroom #137
<http://deepknitting.csail.mit.edu>



Pacific Ballroom #137
<http://deepknitting.csail.mit.edu>

Pacific Ballroom #137, <http://deepknitting.csail.mit.edu>