



Cutting-Plane Training of Non-associative Markov Network for 3D Point Cloud Segmentation



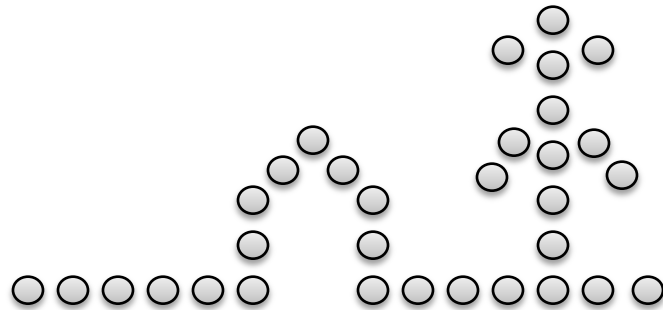
Roman Shapovalov, Alexander Velizhev
Lomonosov Moscow State University



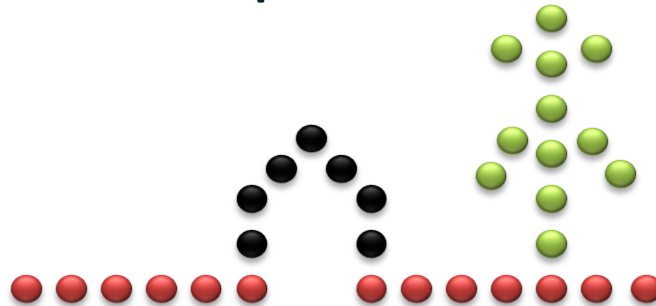
Hangzhou, May 18, 2011

Semantic segmentation of point clouds

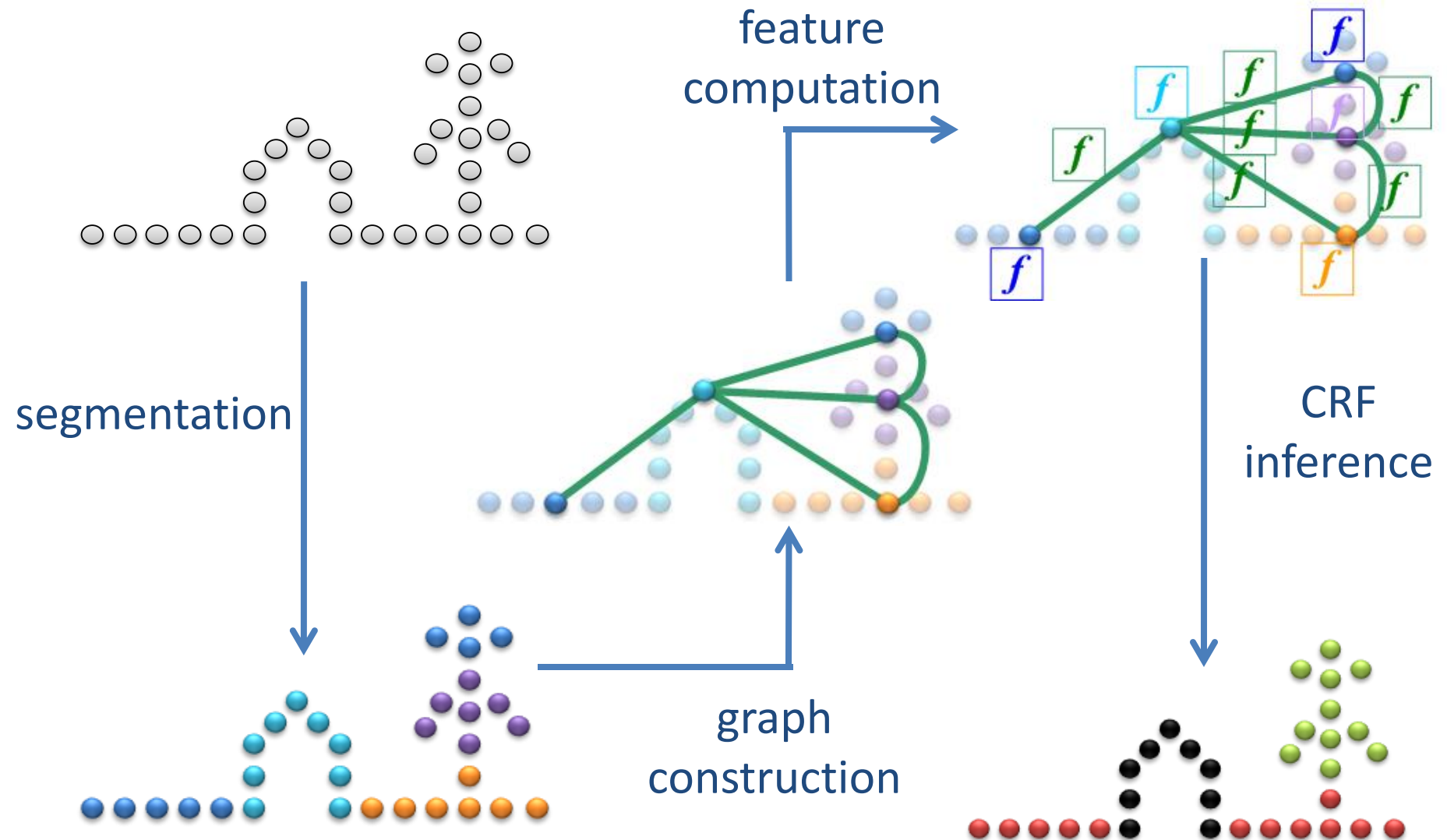
- LIDAR point cloud without color information



- Class label for each point



System workflow

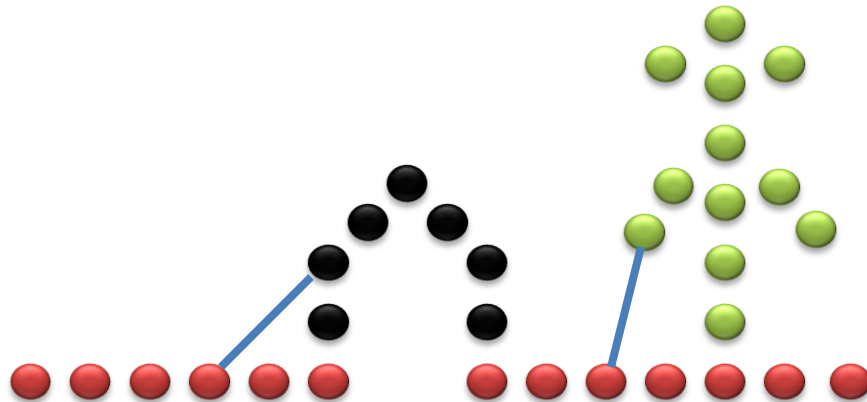


Non-associative CRF

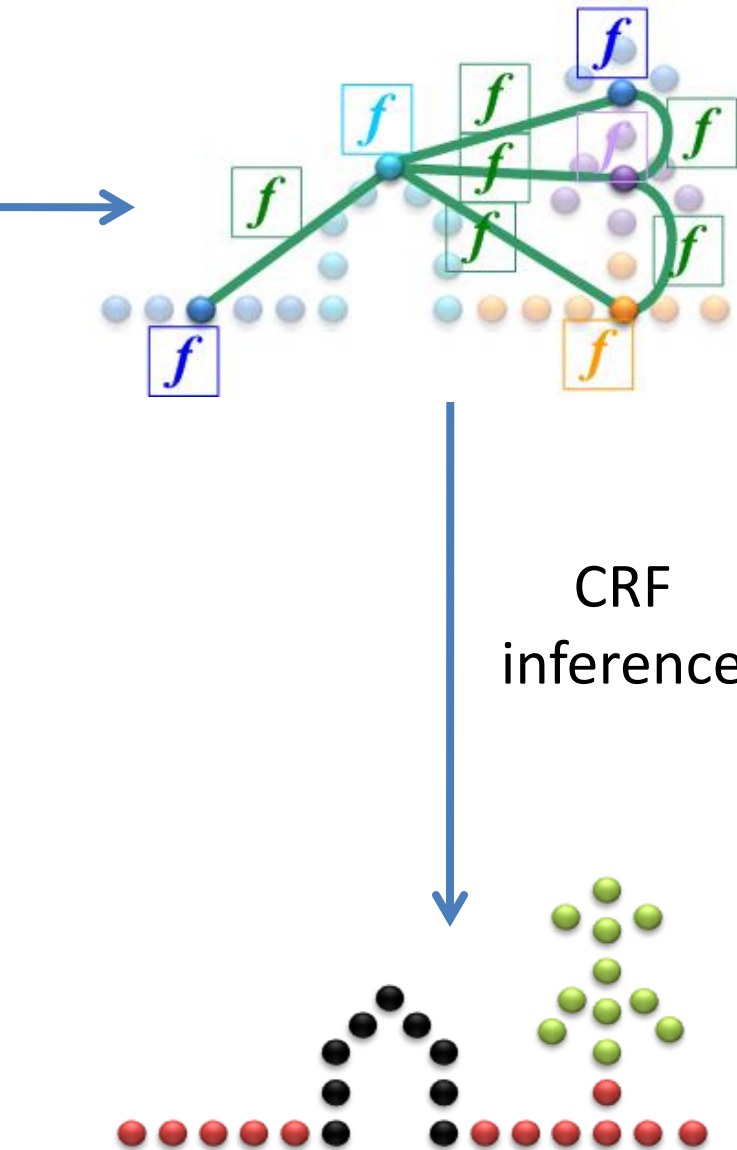
$$\sum_{i \in N} \phi(\mathbf{x}_i, y_i) + \sum_{(i,j) \in E} \phi(\mathbf{x}_{ij}, y_i, y_j) \rightarrow \max_y$$

↑ ↑ ↑
node features edge features point labels

- Associative CRF: $\phi(\mathbf{x}_{ij}, y_i, y_j) \leq \phi(\mathbf{x}_{ij}, k, k)$
- Our model: no such constraints!



CRF training



- parametric model
- parameters need to be learned!

Structured learning

[Anguelov et al., 2005; and a lot more]

- Linear model: $\phi(\mathbf{x}_i, y_i) = \mathbf{w}_{n,i}^T \mathbf{x}_i y$
 $\phi(\mathbf{x}_i, y_i, y_j) = \mathbf{w}_{e,ij}^T \mathbf{x}_{ij} y_{i,k} y_{j,l}$

- CRF negative energy: $\mathbf{w}^T \Psi(\mathbf{x}, \mathbf{y}) \rightarrow \max_y$

- Find $\mathbf{w}$ such that

$$\mathbf{w}^T \Psi(\text{train}, \text{gold}) > \mathbf{w}^T \Psi(\text{train}, \text{argmax})$$

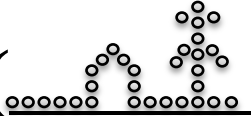
$$\mathbf{w}^T \Psi(\text{train}, \text{gold}) > \mathbf{w}^T \Psi(\text{train}, \text{argmax})$$

$$\mathbf{w}^T \Psi(\text{train}, \text{gold}) > \mathbf{w}^T \Psi(\text{train}, \text{argmax})$$

...

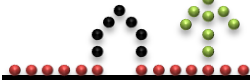
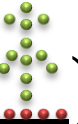
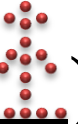
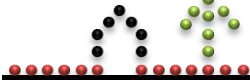
$$\mathbf{w}^T \Psi(\text{train}, \text{gold}) > \mathbf{w}^T \Psi(\text{train}, \text{argmax})$$

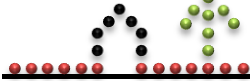
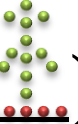
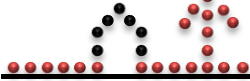
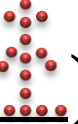
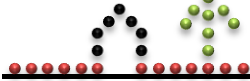
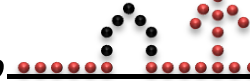
Structured loss

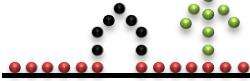
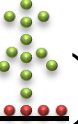
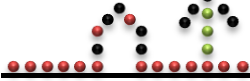
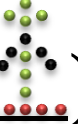
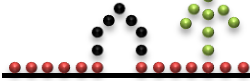
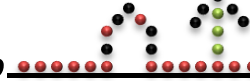
- Define $\mathbf{x} = \text{features}(\text{)$
- Define structured loss, for example:

$$\Delta(\mathbf{y}, \bar{\mathbf{y}}) = \sum_{i \in N} [y_i \neq \bar{y}_i]$$

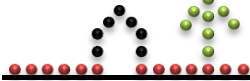
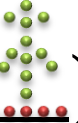
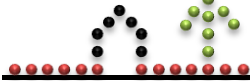
- Find $\mathbf{w}$ such that

$$\mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) > \mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) + \Delta(\text{, \text{, \text{, \text{)})$$

$$\mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) > \mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) + \Delta(\text{, \text{, \text{, \text{)})$$

$$\mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) > \mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) + \Delta(\text{, \text{, \text{, \text{)})$$

...

$$\mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) > \mathbf{w}^T \Psi(\mathbf{x}, \text{, \text{)}) + \Delta(\text{, \text{, \text{, \text{)})$$

Cutting-plane training

- A lot of constraints (K^n)

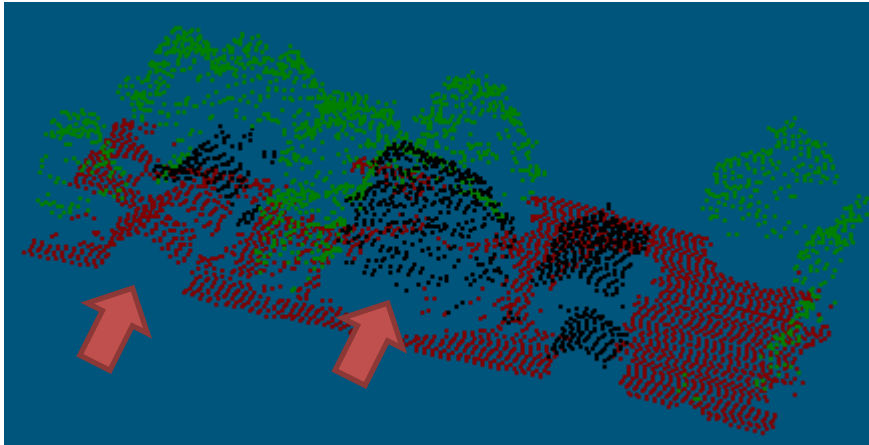
$$\mathbf{w}^T \Psi(\mathbf{x}, \mathbf{y}) > \mathbf{w}^T \Psi(\mathbf{x}, \bar{\mathbf{y}}) + \Delta(\mathbf{y}, \bar{\mathbf{y}}), \forall \bar{\mathbf{y}}$$

- Maintain a working set
- Add iteratively the most violated one:

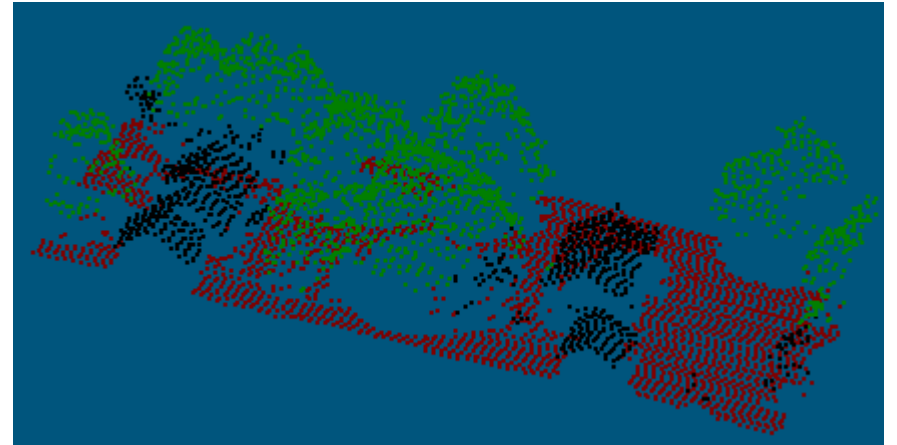
$$\bar{\mathbf{y}} = \arg \max_{\bar{\mathbf{y}}} \left[\mathbf{w}^T \Psi(\mathbf{x}, \bar{\mathbf{y}}) + \Delta(\mathbf{y}, \bar{\mathbf{y}}) \right]$$

- Polynomial complexity
- SVM^{struct} implementation [Joachims, 2009]

Results

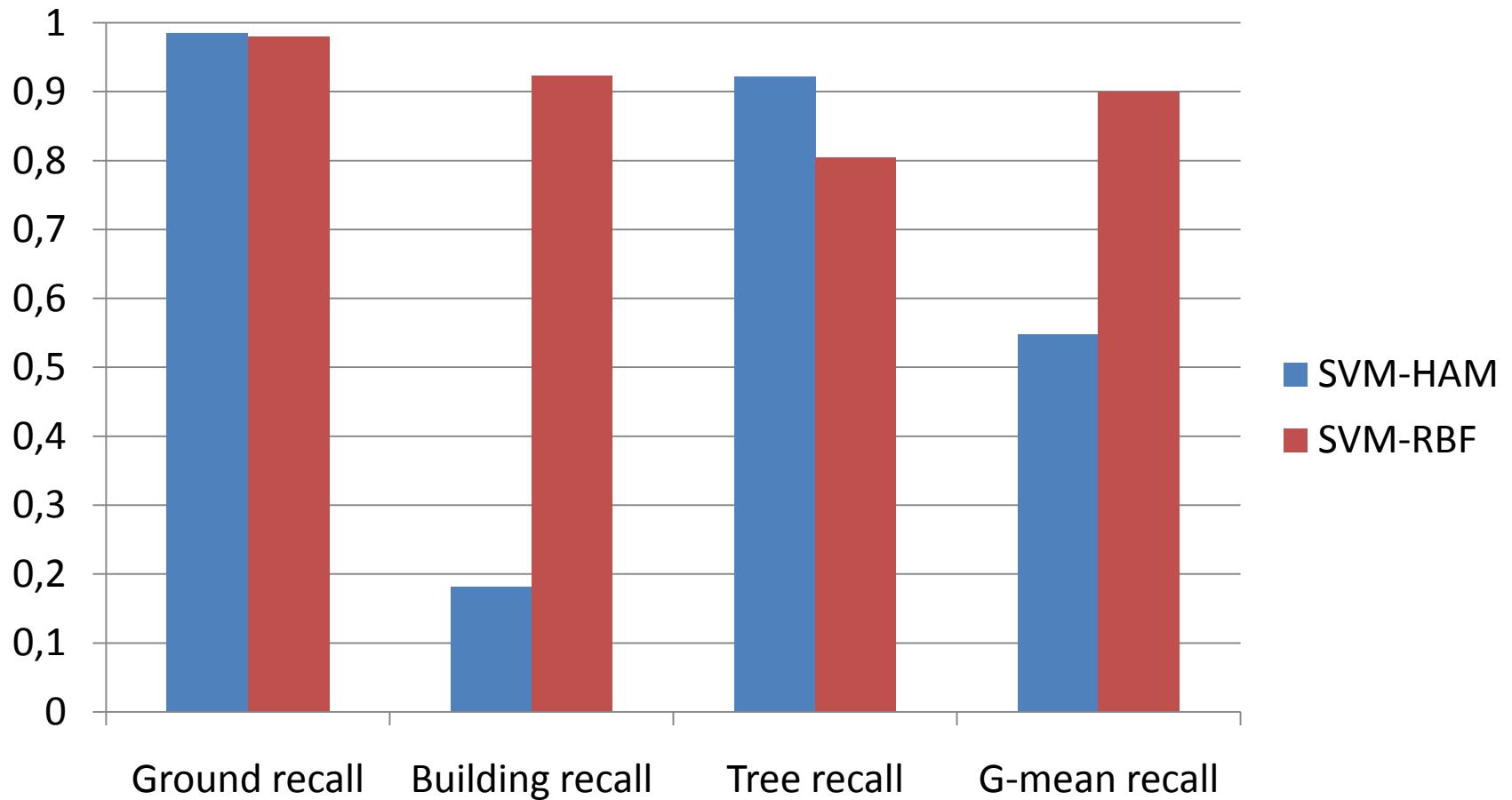


[Munoz et al., 2009]

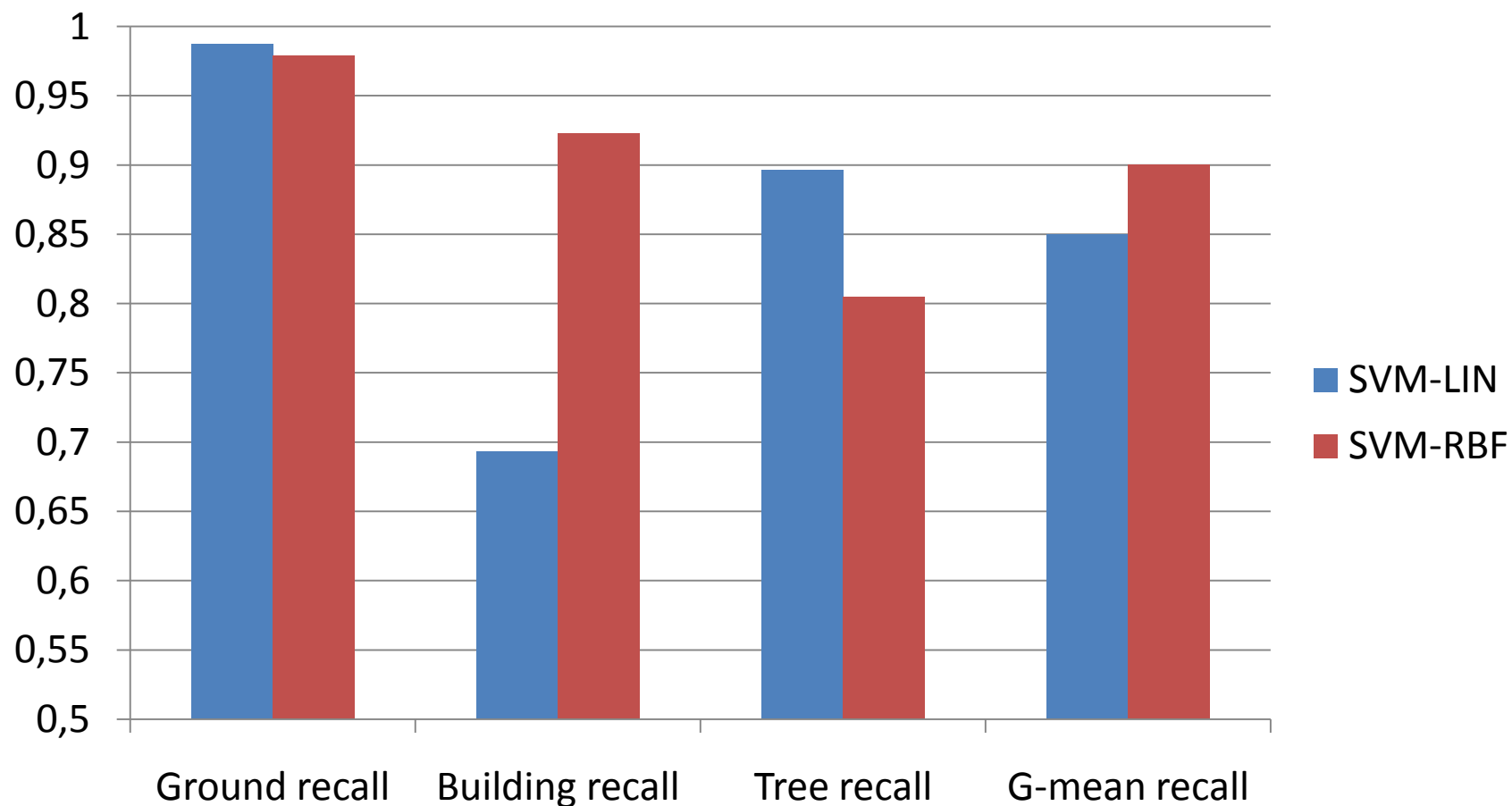


Our method

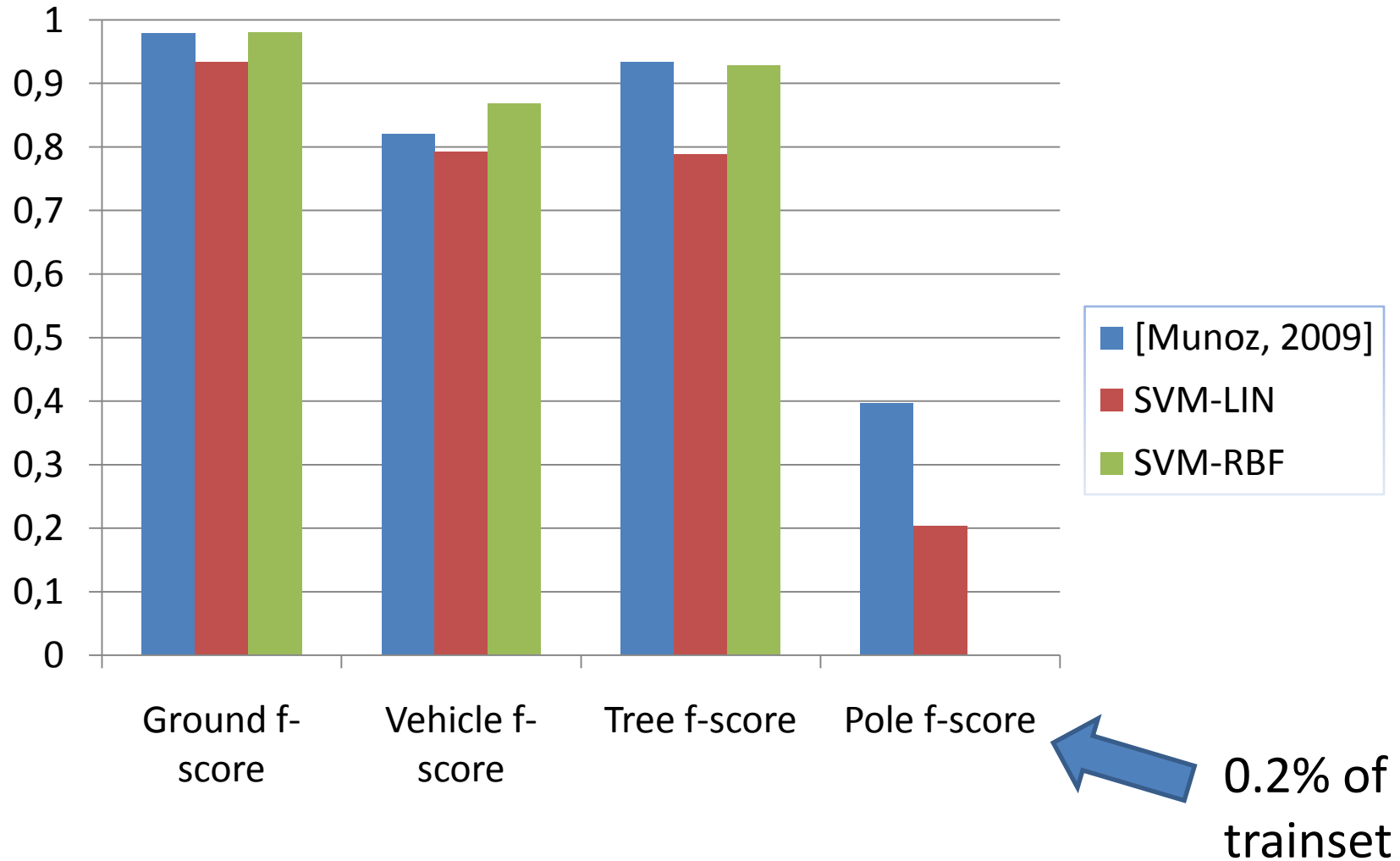
Results: balanced loss better than the Hamming one



Results: RBF better than linear



Results: fails at very small classes



Analysis

- Advantages:
 - more flexible model
 - accounts for class imbalance
 - allows kernelization
- Disadvantages:
 - really slow (esp. with kernels)
 - learns small/underrepresented classes badly