

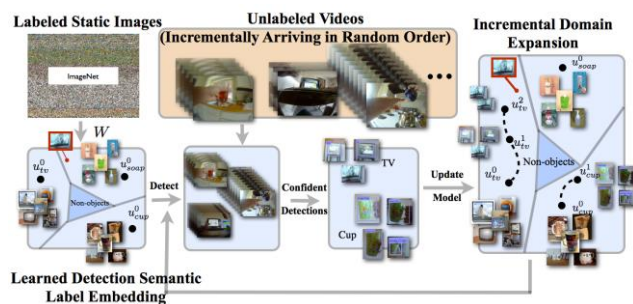
Structured Sparsity

Sung Ju Hwang

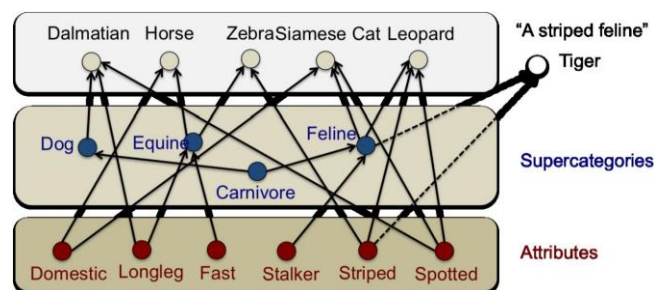
Ulsan National Institute of Science and Technology

Lecturer Introduction

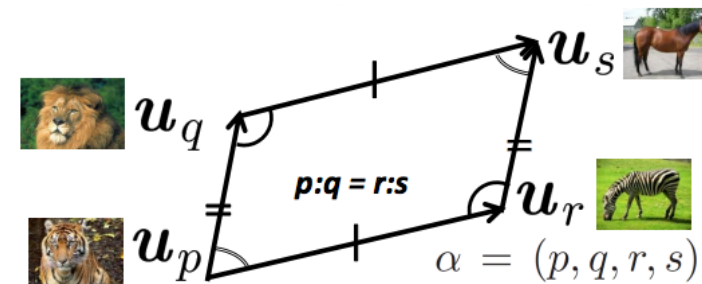
Works on the intersection between Computer Vision and Machine Learning



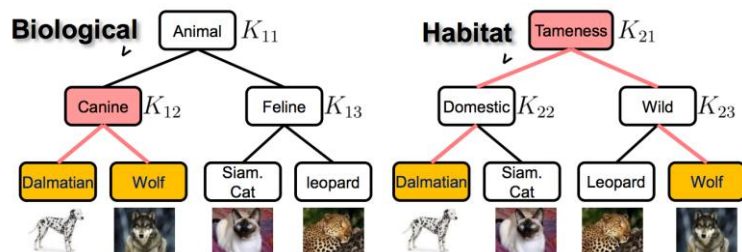
Expanding Object Detector's Horizon:
Incremental Learning Framework for
Object Detection in Videos, **CVPR 2015**



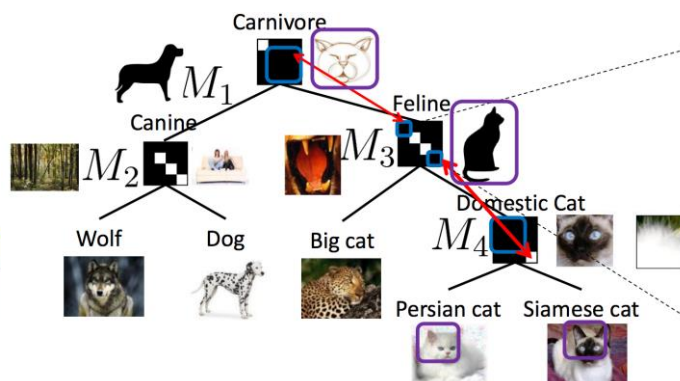
A Unified Semantic Embedding:
Relating Taxonomies with Attributes,
NIPS 2014



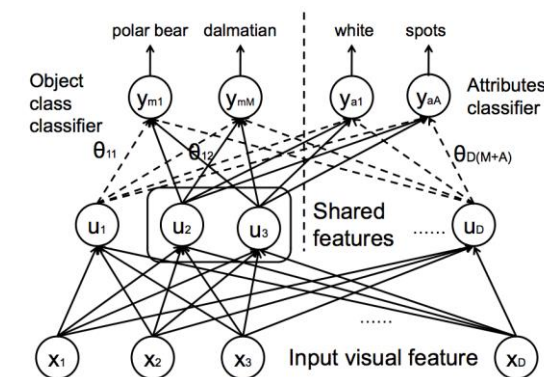
Analogy Preserving Semantic Embedding for Visual
Object Categorization, **ICML 2013**



Semantic Kernel Forests from Multiple
Taxonomies, **NIPS 2012**



Learning a Tree of Metrics with Disjoint Visual
Features, **NIPS 2011**



Sharing Features between Objects and Their
Attributes, **CVPR 2011**

Lecture Note Update

<http://sjhwang.unist.ac.kr/sparsity.pdf>

Current Research Interests

Learning Semantics

- How can we learn high-level semantic knowledge from the given visual and textual data, such that machine's understanding of the world aligns well with how we understand the world?

Interactive Learning

- How can we build a learning model such that machines learn from interacting with humans?

Lifelong Learning

- How can we learn a model that can basically learn forever, while transferring knowledge obtained at earlier stages to the learning for later ones?

Part 1: General Sparsity

- Background: Optimization
- Regularizations
- Sparsity
- Sparsity-inducing Regularization
- Example: lasso
- Numerical optimization – proximal methods
- Alternative sparse methods
- Example: Image restoration
- Example: Self-taught learning

Background: Mathematical Optimization

Selection of the values that minimize/maximize a given function

minimize $f_0(w)$
subject to $f_i(w) \leq b_i, i = 1, \dots, m$

$\min_w f_0(w)$
s.t. $f_i(w) \leq b_i, i = 1, \dots, m$

$w = (w_1, \dots, w_n)$ optimization variables

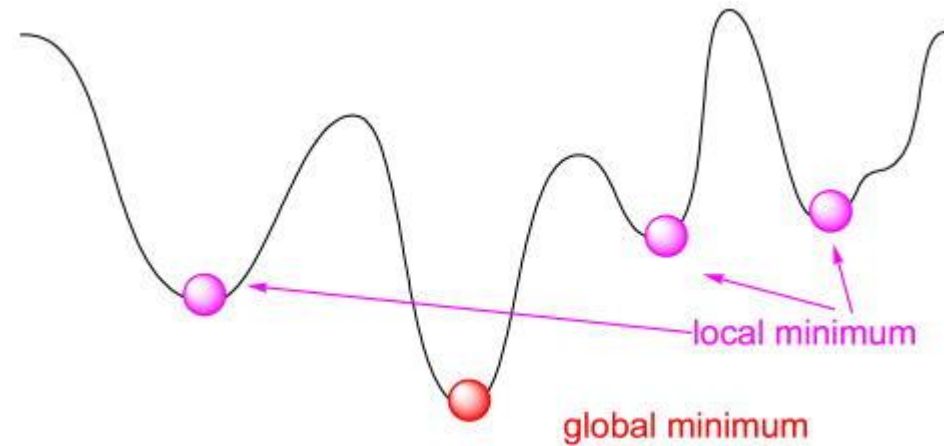
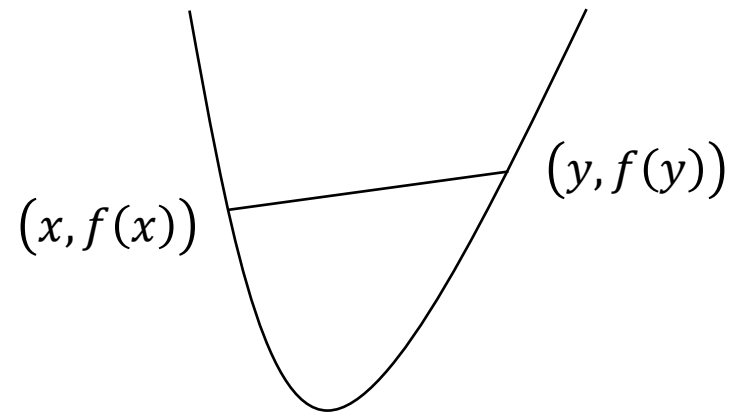
$f_0: R^n \rightarrow R$ objective function

$f_i: R^n \rightarrow R, i = 1, \dots, m$ constraints function

We want to find the optimal solution w^* that has the smallest value of f_0 among all vectors that satisfy the constraints.

Background: Convex Functions

$f: R^n \rightarrow R$ is convex if $\text{dom } f$ is a convex set and
$$f(\theta x + (1 - \theta)y) \leq \theta f(x) + (1 - \theta)f(y)$$



Has a unique global minimum.

f is concave if $-f$ is convex

Regularization

Introduction of bias to better condition the target problem with additional information

$$\min_w f(\mathbf{w}) + \lambda \Omega(\mathbf{w})$$

$$\min_w l(\overset{\text{feature}}{\mathbf{X}}, \overset{\text{label}}{\mathbf{y}}, \overset{\text{parameter, weight}}{\mathbf{w}}) + \overset{\text{regularization parameter}}{\lambda} \Omega(\mathbf{w})$$

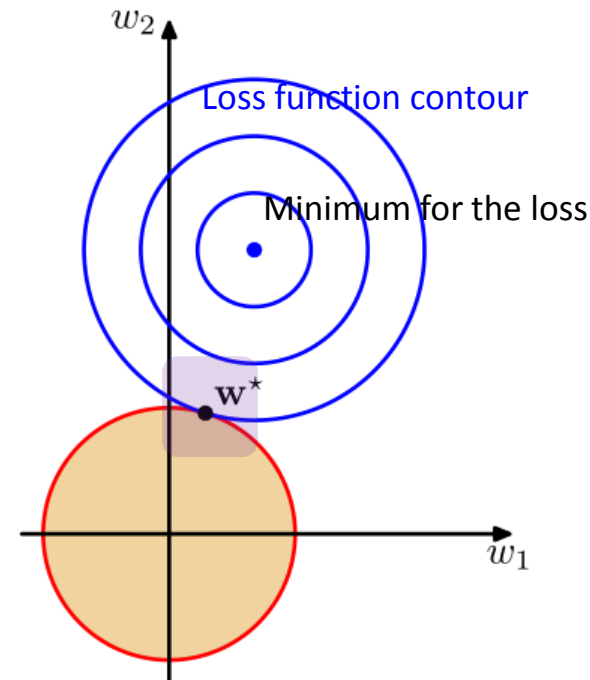
For machine learning, regularization is often used to prevent overfitting.

Regularization

L2 regularizations, based on 2 norm, is one the most common types of regularizations

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \frac{\lambda}{2} \|\mathbf{w}\|_2^2$$

$$\|\mathbf{w}\|_2 = \sqrt{\sum_j w_j^2}$$

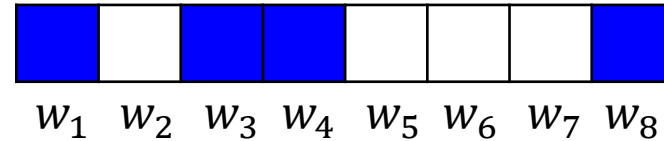


Shrinks the value of the variables while favoring similar weights among them

Sparsity

We define the L0-pseudonorm as follows:

$$\|\mathbf{w}\|_0 = \#\{i = 1, \dots, p \mid w_i \neq 0\}$$

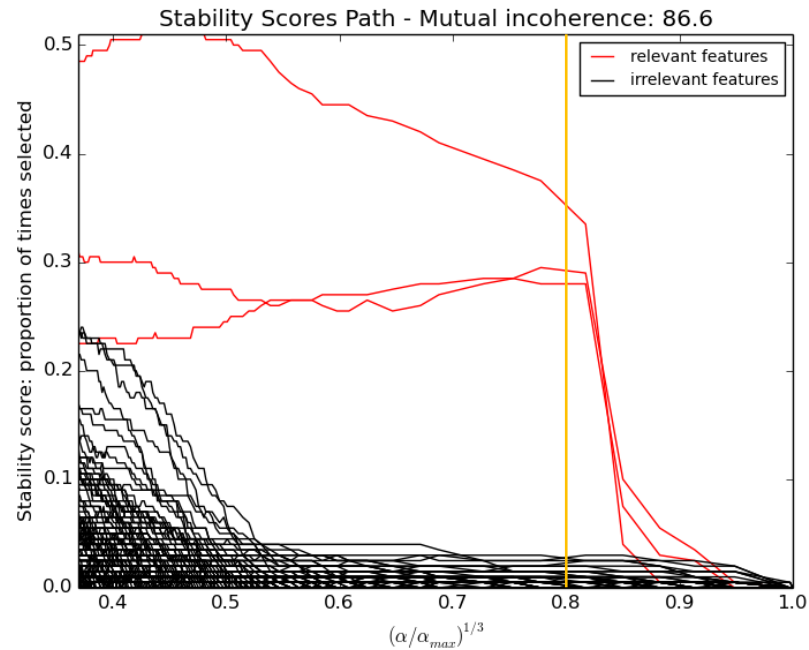


This is a measure of how many of the variables are non-zero.

L0-regularization results in learning **sparse** models, by selecting variables

Why is Sparsity Good?

Feature selection - Identifies features that are truly relevant to the target task.



Useful for high-dimensional learning and data-driven methods

Why is Sparsity Good?

Better Interpretability – The learned model can be better explained in terms of selected non-zero entries.

Category	Ground-truth attributes	Supercategory + learned attributes
 Otter	An animal that swims, fish, water, new world, small, flippers, furry, black, brown, tail, ...	A musteline mammal that is quadrapedal, flippers, furry, ocean
 Skunk	An animal that is smelly, black, stripes, white, tail, furry, ground, quadrapedal, new world, walks, ...	A musteline mammal that has stripes
 Deer	An animal that is brown, fast, horns, grazer, forest, quadrapedal, vegetation, timid, hooves, walks, ...	A deer that has spots, nestspot, longneck, yellow, hooves
 Moose	An animal that has horns, brown, big, quadrapedal, new world, vegetation, grazer, hooves, strong, ground,...	A deer that is arctic, stripes , black
Equine	N/A	An odd-toed ungulate, that is lean and active
Primate	N/A	An animal, that has hands and bipedal

Why is Sparsity Good?

Model compression – With most of the parameters set to zero, sparsity can greatly reduce the memory and computational requirements.

$$\begin{pmatrix} 1.0 & 0 & 5.0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 3.0 & 0 & 0 & 0 & 0 & 11.0 & 0 \\ 0 & 0 & 0 & 0 & 9.0 & 0 & 0 & 0 \\ 0 & 0 & 6.0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 7.0 & 0 & 0 & 0 & 0 \\ 2.0 & 0 & 0 & 0 & 0 & 10.0 & 0 & 0 \\ 0 & 0 & 0 & 8.0 & 0 & 0 & 0 & 0 \\ 0 & 4.0 & 0 & 0 & 0 & 0 & 0 & 12.0 \end{pmatrix}$$

8x8 matrix – requires 512 bytes to store in double precision

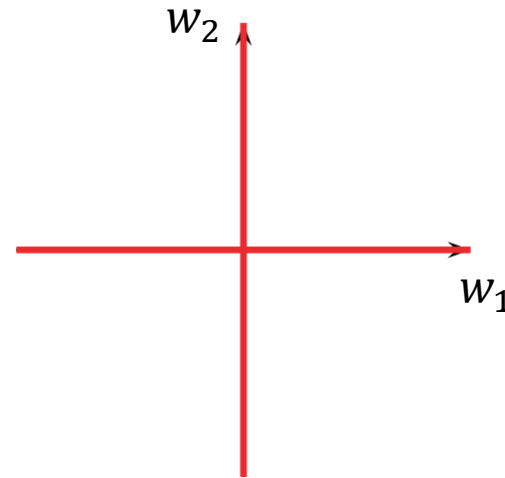
12 nonzero entries – requires 96 bytes + additional memory to store indices.

L0-regularization

Directly solving for L0-regularization is difficult as it should try all possible subsets of variables.

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \|\mathbf{w}\|_0$$

$$\|\mathbf{w}\|_0 = \#\{i = 1, \dots, p \mid w_i \neq 0\}$$



Are there more efficient ways to obtain sparse models?

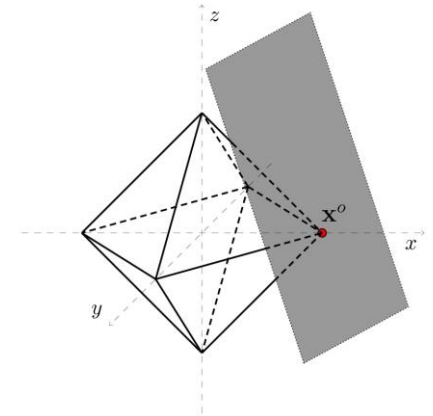
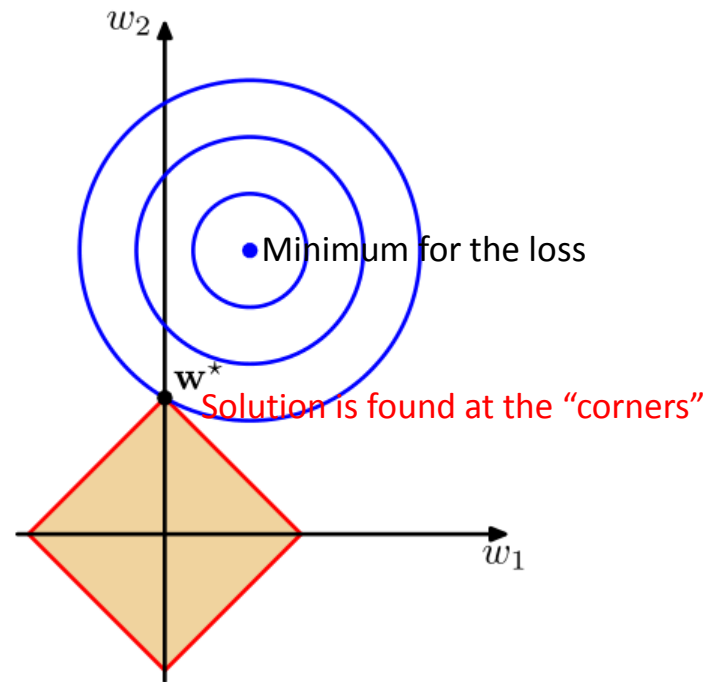
L1 regularization

Convex relaxation of L0 regularization.

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \|\mathbf{w}\|_1$$

$$\|\mathbf{w}\|_1 = \sum_j |w_j|$$

$$\|\mathbf{w}\|_1 \leq \lambda$$



As the dimensionality of $\mathbf{w}$ increases, the norm ball will have increasingly more number of corners.

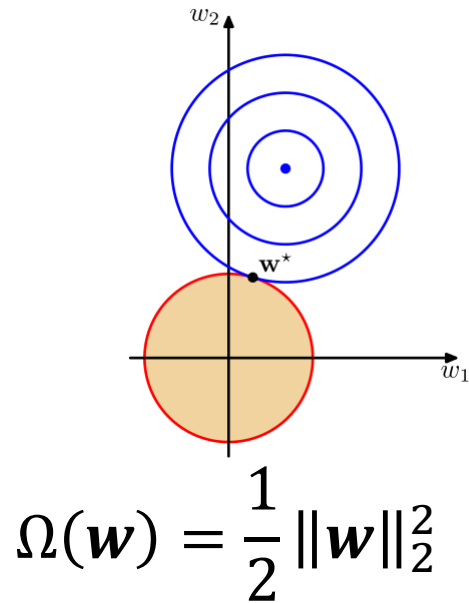
Results in parameter selection

L2- vs L1-regularization

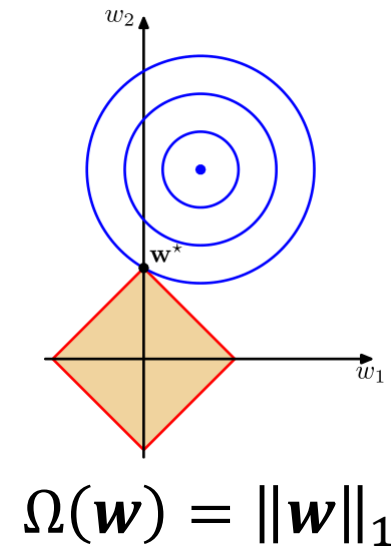
L2 regularization promotes **grouping** – results in equal weights for correlated features

L1 regularization promotes **sparsity** – selects few informative features

L2-regularization

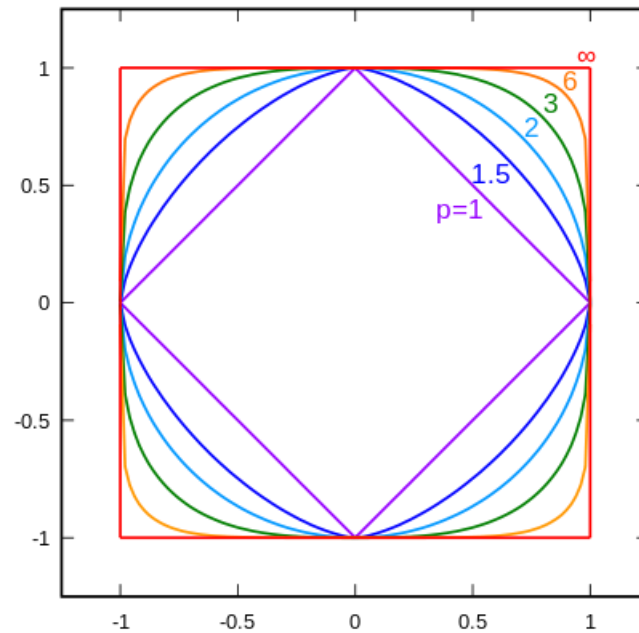


L1-regularization



Lp-Norms

General case $\|\mathbf{w}\|_p$ where p can be any non-negative number



Lp norms with $p < 2$ promotes sparsity

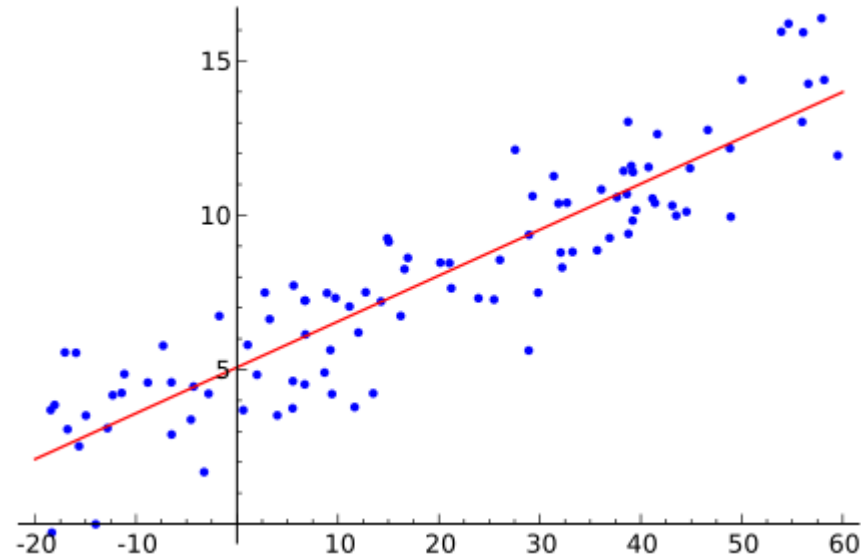
Lp norms with $p > 2$ promotes grouping

Example: Linear Regression

Fit a linear model $\mathbf{w}$, that minimizes the residual between the observed and predicted values

$$\min_{\mathbf{w}} \frac{1}{N} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2$$

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$



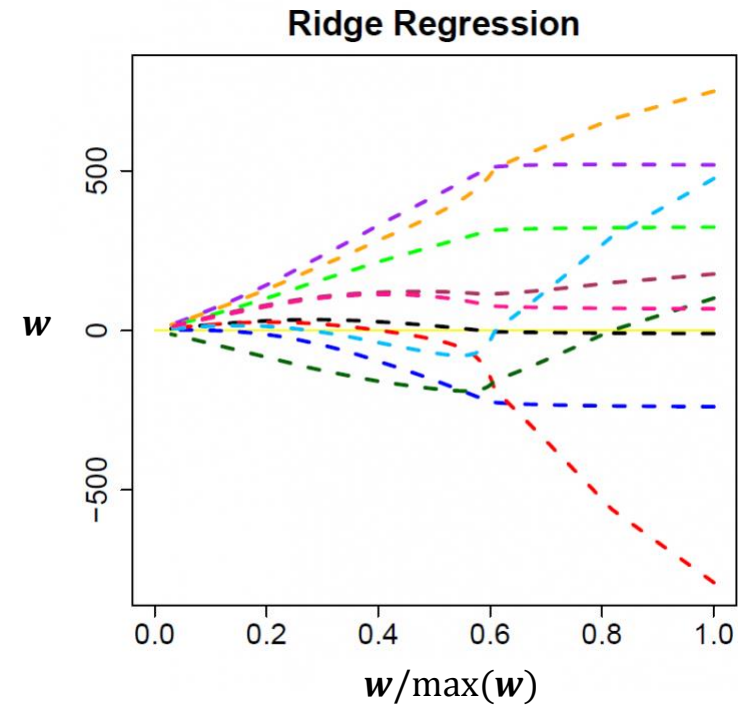
Ridge Regression

Linear regression with L2-regularization

$$\min_{\mathbf{w}} \frac{1}{N} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 + \lambda \Omega(\mathbf{w})$$

$$\Omega(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|_2^2$$

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^T \mathbf{y}$$



Shrinks all variables to zero – reduces variance while introducing bias.

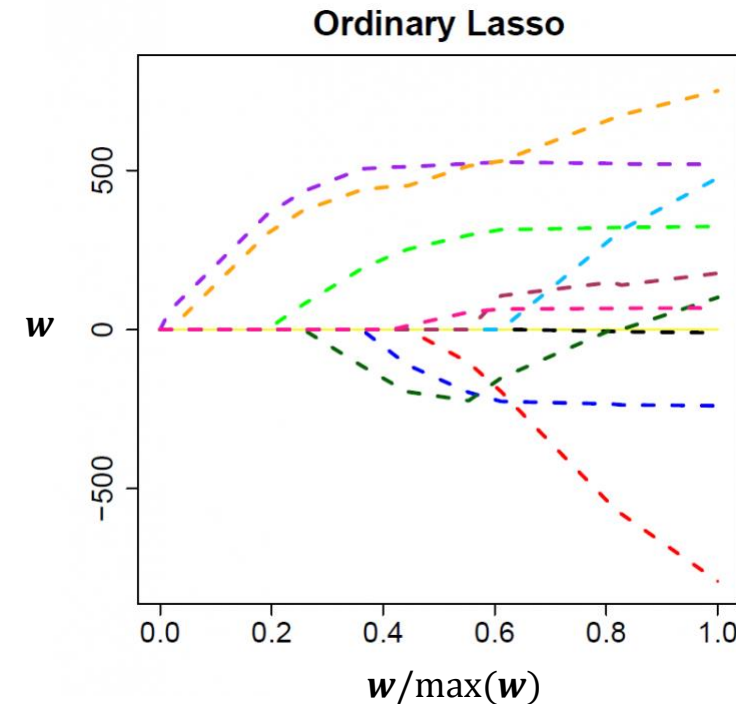
Lasso (L1-regularized Linear Regression)

Shorthand for Least Absolute Shrinkage and Selection Operator

Linear regression with L1-regularization

$$\min_w \frac{1}{N} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 + \lambda \Omega(\mathbf{w})$$

$$\Omega(\mathbf{w}) = \|\mathbf{w}\|_1$$



Not all variables become zero at the same time – results in variable selection

L1-Regularization for General Loss

L1-regularization can be coupled with various types of loss to obtain sparse models
e.g.) logistic regression, SVM

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \Omega(\mathbf{w})$$

$$l(\mathbf{X}, \mathbf{y}, \mathbf{w}) = \frac{1}{N} \sum_{i=1}^N [(1 - y_i) \mathbf{w}^T \mathbf{x}_i + \log(1 + \exp(-\mathbf{w}^T \mathbf{x}_i))]$$

$$\Omega(\mathbf{w}) = \|\mathbf{w}\|_1 \quad \|\mathbf{w}\|_1 \leq \lambda$$

How can we then **solve** for such l1-regularized objectives?

Gradient Descent

Determine the descent direction as the gradient at the point, determine the step size t , and then move to that direction.

given a starting point $x \in \text{dom } f$.

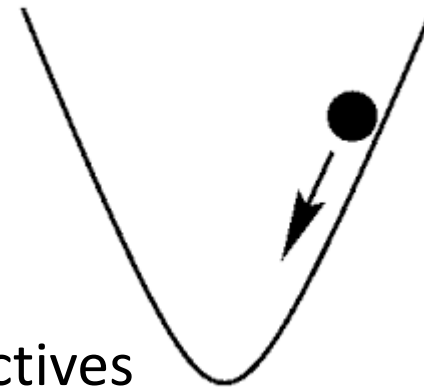
repeat

1. $\Delta x := -\nabla f(x)$.
2. *Line search.* Choose step size t via exact or backtracking line search.
3. *Update.* $x := x + t\Delta x$.

until stopping criterion is satisfied.

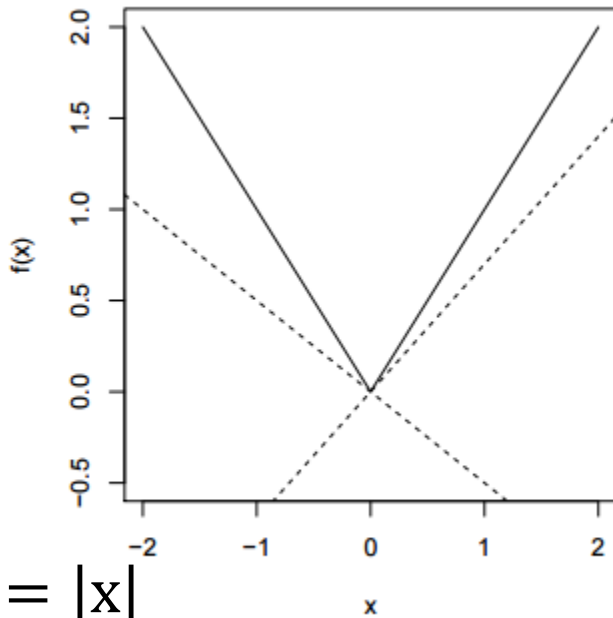
Repeat this process until a stopping criterion is met.

Cannot be applied to optimization of non-smooth objectives



Subgradient Method

Compute a set of gradient, defining the gradient on non-smooth points



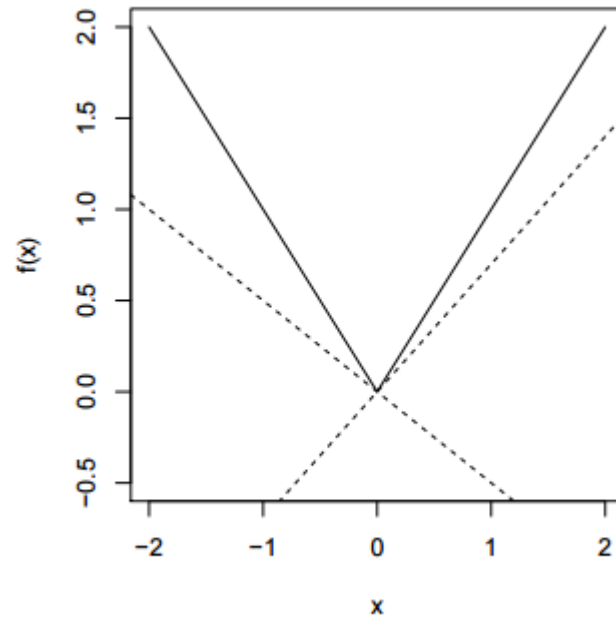
$$f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = |x|$$

For $x \neq 0$, the subgradient $g = \text{sign}(x)$

For $x = 0$, subgradient g is any element of $[-1,1]$.

Subgradient Method

A vector g is a subdifferential of f at x , if $f(y) = f(x) + d^T(y - x), \forall y$



Suffers from slow convergence, and solutions obtained are usually non-sparse

Proximal Gradient

Specifically tailored for regularized optimization problems where $g(\mathbf{w})$ is differentiable but $h(\mathbf{w})$ is a general convex function.

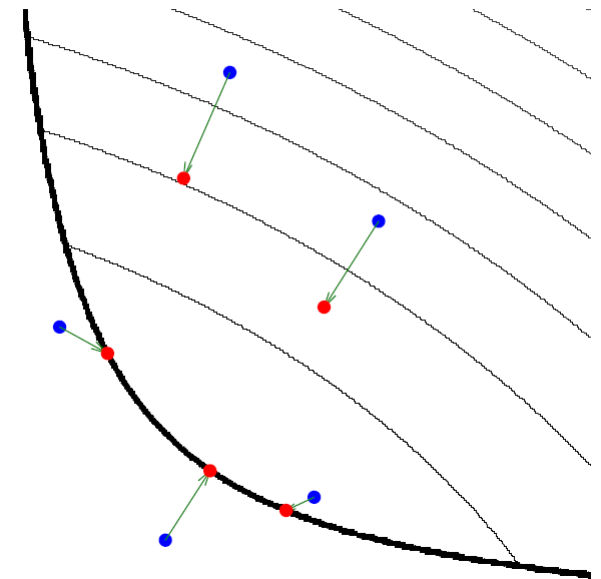
$$\min_{\mathbf{w}} \overset{\text{smooth}}{g(\mathbf{w})} + \overset{\text{smooth or nonsmooth}}{h(\mathbf{w})}$$

$$\mathbf{w}_{t+1} = \text{prox}(\mathbf{w}_t - \alpha \nabla g(\mathbf{x}_t))$$

Solution obtained by taking a gradient step only on $g(\mathbf{w})$

$$\text{prox}(\mathbf{w}) = \underset{\mathbf{w}}{\text{argmin}} \frac{1}{2} \|\mathbf{w} - \mathbf{w}^*\|_2^2 + \alpha h(\mathbf{w})$$

Euclidean Projection



The proximal operator, $\text{prox}(\mathbf{w})$ should be efficient – closed form solution preferred

Proximal Operator for L1-norm regularization

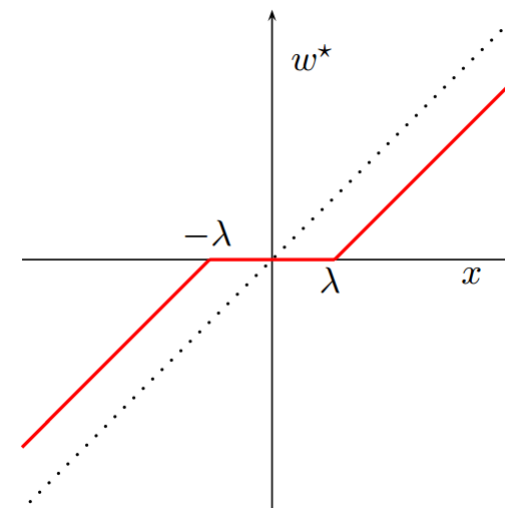
Iterative soft-thresholding operator – reduce the absolute value of each component of $\mathbf{w}$ by λ , and if the resulting value is below zero, set it to zero.

$$\min_{\mathbf{w}} g(\mathbf{w}) + h(\mathbf{w})$$

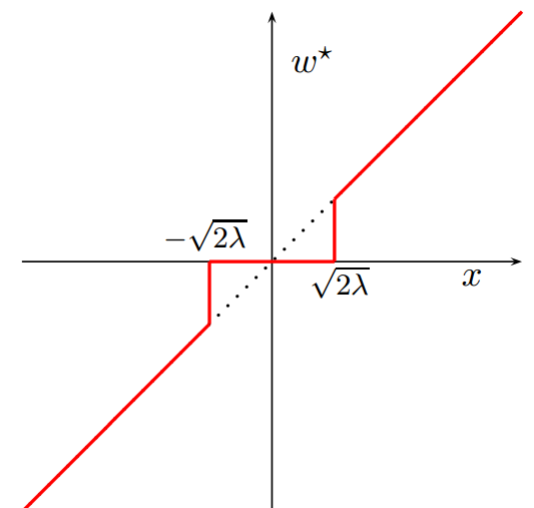
$$h(\mathbf{w}) = \lambda \|\mathbf{w}\|_1$$

$$\mathbf{w}^* = \mathbf{w}_t - \alpha \nabla g(\mathbf{x}_t)$$

$$\text{prox}(\mathbf{w}) = \text{sign}(\mathbf{w}^*) [|\mathbf{w}^*| - \lambda]_+$$



(a) soft-thresholding operator,



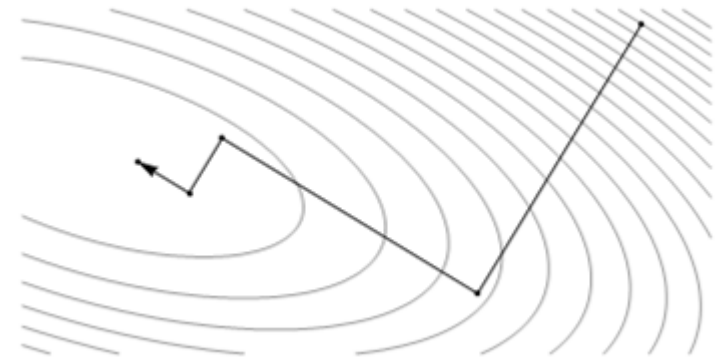
(b) hard-thresholding operator

Coordinate Descent

Optimize one variable at a time while fixing all others.

$$\min_{\mathbf{w}_j} \nabla_j f(\mathbf{w}^t)(\mathbf{w}_j - \mathbf{w}_j^t) + \frac{1}{2} \nabla_{jj}^2 f(\mathbf{w}^t)(\mathbf{w}_j - \mathbf{w}_j^t) + \lambda |\mathbf{w}_j|$$

$$\mathbf{w}_j^* = \text{prox}_{\frac{\lambda}{\nabla_{jj}^2 f} |\cdot|} \left(\mathbf{w}_j^t - \frac{\nabla_j f(\mathbf{w}_j^t)}{\nabla_{jj}^2 f} \right)$$



$\mathbf{w}_j^*$ can be obtained by solving for the loss with coordinate j and soft-thresholding the solution.

Stochastic Subgradient Descent

Subgradient method using noisy unbiased subgradients for scalable optimization

step size at iteration t

$$\mathbf{w}_{t+1} = \mathbf{w}_t - a_t \widetilde{\mathbf{g}}_t$$

gradient direction at iteration t

$$E(\widetilde{\mathbf{g}}_t | \mathbf{w}_t) = \mathbf{g}_t \in \partial f(\mathbf{w}_t)$$

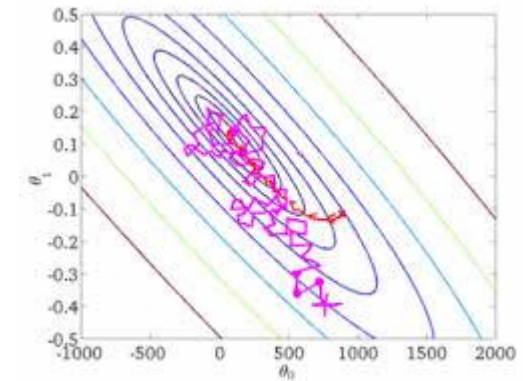
A random vector is a noisy unbiased subgradient for $f: R \rightarrow R$ at $\mathbf{x}$, if for all $\mathbf{z}$

$$f(\mathbf{z}) \geq f(\mathbf{w}) + (E \widetilde{\mathbf{g}})^T (\mathbf{z} - \mathbf{w})$$

Define the optimal value as $f^* = \min\{f(\mathbf{w}_1), \dots, f(\mathbf{w}_t)\}$

Scales and works well for ML problems

However, SGD often does not generate sparse solutions for l1-regularized objectives



Regularized Dual Averaging Method

Consider all past subgradients of the loss and the whole regularization term when solving for a regularized objective

$$\mathbf{w}_{t+1} = \operatorname{argmin}_{\mathbf{w}} \left\{ \langle \overline{\mathbf{g}}_t, \mathbf{w} \rangle + \Omega(\mathbf{w}) + \frac{\beta_t}{t} h(\mathbf{w}) \right\}$$

$\overline{\mathbf{g}}_t = \frac{t-1}{t} \overline{\mathbf{g}}_{t-1} + \frac{1}{t} \mathbf{g}_t$
Average of all gradients over t iterations

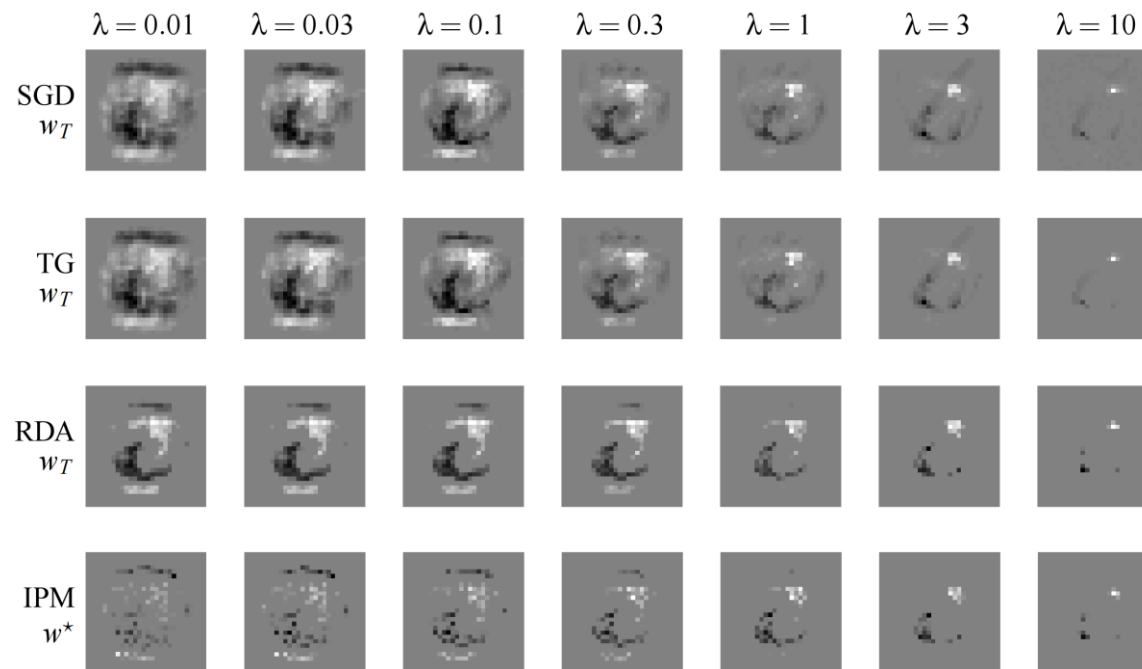
Positive sequence - $\gamma\sqrt{t}$
auxiliary Strongly convex function

$$w_{t+1}^i = \begin{cases} 0, & |g_t^i| \leq \lambda \\ -\frac{\sqrt{t}}{\gamma} (g_t^i - \lambda \operatorname{sign}(g_t^i)), & \text{otherwise} \end{cases}$$

For l1-regularization RDA has a closed form solution

Regularized Dual Averaging Method

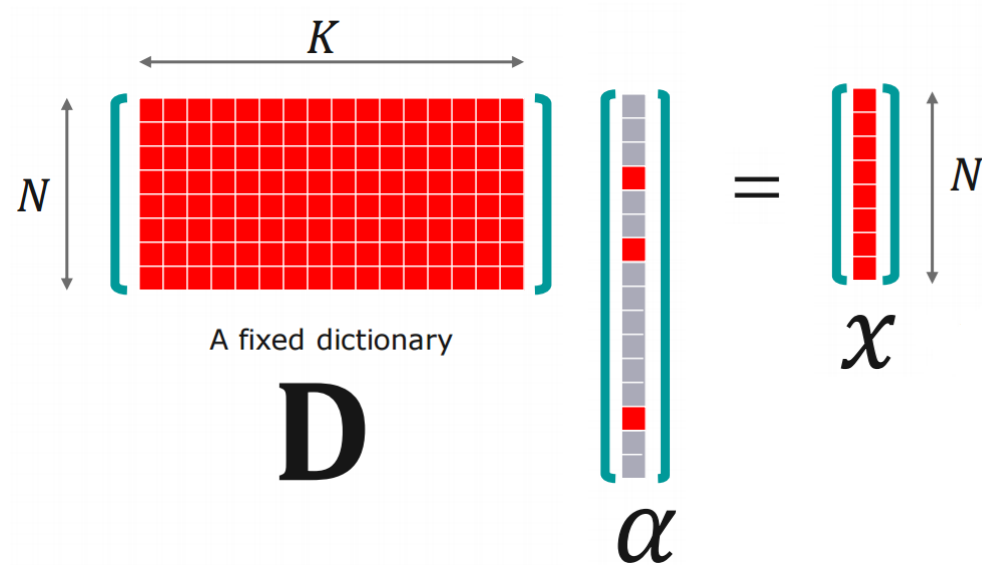
RDA obtains sparser solutions compared to those obtained by SGD and proximal gradient



RDA has a convergence rate of $1/\sqrt{T}$ for general convex regularizers.

Sparse Coding

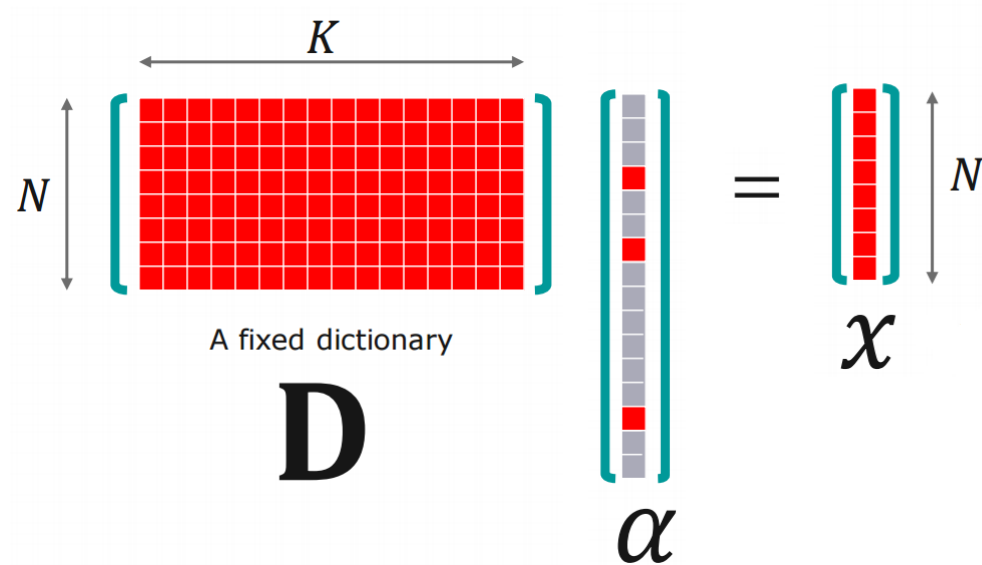
Find sparse decomposition of each data instance x_i , using some known dictionary atoms



$$\min_{\alpha} \frac{1}{N} \sum_{i=1}^N \| \overset{\text{data instance}}{x_i} - \underset{\text{dictionary}}{D} \overset{\text{sparse code}}{\alpha_i} \|_2^2 + \lambda \underset{\text{Sparsity-inducing regularization}}{\Omega(\alpha)}$$

Sparse Coding and Dictionary Learning

Find both a sparse decomposition α_i of each data instance x_i , and the bases D



$$\min_{\alpha, D} \frac{1}{N} \sum_{i=1}^N \| \overset{\text{data instance}}{x_i} - \overset{\text{dictionary}}{D} \overset{\text{sparse code}}{\alpha_i} \|_2^2 + \lambda \overset{\text{Sparsity-inducing regularization}}{\Omega(\alpha)}$$

Can use alternating optimization to solve for each variable at a time.

Alternative Sparse Methods

Use greedy algorithms to directly solve for l0-norm

$$\min_{\mathbf{w}} \frac{1}{N} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 \quad s.t. \|\mathbf{w}\|_0 \leq \lambda$$

$$\min_{\boldsymbol{\alpha}} \|\mathbf{x} - \mathbf{D}\boldsymbol{\alpha}\|_2^2 \quad s.t. \|\boldsymbol{\alpha}\|_0 \leq \lambda$$

Enforce sparsity by setting the desired number of nonzero variables.

- Matching Pursuit, Orthogonal Matching Pursuit

Matching Pursuit

At each step, select the column that is most correlated with the input

$$\min_{\alpha} \|\mathbf{x} - \mathbf{D}\alpha\|_2^2 \quad \text{s.t. } \|\alpha\|_0 \leq \lambda$$

$$\alpha \leftarrow 0$$

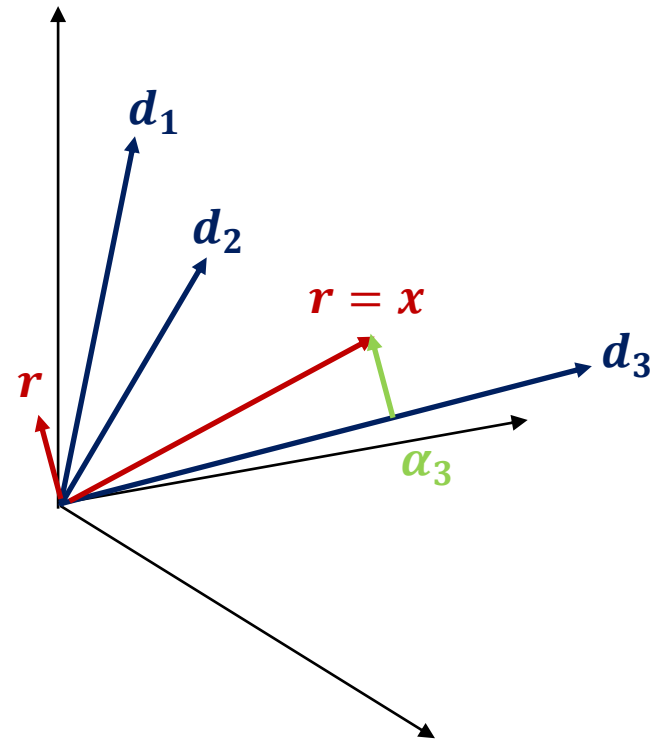
$$\mathbf{r} \leftarrow \mathbf{x}$$

while $\|\alpha\|_0 \leq \lambda$

$$i \leftarrow \operatorname{argmax}_{\{i=1,\dots,p\}} |\mathbf{d}_i^T \mathbf{r}|$$

$$\alpha_i \leftarrow \alpha_i + \mathbf{d}_i^T \mathbf{r}$$

$$\mathbf{r} \leftarrow \mathbf{r} - (\mathbf{d}_i^T \mathbf{r}) \mathbf{d}_i$$



Orthogonal Matching Pursuit

At each step, select the column that helps reduce the objective the most

$$\min_{\alpha} \|\mathbf{x} - \mathbf{D}\alpha\|_2^2 \quad s.t. \quad \|\alpha\|_0 \leq \lambda$$

$$\Gamma = \emptyset$$

for $i = 1, \dots, \lambda$ do

$$i = \operatorname{argmin}_{i \in \Gamma^c} \left\{ \min_{\alpha'} \|\mathbf{x} - \mathbf{D}_{\Gamma \cup \{i\}} \alpha'\|_2^2 \right\}$$

$$\Gamma \leftarrow \Gamma \cup \{i\}$$

$$\mathbf{r} \leftarrow \left(\mathbf{I} - \mathbf{D}_{\Gamma} (\mathbf{D}_{\Gamma}^T \mathbf{D}_{\Gamma})^{-1} \mathbf{D}_{\Gamma}^T \right) \mathbf{x}$$

$$\alpha_{\Gamma} \leftarrow (\mathbf{D}_{\Gamma}^T \mathbf{D}_{\Gamma})^{-1} \mathbf{D}_{\Gamma}^T \mathbf{x}$$

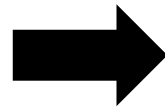
All coefficients extracted so far are updated at each step

Image Restoration

Eliminate noise from the original image



Original

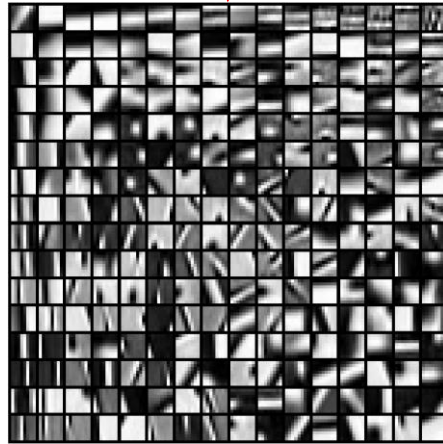
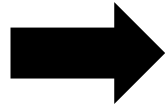
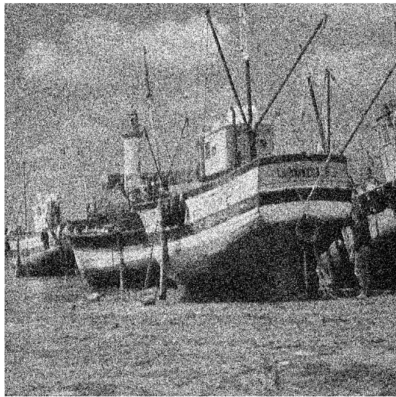


Restored

Image Restoration with Sparse Coding

Reconstruct the image with sparse combination of learned bases

$$\min_{\alpha} \frac{1}{N} \sum_{i=1}^N \|x_i - \mathbf{D}\alpha_i\|_2^2 + \lambda \|\alpha\|_1$$



Bases

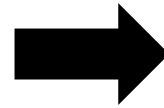
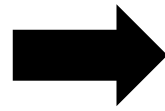


Image Inpainting

Works with even larger level of noise



Original



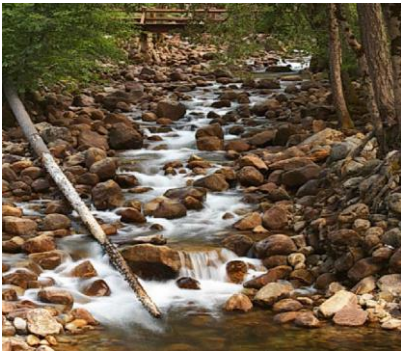
Restored

Self-taught Learning

Transfer learning when source domain is **unlabeled** and only remotely related to the target task – based on the intuition that human learning is largely unsupervised.

Source

$$\{x_u^{(i)}\}_{i=1}^k \quad x_u^{(i)} \in R^n, k \gg m$$



Target

$$\{(x_l^{(i)}, y^{(i)})\}_{i=1}^m \quad x_l^{(i)} \in R^n, y^{(i)} \in \{1, \dots, T\}$$



Learn a higher-level, more abstract feature for the given feature type (modality)

Self-taught Learning with Sparse Coding

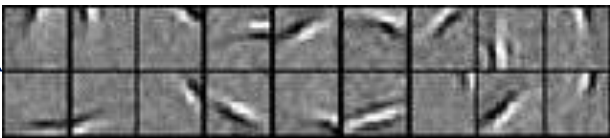
Given unlabeled data x , find good bases b using sparse coding and dictionary learning.

$$\min_{a,b} \sum_i \underbrace{\|x_u^{(i)} - \sum_j \underbrace{a_j^{(i)}}_{\text{Sparse codes}} \underbrace{b_j}_{\text{Bases}}\|_2^2}_{\text{Reconstruction error}} + \beta \sum_i \underbrace{\|a^{(i)}\|_1}_{\text{Sparsity Regularization}} \quad \underbrace{\|b_j\|_2^2 \leq 1, \forall j \in 1, \dots, s}_{\text{L2 regularization}}$$

Natural Image



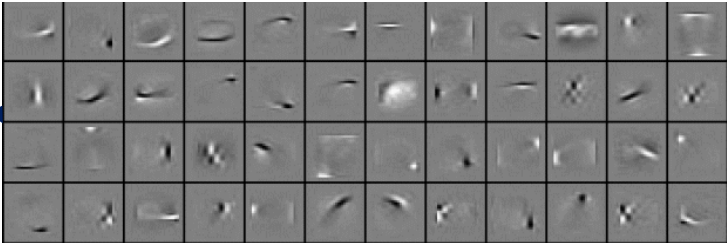
Edges



Handwritten Characters

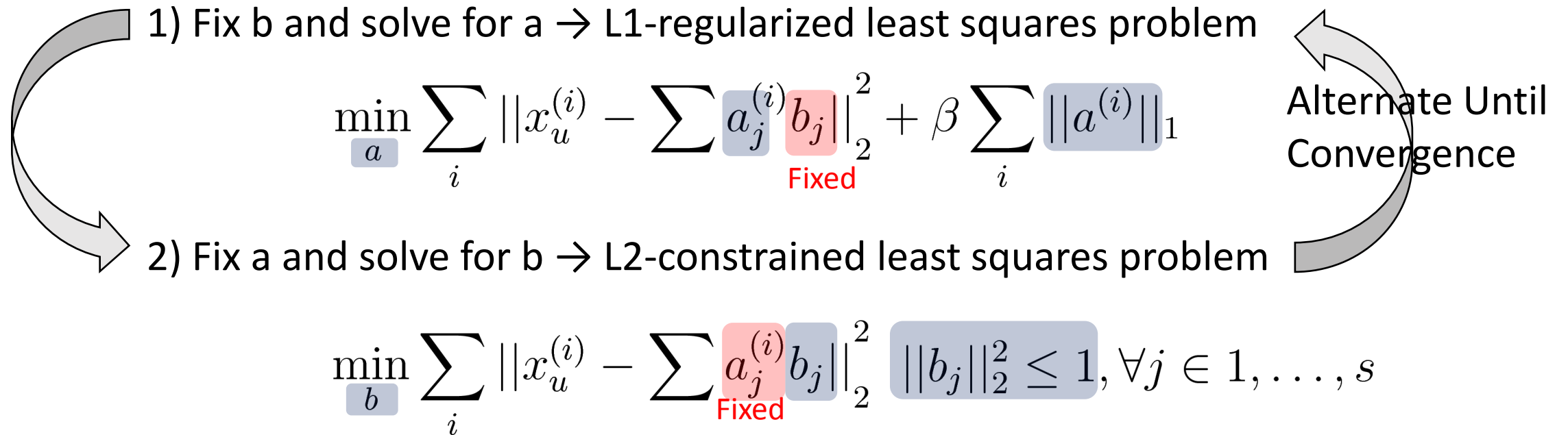


Strokes



Self-taught Learning with Sparse Coding

Use alternating optimization - Solve for each variable while fixing the other, and alternate the process until convergence (Efficient algorithm introduced in 07)



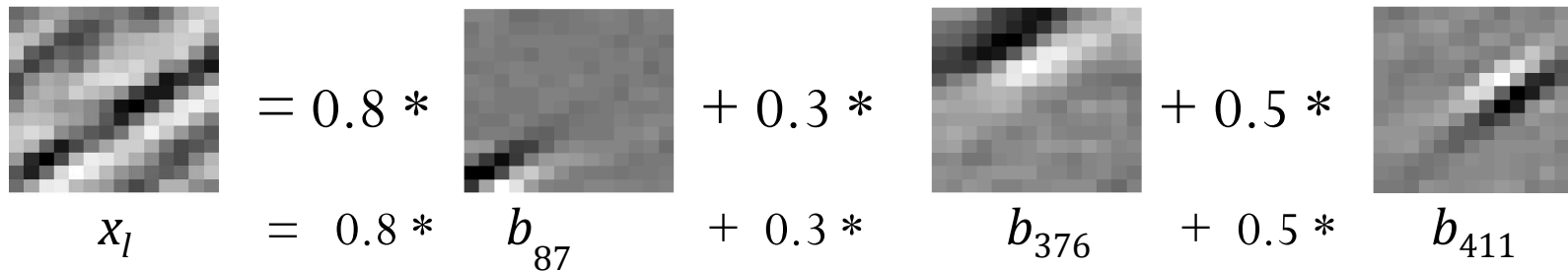
[Raina07] R. Raina, A. Battle, H. Lee, B. Packer, A. Y. Ng, Self-taught Learning: Transfer Learning from Unlabeled Data, ICML 2007

[Lee07] H. Lee, A. Battle, R. Raina, A. Y. Ng, Efficient Sparse Coding Algorithms, NIPS 2007

Self-taught Learning with Sparse Coding

Then reconstruct the examples in the target dataset using the learned bases.

$$\min_{a,b} \sum_i \underbrace{\left\| x_u^{(i)} - \sum_j \underbrace{a_j^{(i)}}_{\text{Sparse codes}} \underbrace{b_j}_{\text{Bases}} \right\|_2^2}_{\text{Reconstruction error}} + \beta \sum_i \|a^{(i)}\|_1 \quad \|b_j\|_2^2 \leq 1, \forall j \in 1, \dots, s$$


$$x_l = 0.8 * b_{87} + 0.3 * b_{376} + 0.5 * b_{411}$$

Then train a classifier (such as a SVM) over the obtained sparse codes

Self-taught Learning with Sparse Coding

Results show that information gained from unlabeled data is actually useful.

Image Classification



Method	Accuracy
Baseline	16%
PCA	37%
Sparse coding	47%

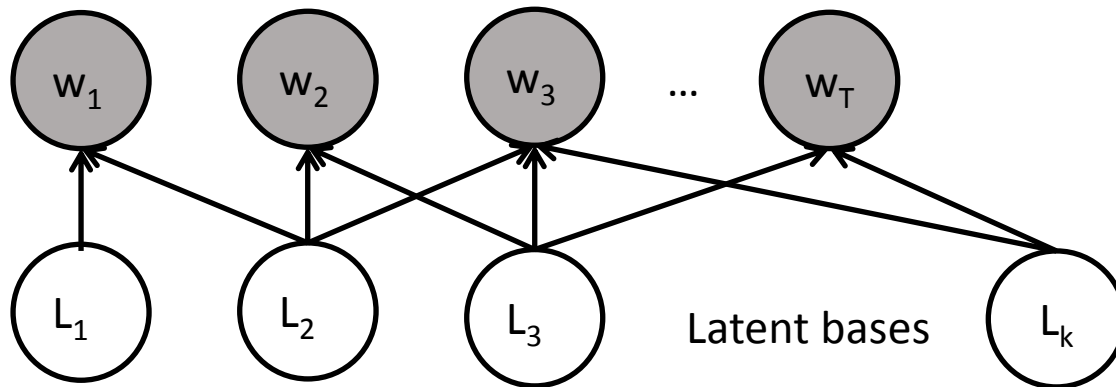
Handwritten Character Recognition



Method	Accuracy
Raw	54.8%
PCA	54.8%
Sparse coding	58.5%

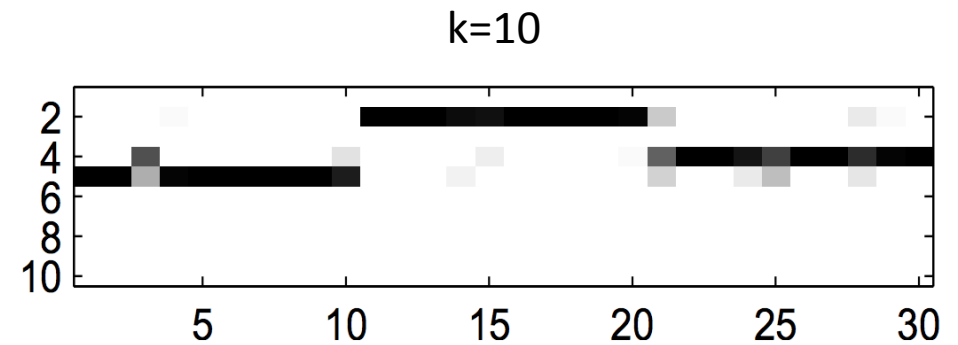
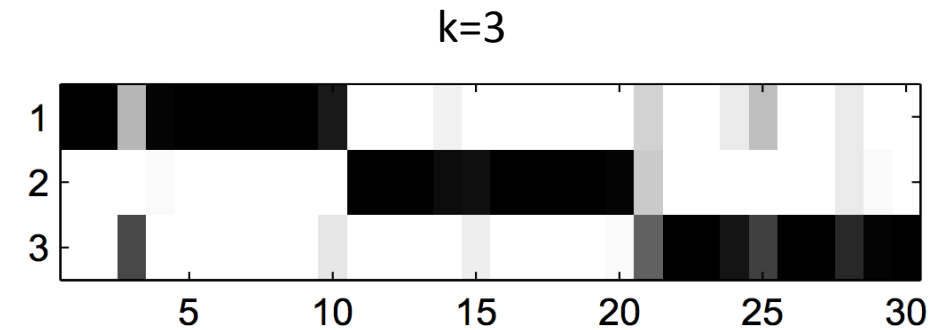
Learning Task Grouping in Multitask Learning

Allow the learners to selectively share the information across the tasks.



$$\sum_t \sum_{(x_{ti}, y_{ti}) \in Z_t} \mathcal{L}(y_{ti}, \underbrace{s_t'}_{\text{Task-specific sparse weight}} \underbrace{L'}_{\text{Shared parameter bases}} x_{ti}) + \mu ||S||_1 + \lambda ||L||_F^2$$

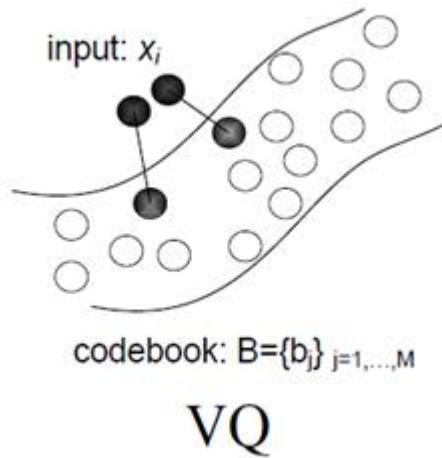
$$w_t = L s_t$$



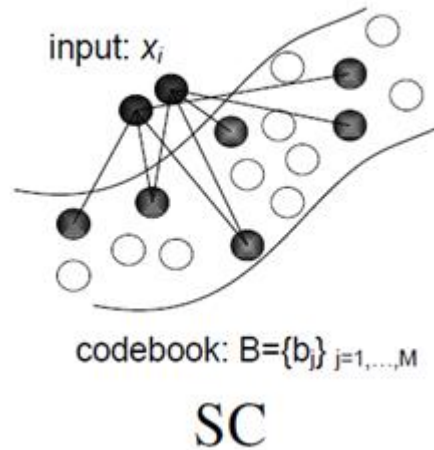
Can discover true support without having to provide the number of groups

Sparse Feature Encoding

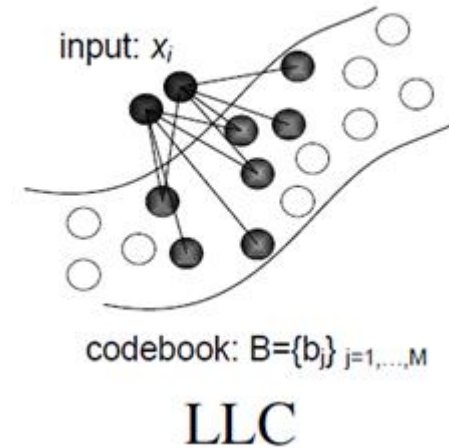
Soft-Encoding – encodes each feature as a sparse combination of visual words.



$$\arg \min_{\mathbf{C}} \sum_{i=1}^N \|\mathbf{x}_i - \mathbf{B}\mathbf{c}_i\|^2$$
$$s.t. \|\mathbf{c}_i\|_{\ell^0} = 1, \|\mathbf{c}_i\|_{\ell^1} = 1, \mathbf{c}_i \succeq 0, \forall i$$



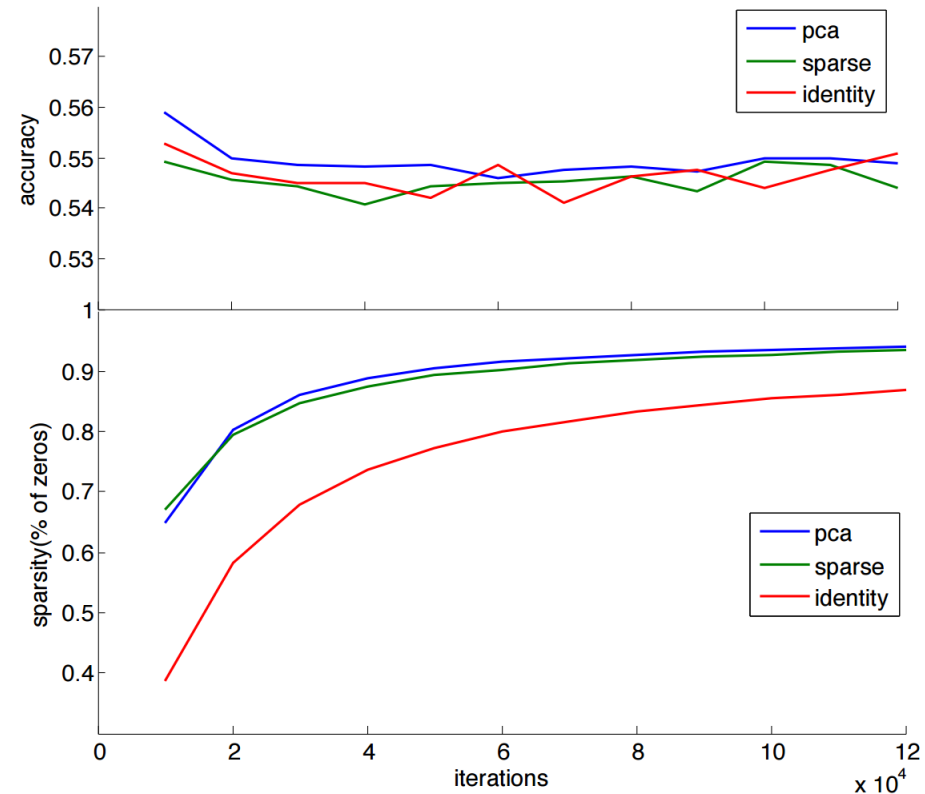
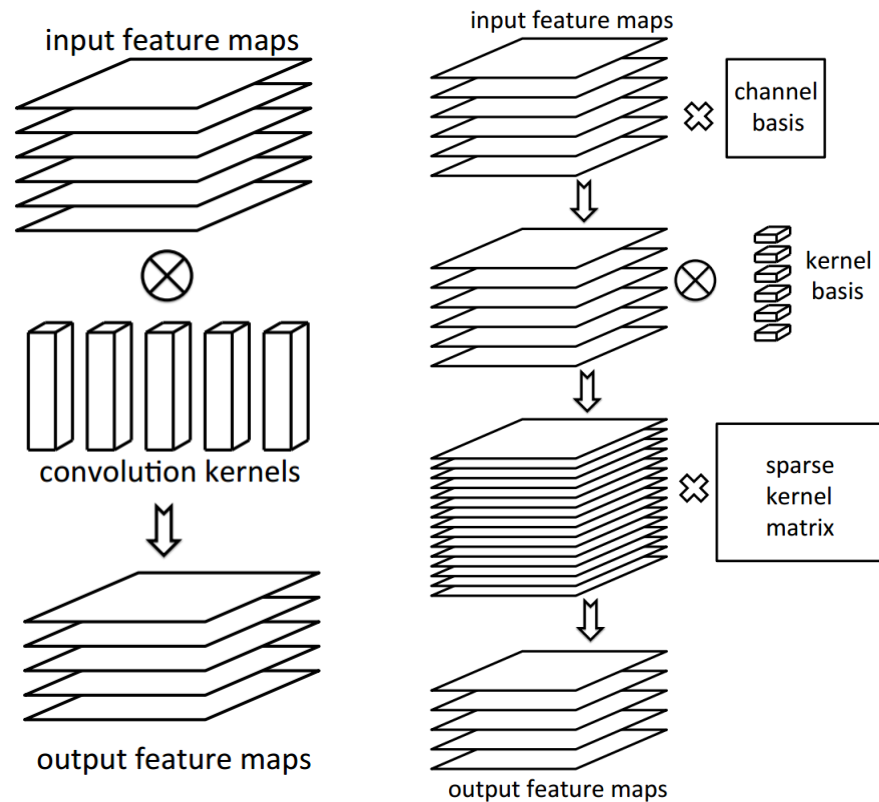
$$\arg \min_{\mathbf{C}} \sum_{i=1}^N \|\mathbf{x}_i - \mathbf{B}\mathbf{c}_i\|^2 + \lambda \|\mathbf{c}_i\|_{\ell^1}$$



$$\min_{\mathbf{C}} \sum_{i=1}^N \|\mathbf{x}_i - \mathbf{B}\mathbf{c}_i\|^2 + \lambda \|\mathbf{d}_i \odot \mathbf{c}_i\|^2$$
$$s.t. \mathbf{1}^\top \mathbf{c}_i = 1, \forall i$$

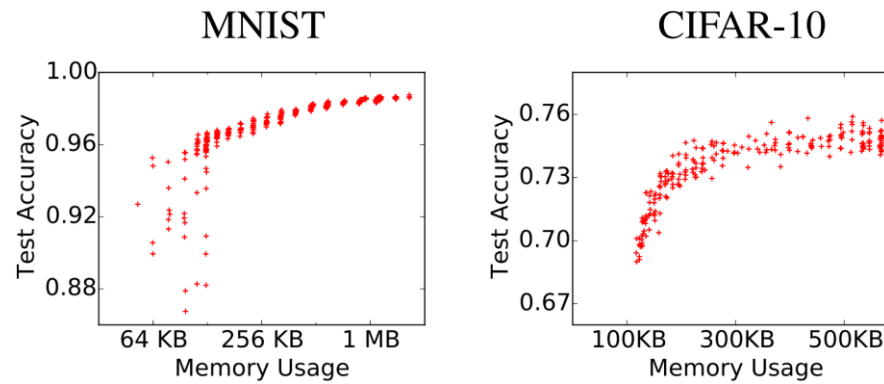
Learning Sparse Filters for a CNN

Use sparse coding to reduce the number of convolutional kernels



Learning Sparse Weights for a CNN

Learn sparse network weights to reduce memory usage



ImageNet			
Model	Top-1	Top-5	Memory
Reported [11]	59.3%	81.8%	-
Caffe Version [9]	57.28%	80.44%	233 MB
Sparse (ours)	55.60%	80.40%	58 MB

Summary

Sparsity regularizations result in learning a model with ***few nonzero*** parameters. It is beneficial for ***feature selection***, ***model analysis***, and ***model compression***.

L1-regularization uses 1-norm for regularization, which enduces sparsity. Since 1-norm is ***non-differentiable***, we can solve it through optimization methods such as ***subgradient descent***, ***proximal gradient***, and ***regularized dual averaging method***.

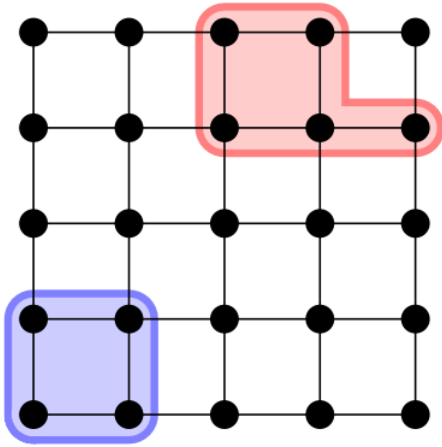
Sparse coding encodes an input variable as a ***sparse combination of some bases***. It is useful for multiple applications such as ***image restoration*** and ***transfer learning***.

Part 2: Structured Sparsity

- Introduction to Structured Sparsity
- Group-structured sparsity
- Block-structured sparsity
- Optimizing for $(2,1)$ -norm
- Tree-structured sparsity
- Graph-structured sparsity
- Example: Image inpainting with hierarchical dictionary learning
- Example: Multi-task learning
- Example: Learning decorrelated attributes

Structured Sparsity

Sparsity regularizations that prefer *certain structure* among the variables over others.



Enables to exploit known structure – e.g.) To reconstruct handwritten characters, we can exploit the fact that they form connected components

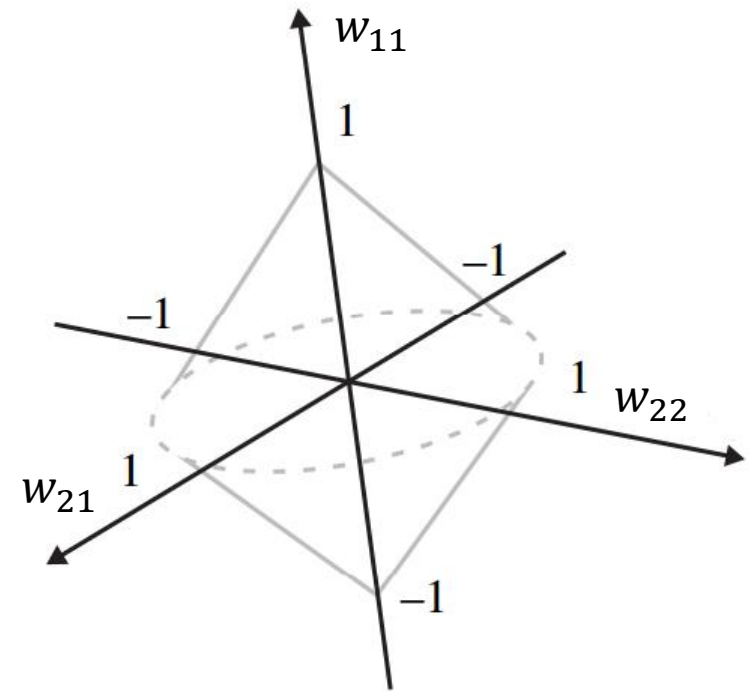
(2,1)-norm

Mixed norm that has 1-norm over 2-norm groups.

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \|\mathbf{w}\|_{2,1}$$

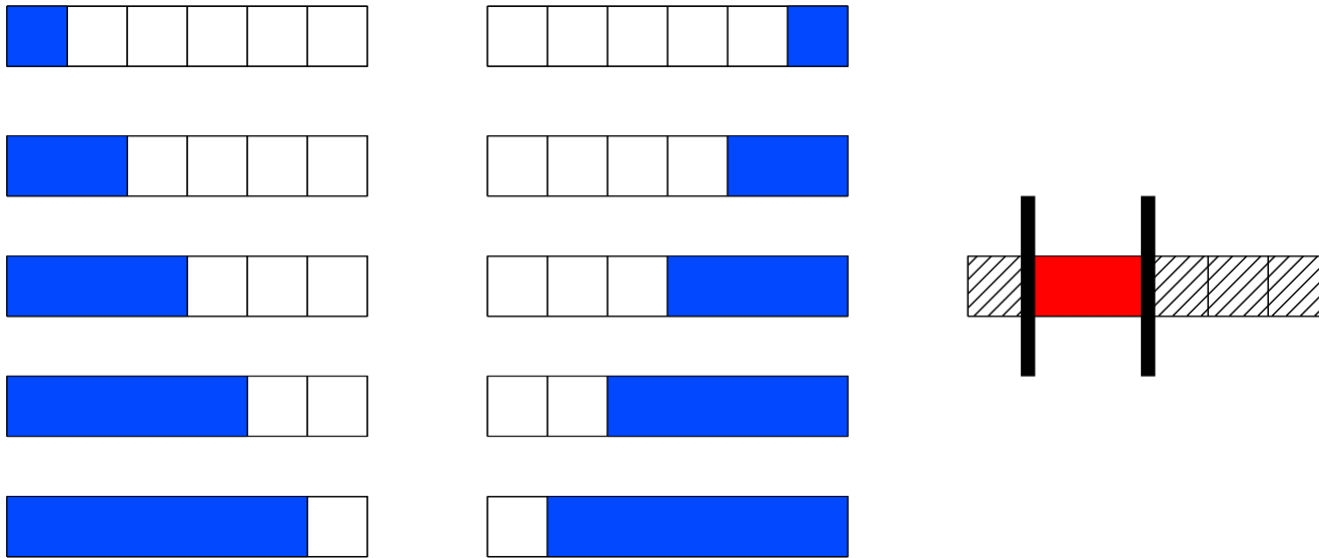
$$\|\mathbf{w}\|_{2,1} = \sum_l^L \underbrace{\|\mathbf{w}_l\|_2}_{\text{2-norm}} \quad \text{1-norm}$$

Used to promote sparsity at group level



Group Sparsity

Structured sparsity with groups of variables



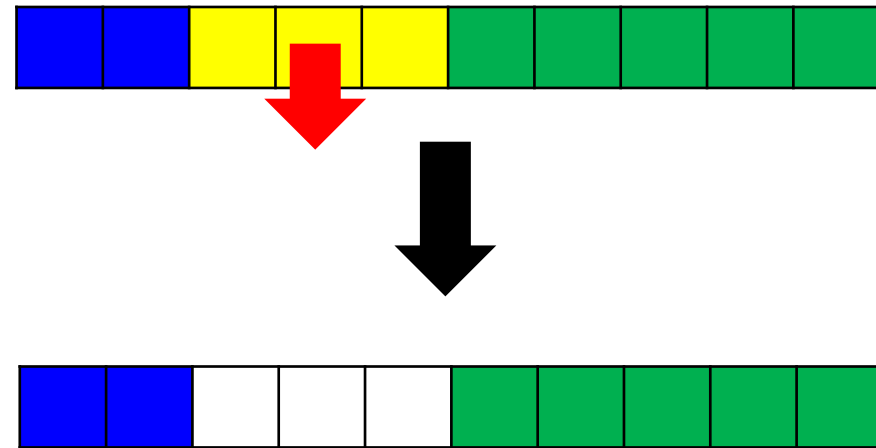
Group Sparsity

Use L2/L1-regularization to select / drop variables in a group

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \|\mathbf{w}\|_{2,1}$$

$$\|\mathbf{w}\|_{2,1} = \sum_l^L \|\mathbf{w}_l\|_2$$

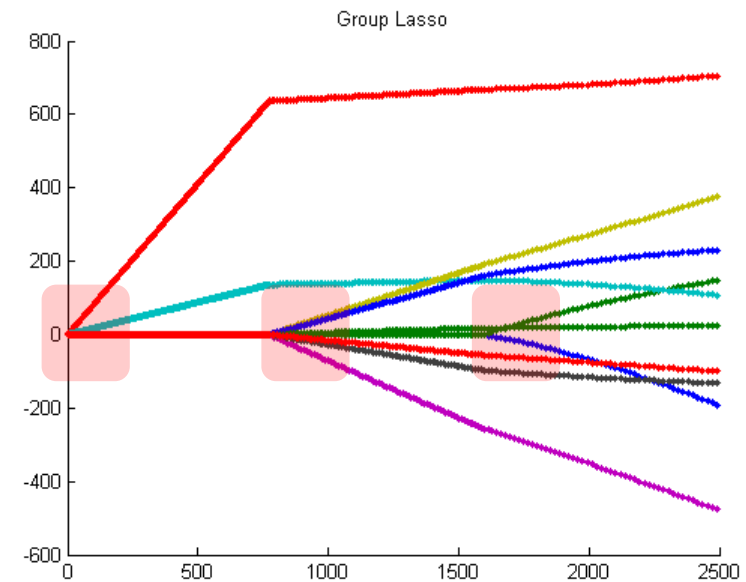
Used to promote sparsity at group level



Group Lasso

Correlated variables gets selected / dropped at the same time

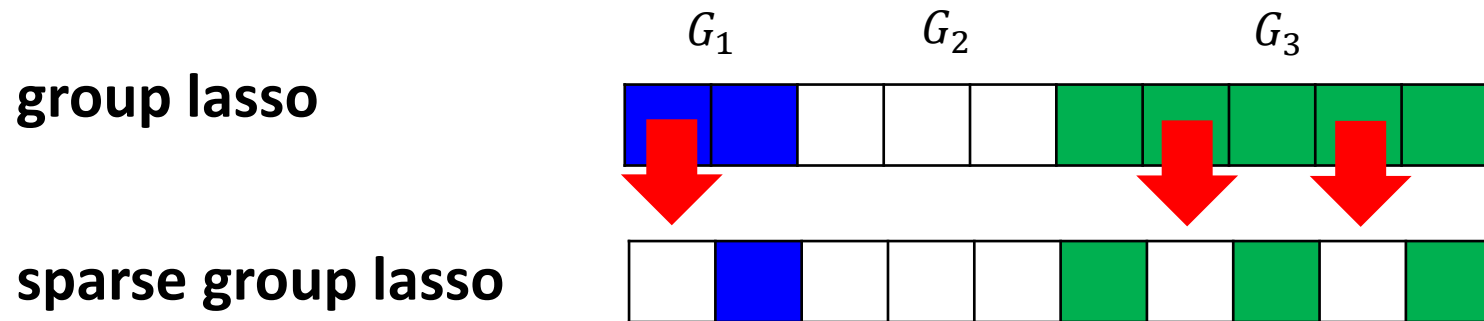
$$\min_{\mathbf{w}} \frac{1}{N} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 + \lambda \sum_{l=1}^L \|\mathbf{w}_l\|_2$$



Sparse Group Lasso

Group lasso does not yield sparsity within a group. All variables selected will have non-zero values.

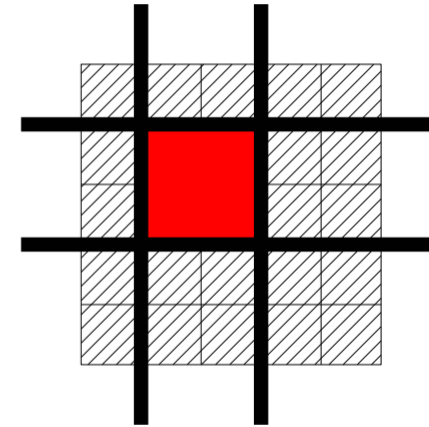
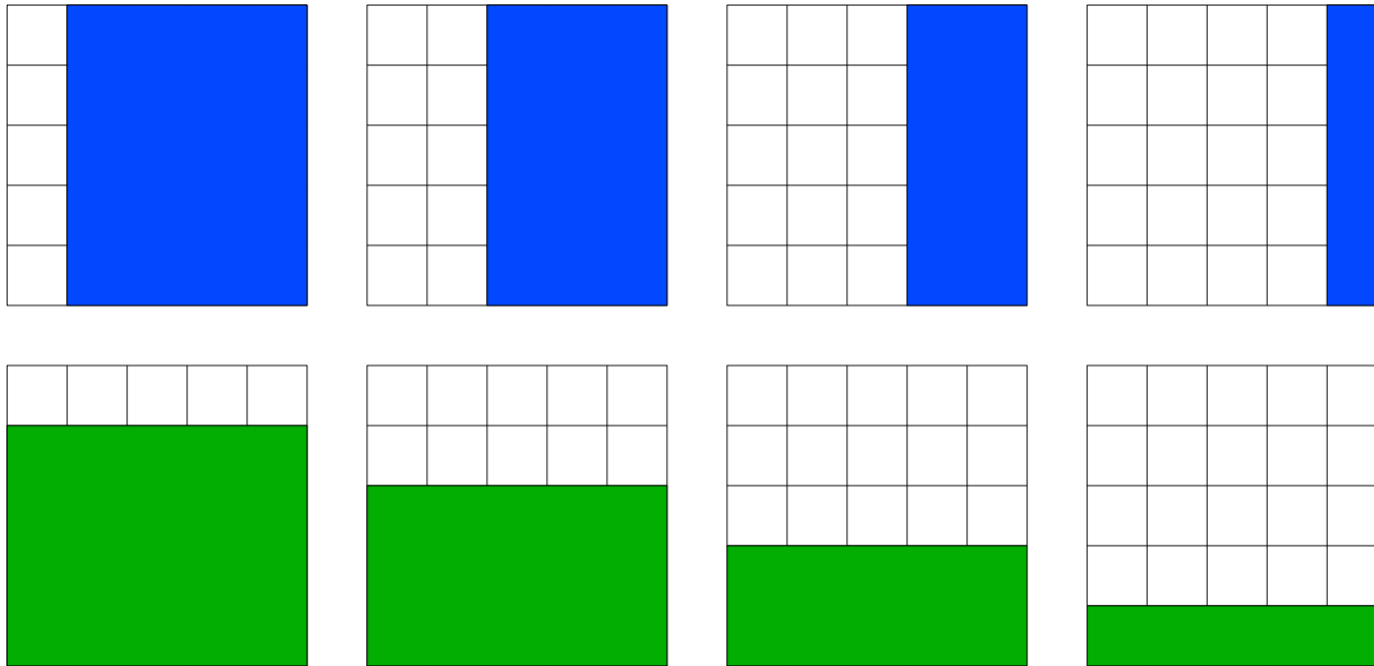
$$\min_{\mathbf{w}} \frac{1}{N} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 + \lambda_1 \sum_{l=1}^L \|\mathbf{w}_l\|_2 + \lambda_2 \|\mathbf{w}\|_1$$



Use additional l1-regularization to yield sparsity at individual feature level.

Block Sparsity

Structured Sparsity on a 2D Grid

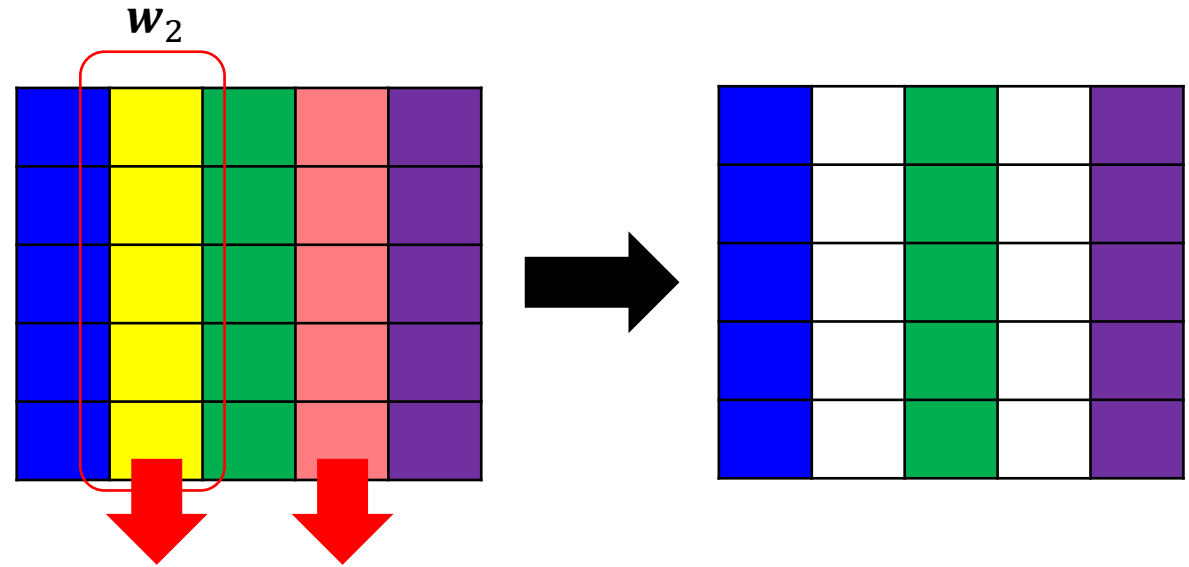


Block Sparsity

Same as group sparsity applied to a sequence. Group each row or columns using 2-norm, and apply 1-norm on the groups

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \|\mathbf{w}\|_{2,1}$$

$$\|\mathbf{w}\|_{2,1} = \sum_l^L \|\mathbf{w}_l\|_2$$



Exclusive Lasso

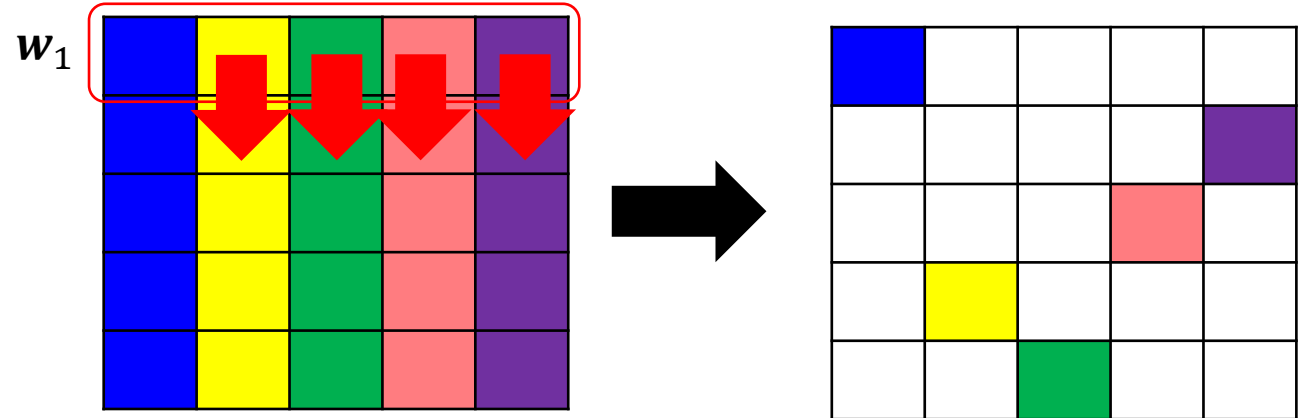
Promote competition for the features between tasks

$$\min_{\mathbf{w}} l(\mathbf{X}, \mathbf{y}, \mathbf{w}) + \lambda \Omega(\mathbf{w})$$

2-norm over the parameter for the same task

$$\Omega(\mathbf{w}) = \sum_j^d \left(\sum_k^m |w_k^j| \right)^2$$

1-norm over the parameters for m tasks on the same feature dimension



Optimization for L2/L1-Regularized Objectives

(2,1)-norm is nonsmooth.

proximal operator

$$\text{prox}(\mathbf{w})_l = \left[1 - \frac{\lambda}{\|\mathbf{w}_l\|} \right]_+ \mathbf{w}_l$$

Regularized dual averaging

$$w_{t+1}^i = \begin{cases} 0, & \text{if } |g_t^i| \leq \lambda \\ -\frac{1}{\sigma} \left(g_t^i - \lambda \text{sign}(g_t^i) \right), & \text{otherwise} \end{cases}$$

Can use proximal gradient and regularized dual averaging method

Block Coordinate Descent

Optimize a group of variable at a time, while fixing all others.

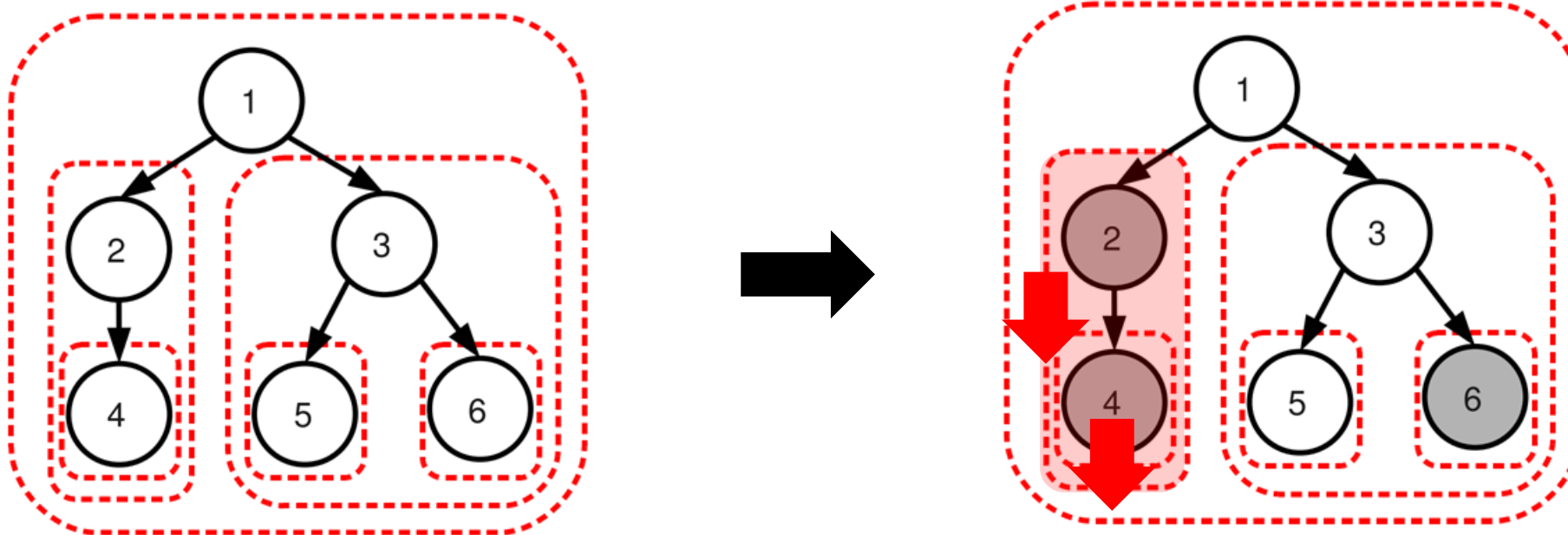
$$\mathbf{w}_g^* = \text{prox}_{\frac{\lambda}{L_t}} \left(\mathbf{w}_g - \frac{1}{L_t} \nabla_g f(\mathbf{w}_g^t) \right)$$

$$\text{prox}(\mathbf{w}) = \left[1 - \frac{\lambda}{\|\mathbf{w}\|} \right]_+ \mathbf{w}$$

The solution $\mathbf{w}_g^*$ can be obtained through group-soft thresholding

Tree Sparsity

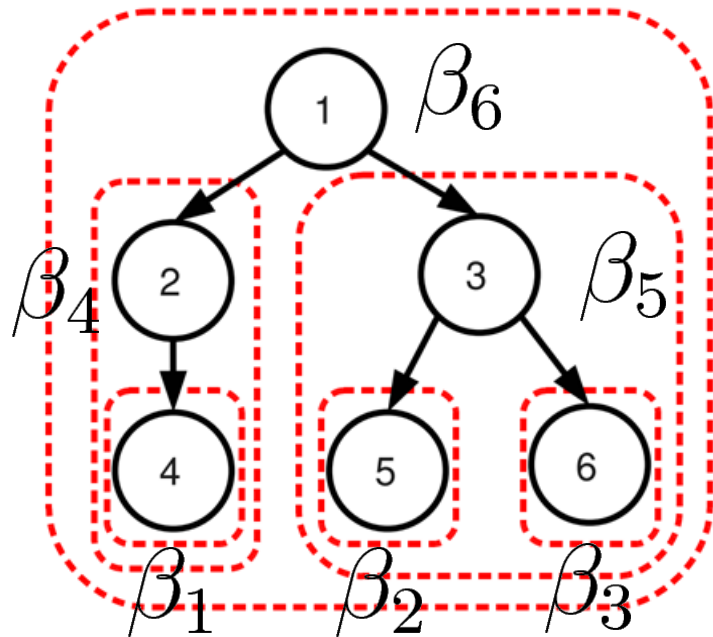
Structural sparsity on a tree



If a node is removed, then all its descendant nodes are dropped.

Hierarchical Sparsity-Inducing Norm

Perform group lasso where each group $\mathcal{G}_j$ contains node j and all its descendants.



$$\Omega(\beta) = \sum_{g \in \mathcal{G}} w_g \|\beta_g\|_2$$

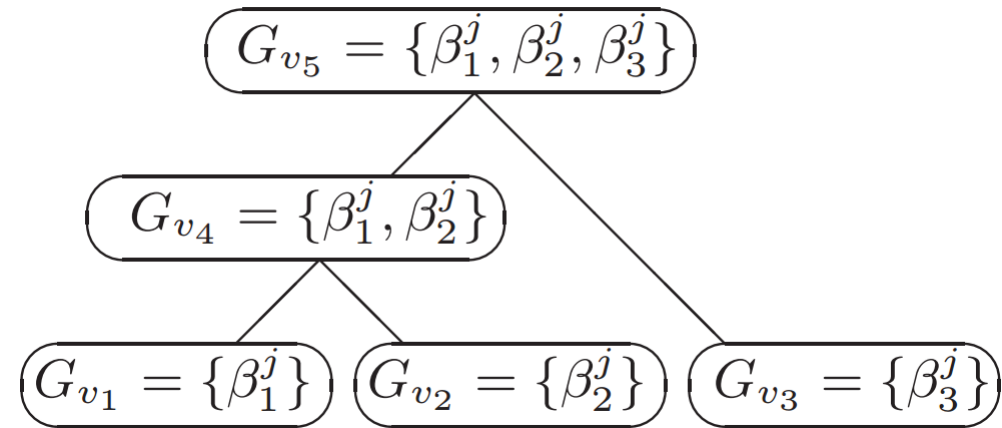
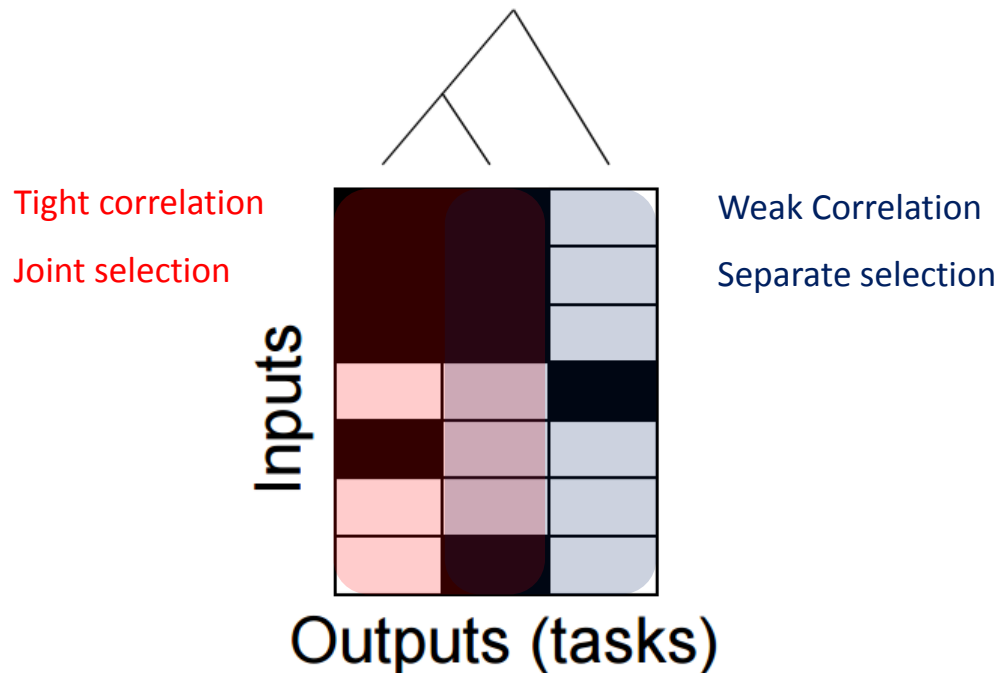
$$\beta_1 = \beta_1 \quad \beta_2 = \beta_2 \quad \beta_3 = \beta_3$$

$$\beta_4 = \{\beta_1, \beta_4\} \quad \beta_5 = \{\beta_2, \beta_3, \beta_5\}$$

$$\beta_6 = \{\beta_1, \beta_2, \beta_3, \beta_4, \beta_5\}$$

Tree-guided Group Lasso

Exploit a given tree structure on the output values to recover the true structure in the parameters



Idea: use overlapping groups in group lasso

Tree-guided Group Lasso

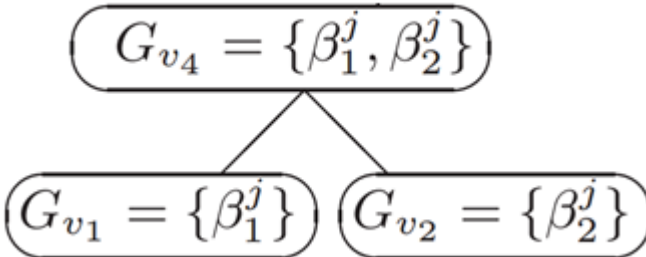
Promotes sharing and competition between tasks in hierarchical manner.

L1 regularization – promotes competition

$$\sum_j \left[h(|\beta_1^j| + |\beta_2^j|) + (1 - h) \left(\sqrt{(\beta_1^j)^2 + (\beta_2^j)^2} \right) \right]$$

L2 regularization – promotes sharing

Decides the degree of competition and sharing



```
graph TD; G4("G_{v_4} = {\beta_1^j, \beta_2^j}") --- G1("G_{v_1} = {\beta_1^j}"); G4 --- G2("G_{v_2} = {\beta_2^j}")
```

Simplest case where there are only two outputs

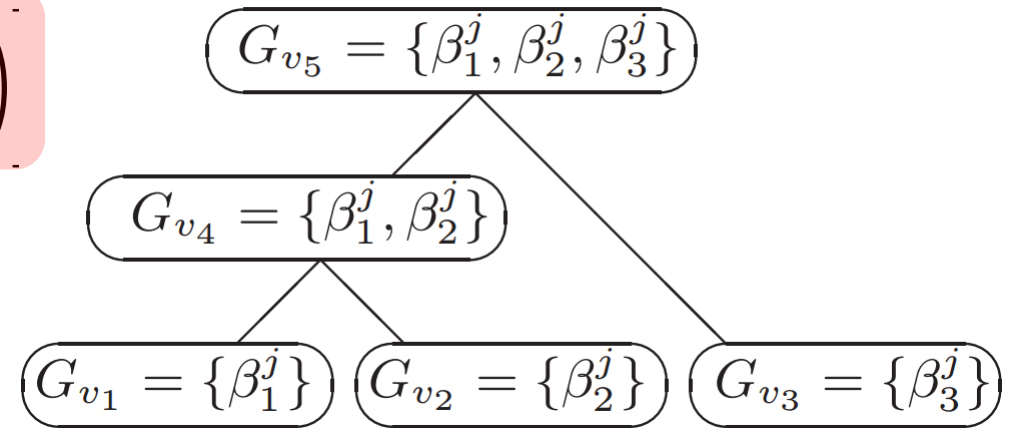
Tree-guided Group Lasso

Promotes sharing and competition between tasks in hierarchical manner.

$$\sum_j \left[h_2(|G_{v_4}| + |\beta_3^j|) + (1 - h_2) \left(\sqrt{(\beta_1^j)^2 + (\beta_2^j)^2 + (\beta_3^j)^2} \right) \right]$$

competition sharing

$$\sum_j \left[h(|\beta_1^j| + |\beta_2^j|) + (1 - h) \left(\sqrt{(\beta_1^j)^2 + (\beta_2^j)^2} \right) \right]$$



True coefficients

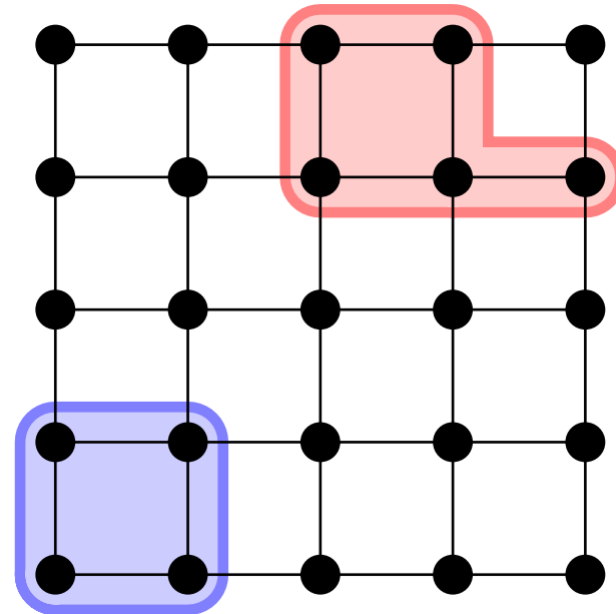
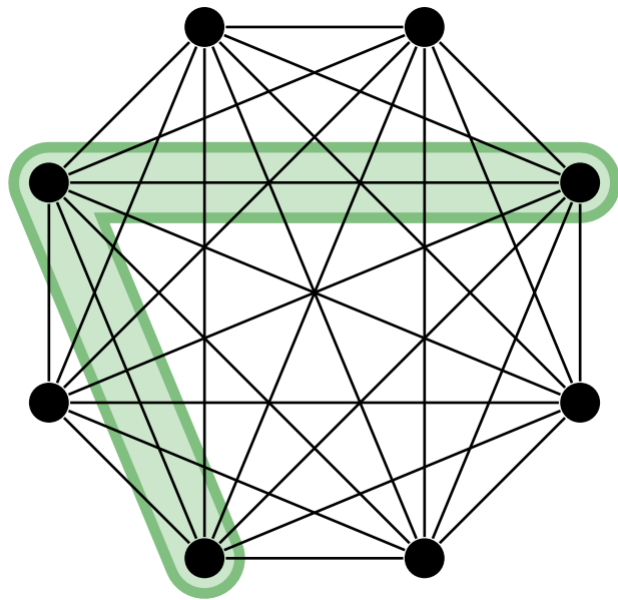
Lasso

Group lasso

Tree-guided Group lasso

Graph-Structured Sparsity

Structured sparsity on a generic graph

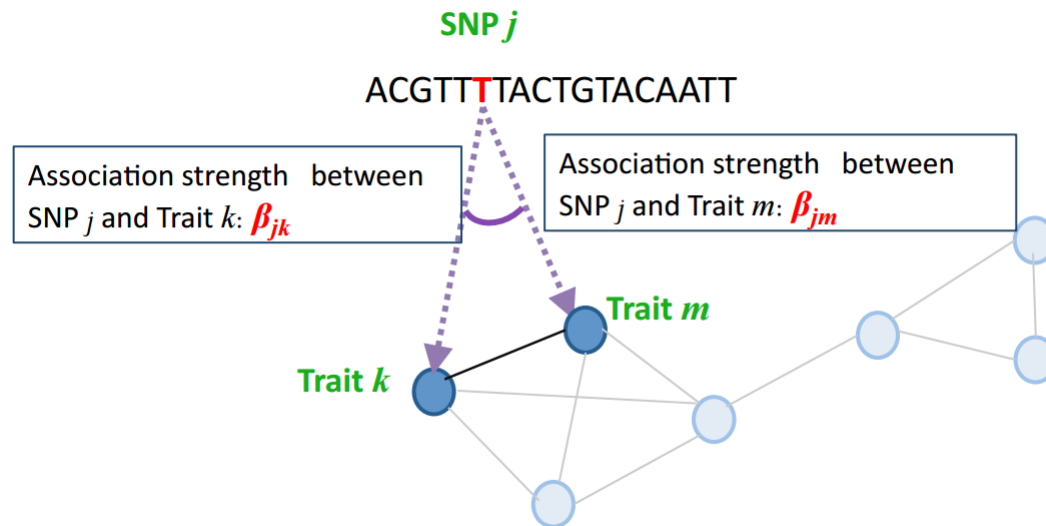


Graph-Guided Fused Lasso

Use known graph structure among the output to constrain correlated variables to have similar parameters

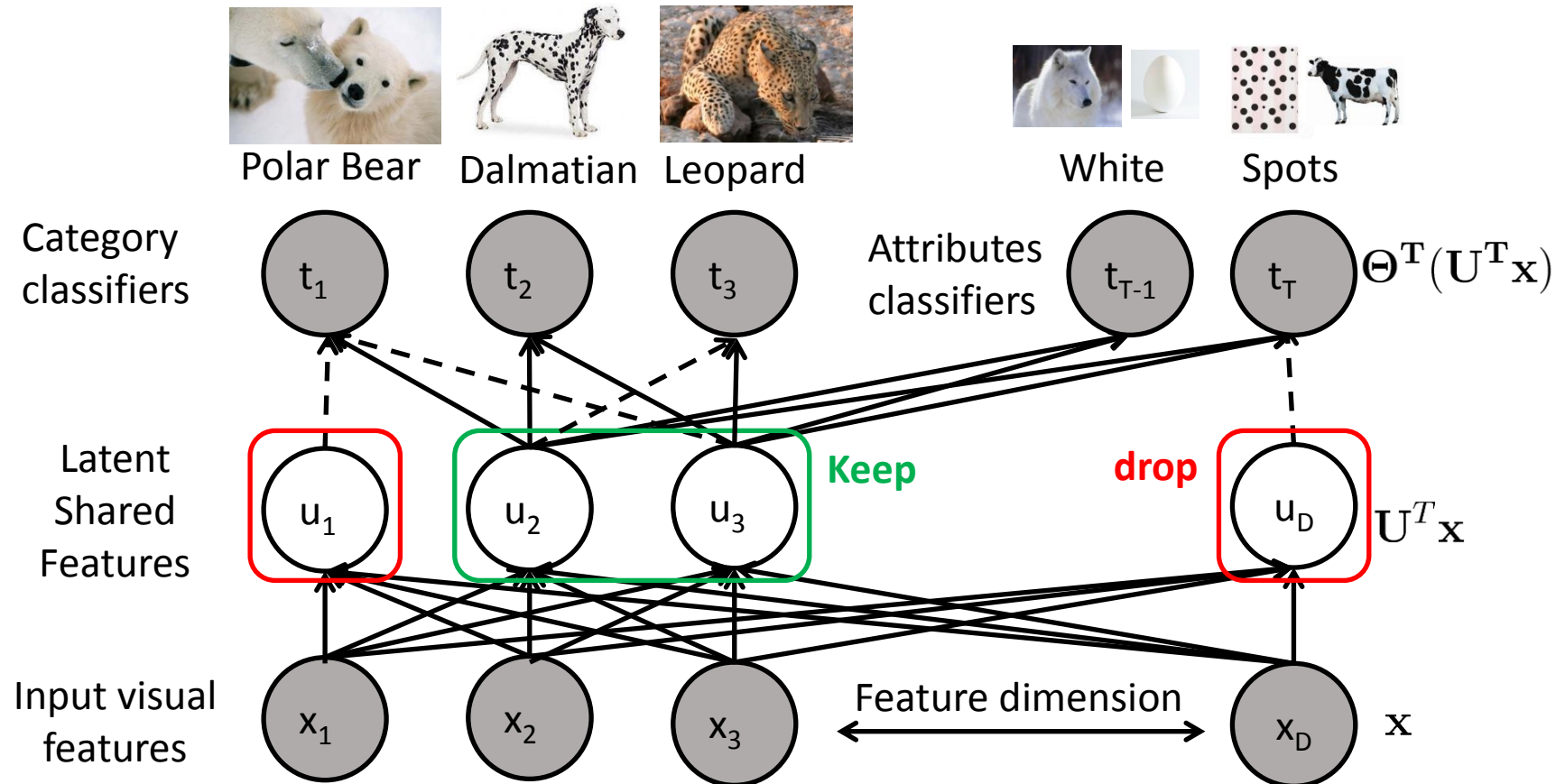
$$\Omega(B) = \sum_k \sum_j |\beta_{jk}| + \gamma \sum_{m,l \in E} \sum_j |\beta_{jm} - \text{sgn}(r_{ml})\beta_{jl}|$$

sparsity Fusion penalty



Multitask Feature Learning

Assume that there exist some latent shared features that are shared across multiple tasks, and learn them.



Multitask Feature Learning

Given N input features x and label y_t for each task t , Simultaneously learn the transformation U and model parameter θ_t for each class t

The diagram shows the equation $\Theta^*, U^* = \arg \min \sum_t \sum_n \ell(\theta_t^T U^T x_n, y_{nt})$ with several annotations. A green box highlights the term t_{th} classifier above the summation index t . A yellow box highlights the term n_{th} training example above the summation index n . A blue arrow points from the text 'Classifier' to the term θ_t^T . A purple arrow points from the text 'Transformation to a shared feature space' to the term U^T . A purple arrow points from the text 'Prediction loss' to the loss function ℓ . The entire expression $\ell(\theta_t^T U^T x_n, y_{nt})$ is enclosed in a light green rounded rectangle.

$$\Theta^*, U^* = \arg \min \sum_t \sum_n \ell(\theta_t^T U^T x_n, y_{nt})$$

t_{th} classifier n_{th} training example Classifier

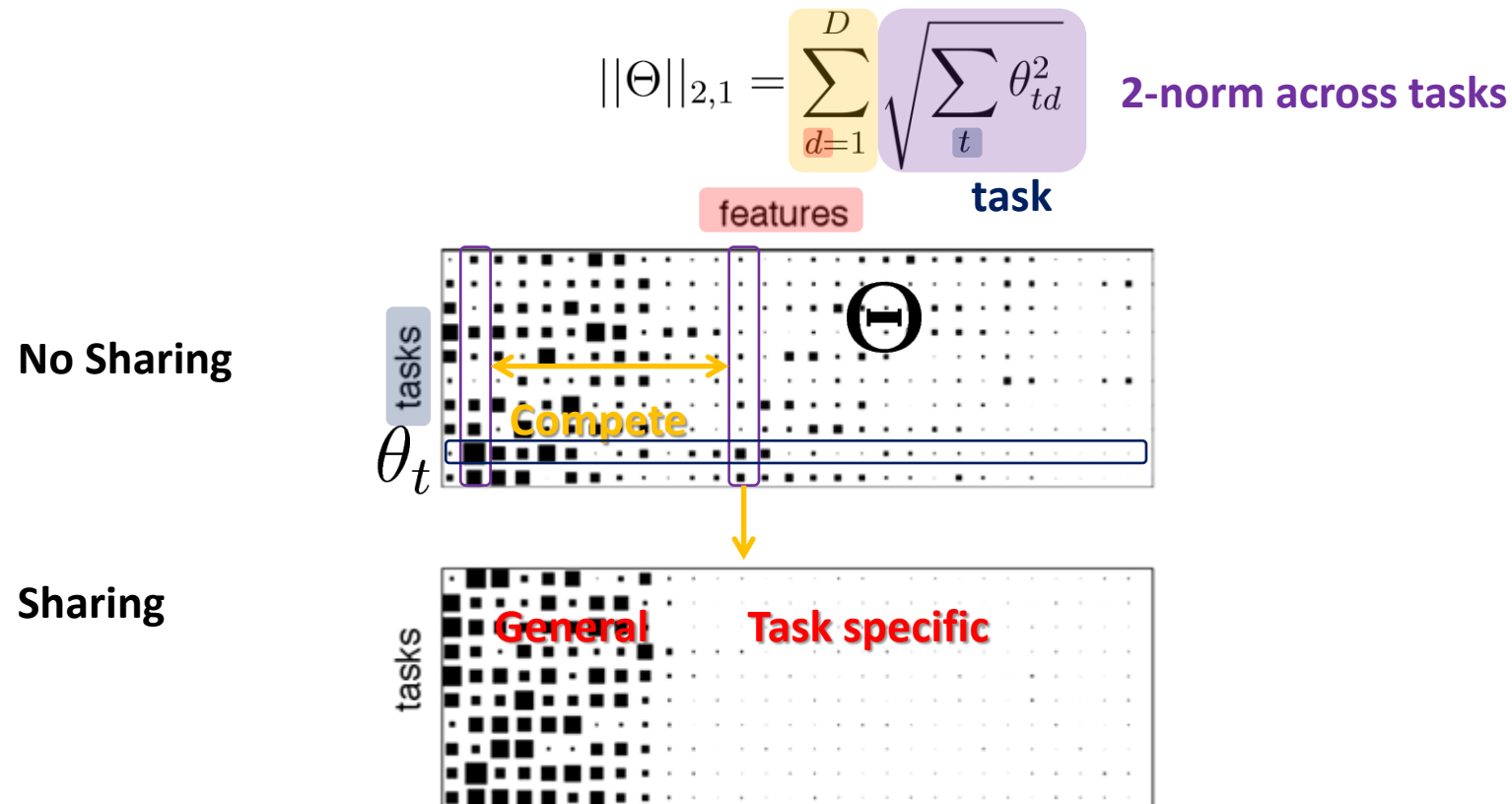
Prediction loss

Transformation to a shared feature space

1. Learn classifiers on the transformed features in shared feature space
2. Promote a common sparsity pattern in the new parameters

Sharing features via Sparsity Regularization

Using group sparsity regularization, enforce each learner to use features that are informative across multiple tasks.



Convex Multitask Feature Learning

Previous formulation is non-smooth, which makes it challenging to solve. Thus we solve an equivalent form instead. → Replace features with a covariance matrix Ω that measures the relative effectiveness of each dimension

w_t : model parameter for task t ,
on **original** features. Ω : covariance matrix

$$\begin{aligned} W^*, \Omega^* = \arg \min & \sum_t \sum_n \ell(w_t^T x_n, y_{nt}) \quad \text{Model prediction loss in} \\ & \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \text{the **original** feature space} \\ & + \gamma \sum_t w_t^T \Omega^{-1} w_t + \gamma \epsilon \text{Trace}(\Omega^{-1}) \\ & \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \text{Trace norm} \\ & \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \quad \text{(sum of the diagonal entries for a PSD matrix)} \end{aligned}$$

Multitask Feature Learning - Result

Dataset

Animals with Attributes (AWA)

30,475 images, 50 classes, 85 attributes

White



Polar bear

Spotted



Dalmatian



Leopard



Horse



Lion



Cow

Outdoor Scene Recognition (OSR)

2,688 images, 8 classes, 6 attributes

Open



Coast

Natural



Mountain



Forest



Highway



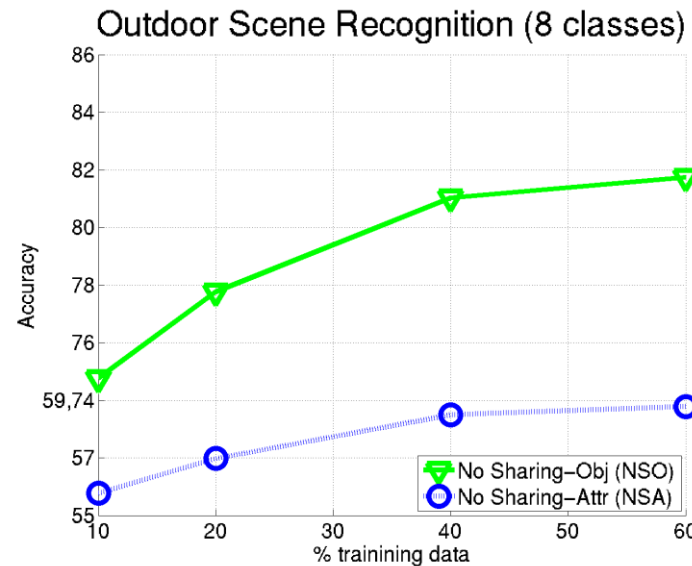
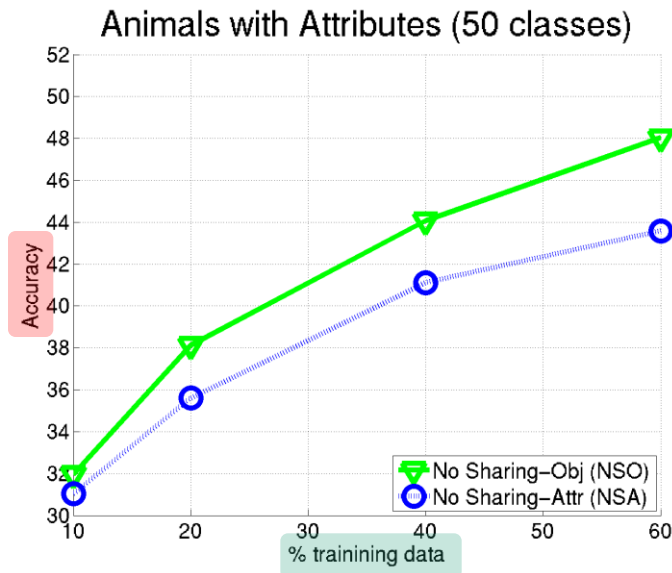
Street



Tall
Buildings

Multitask Feature Learning - Result

No Sharing (Visual Features) > Attributes-based Prediction

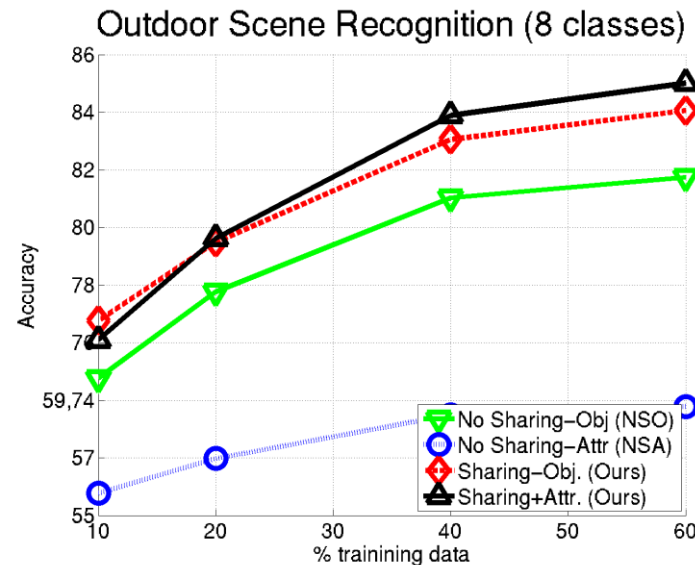
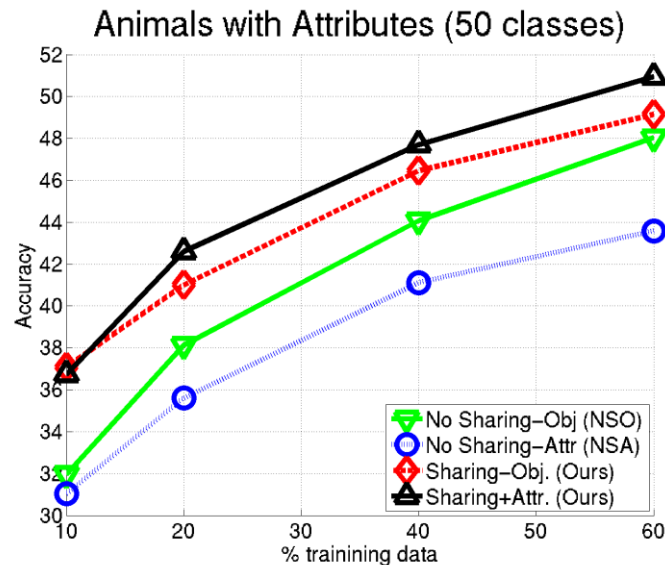


No sharing-Obj.: Independent SVMs trained on visual features

No sharing-Attr. : Object recognition on predicted attributes as in [Lampert09].

Multitask Feature Learning - Result

Shared Features (Object & Attributes) > Shared Features (Objects) > No Sharing



Sharing-Obj.: Multitask Feature Learning on object classifiers

Sharing-Attr. : Multitask Feature Learning on object + attribute classifiers

Multitask Feature Learning - Result

Predicted object categories **Red: incorrect prediction**

More robustness to background clutter from features refined with attributes



No Sharing **Dolphin**

Ours Grizzly Bear



No Sharing **Polar Bear**

Ours Dalmatian

Even when our method fails, it often makes more semantically “close” predictions.



No Sharing **Giant Panda**

Ours Rhinoceros

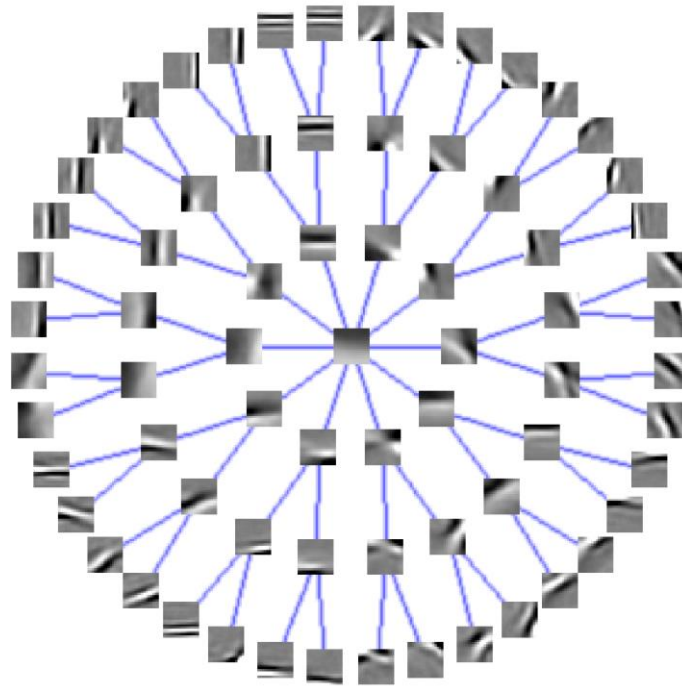


No Sharing **Cow**

Ours Wolf

Image Inpainting with Hierarchical Sparse Coding

Use hierarchical sparsity inducing-norm to reconstruct an input image with a hierarchy of image patches



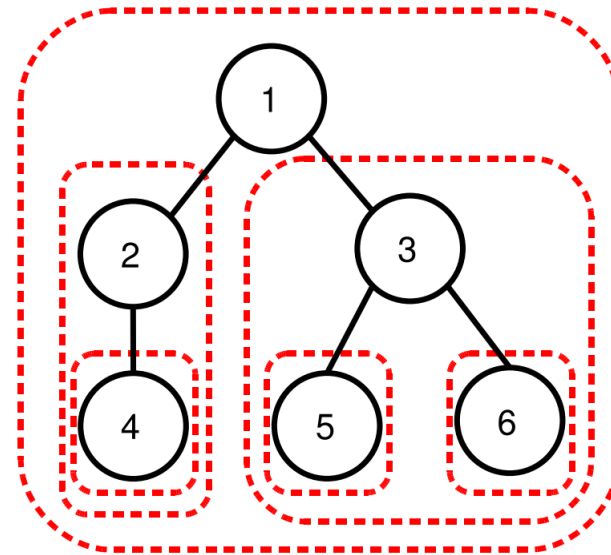
Hierarchical Sparse Coding

Use proximal operator to solve for the hierarchical sparsity-inducing norm

$$\min_{\mathbf{D} \in \mathcal{D}, \mathbf{A} \in \mathcal{A}} \frac{1}{n} \sum_{i=1}^n \left[\frac{1}{2} \|\mathbf{x}^i - \mathbf{D}\boldsymbol{\alpha}^i\|_2^2 + \lambda \Omega(\boldsymbol{\alpha}^i) \right]$$

$$\Omega(\boldsymbol{\alpha}) \triangleq \sum_{g \in \mathcal{G}} w_g \|\boldsymbol{\alpha}_{|g}\|$$

$$\min_{\mathbf{v} \in \mathbb{R}^p} \frac{1}{2} \|\mathbf{u} - \mathbf{v}\|_2^2 + \lambda \Omega(\mathbf{v})$$



Problem – No closed form solution exists for group lasso with overlapping groups

Hierarchical Sparse Coding

Use primal-dual approach to solve for the dual problem

Primal
$$\min_{\mathbf{v} \in \mathbb{R}^p} \frac{1}{2} \|\mathbf{u} - \mathbf{v}\|_2^2 + \lambda \Omega(\mathbf{v})$$

Dual
$$\max_{\boldsymbol{\xi} \in \mathbb{R}^{p \times |\mathcal{G}|}} -\frac{1}{2} \left(\|\mathbf{u} - \sum_{g \in \mathcal{G}} \boldsymbol{\xi}^g\|_2^2 - \|\mathbf{u}\|_2^2 \right)$$

$$s.t. \forall g \in \mathcal{G}, \|\boldsymbol{\xi}^g\|_* \leq \lambda w_g \text{ and } \xi_j^g = 0 \text{ if } j \notin g$$

Algorithm 1 Block coordinate ascent in the dual

Inputs: $\mathbf{u} \in \mathbb{R}^p$ and set of groups $\mathcal{G}$.

Outputs: $(\mathbf{v}, \boldsymbol{\xi})$ (primal-dual solutions).

Initialization: $\mathbf{v} = \mathbf{u}, \boldsymbol{\xi} = 0$.

while (*maximum number of iterations not reached*) **do**

for $g \in \mathcal{G}$ **do**

$\mathbf{v} \leftarrow \mathbf{u} - \sum_{h \neq g} \boldsymbol{\xi}^h$.

$\boldsymbol{\xi}^g \leftarrow \Pi_{\lambda w_g}^*(\mathbf{v}_{|g})$.

end for

end while

$\mathbf{v} \leftarrow \mathbf{u} - \sum_{g \in \mathcal{G}} \boldsymbol{\xi}^g$.

Convex constraints for each vector $\boldsymbol{\xi}^g$ are separable

- Can be efficiently solved using block-coordinate ascent

Image Inpainting Results

Hierarchical dictionary learning yields much less reconstruction noise compared to flat dictionary learning with flat l1-regularization.



L1



Tree-sparsity

Noise	50%	90%
Flat	19.3	72.1
Hierarchical	18.6	65.9

Decorrelating Visual Semantic Attributes

Problem: difficult to distinguish between co-occurring attributes



Forest animal? Brown? Has ears? Combinations?

Decorrelating Visual Semantic Attributes

Motivation: learning correct attributes is crucial for applications such as image search and zero-shot learning.

silver sandals with high heels



Image search

How to identify a band-tailed pigeon

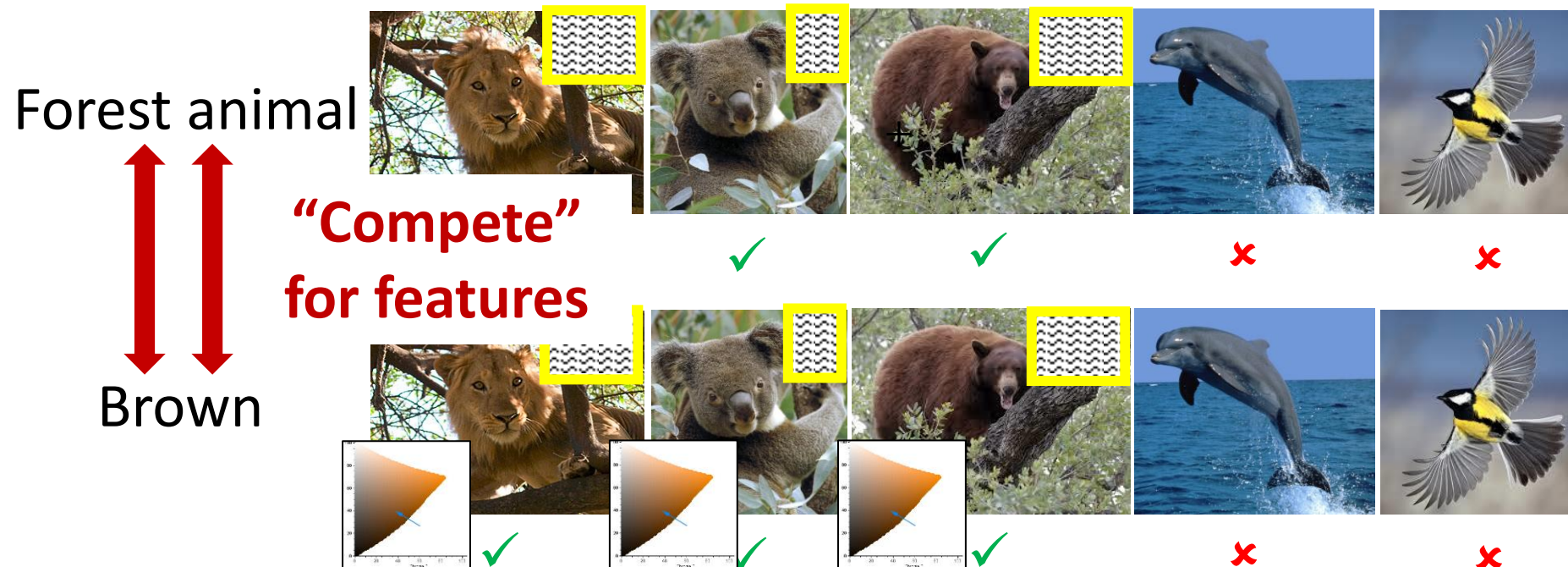
- ✓ White collar
- ✓ Yellow feet
- ✓ Yellow bill
- ✗ Red breast



Zero-shot learning

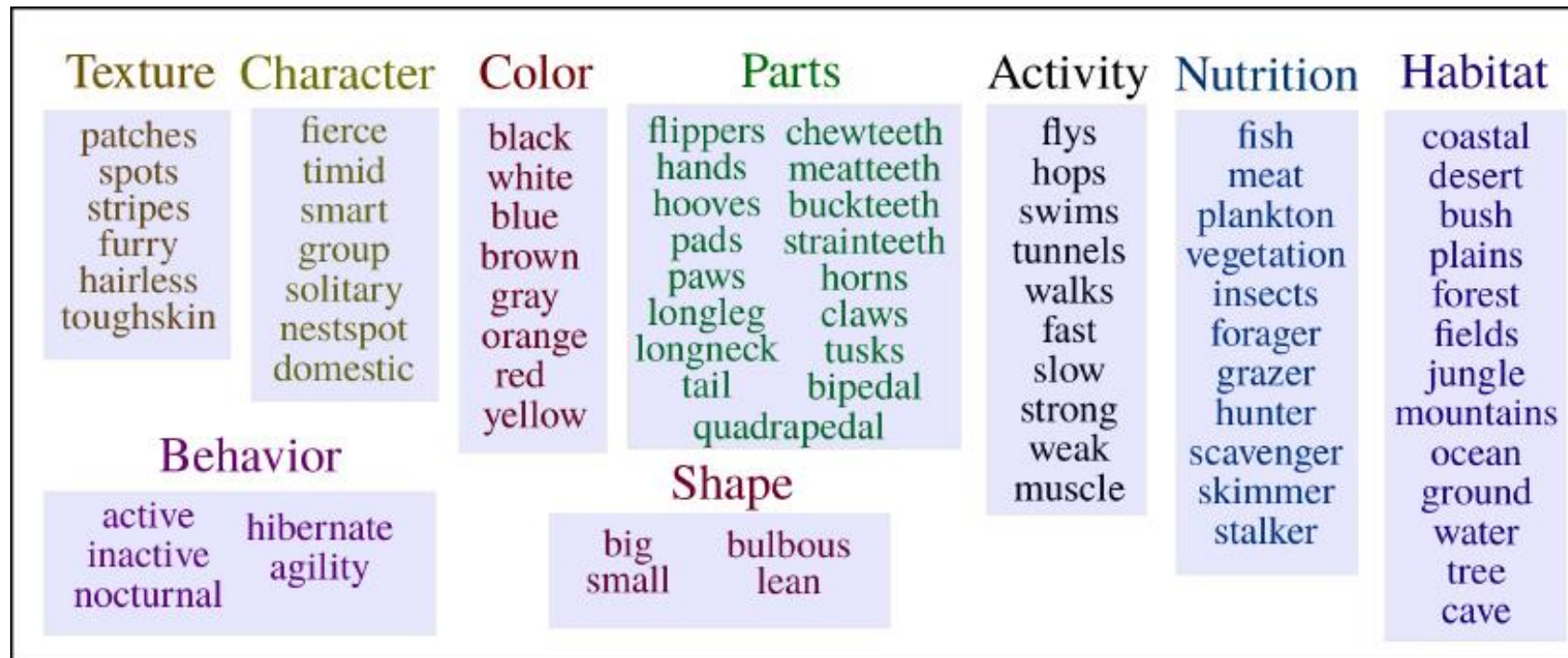
Decorrelating Visual Semantic Attributes

Solution: promote competition between



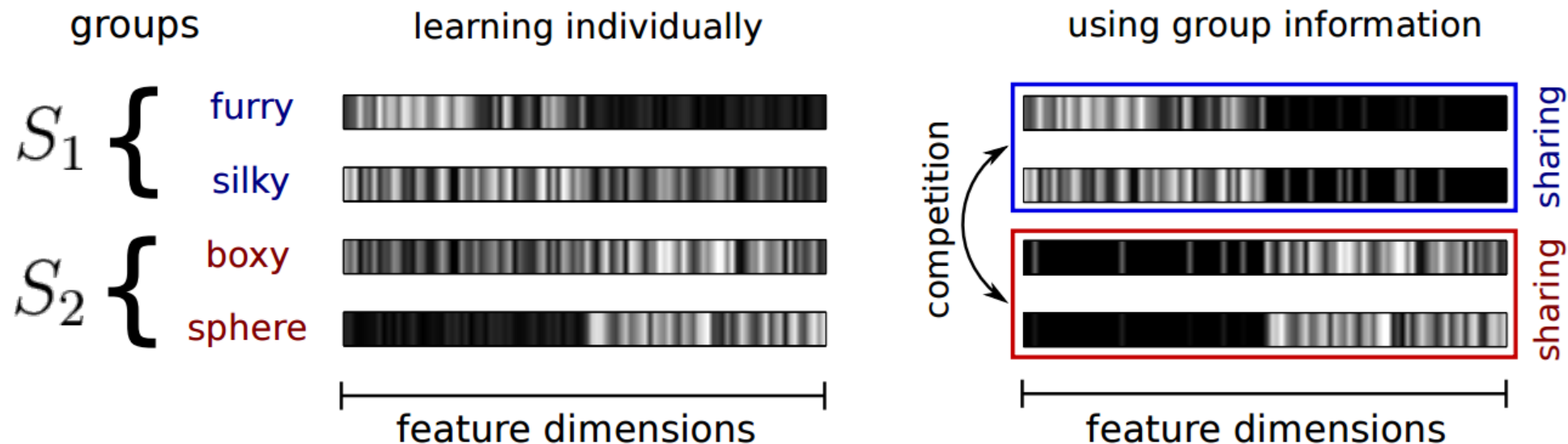
Decorrelating Visual Semantic Attributes

Attributes can be grouped into multiple semantic categories



Decorrelating Visual Semantic Attributes

Promote sharing between attributes in the same group, while promoting competition for features in different groups



$$\operatorname{argmin}_W L(W|X, Y) + \sum_d \sum_l \|w_d^{S_l}\|_2$$

Classification loss

Promotes competition for features support between different groups

Promotes feature sharing between attributes within a same group

Decorrelating Visual Semantic Attributes

Decorrelation model predicts attributes much better in unusual cases.



Not brown
underparts



No eye



Not boxy



No mouth



No ear

Decorrelating Visual Semantic Attributes

This model can also accurately localize the part-based attributes.



SPAMS (SPArse Modeling Software)

A popular optimization toolbox for solving most of the methods introduced in this presentation

<http://spams-devel.gforge.inria.fr/>

- Implements OMP, Lasso, LARS, coordinate descent and proximal methods for l_1 -regularization.
- Provides proximal toolbox for solving l_2/l_1 -reg. , sparse group lasso, tree-structured regularization, and structured sparsity with overlapping groups.

Summary

Structured sparsity is a kind of sparsity that prefers *certain structure* among the variables, based on some *prior knowledge*.

There are various types of structured sparsity, including *group sparsity*, *block sparsity*, *tree-structured sparsity*, and *graph-based sparsity*. Structured sparsity is often enforced as *mixed norm* regularization, such as (2,1)-norm

Structured sparsity is useful for multiple applications, including *multi-task learning*, and hierarchical dictionary learning for *image denoising*.