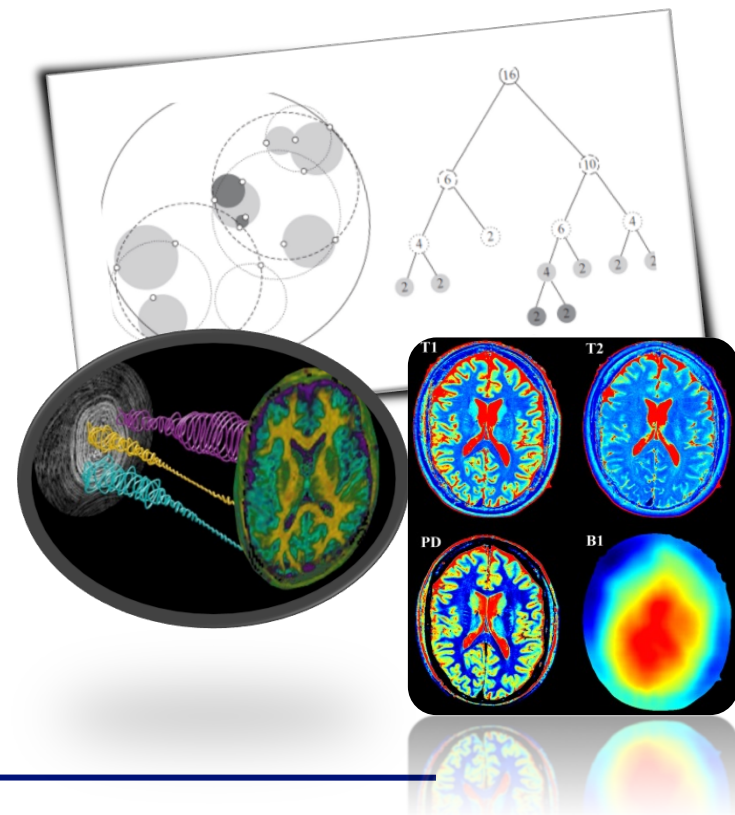


Fast Data Driven Compressed Sensing

and application to
compressed quantitative MRI

Mike Davies
& Mohammad Golbabaee

Joint work with
Zhouye Chen, Yves Wiaux
Zaid Mahbub, Ian Marshall



Outline

- Iterative Projected Gradients (IPG)
- Approximate/inexact oracles
- Robustness of inexact IPG
- Application to data driven Compressed Sensing
 - Fast MR Fingerprinting reconstruction
 - IPG with Approximate Nearest Neighbour searches
 - Cover trees for fast ANN
- Numerical results

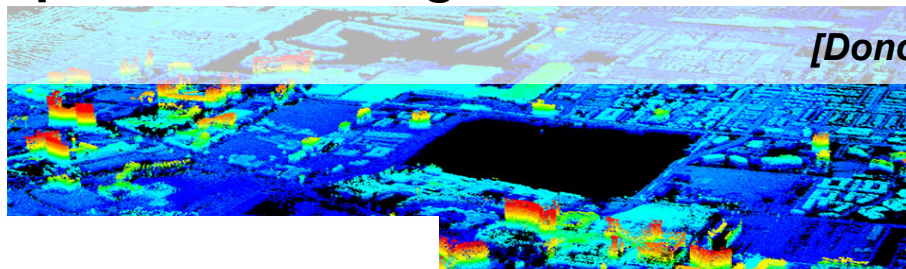
Inverse problems

$$y = Ax + w \in \mathbb{R}^m, \quad A: \mathbb{R}^n \rightarrow \mathbb{R}^m \text{ where } m \ll n$$

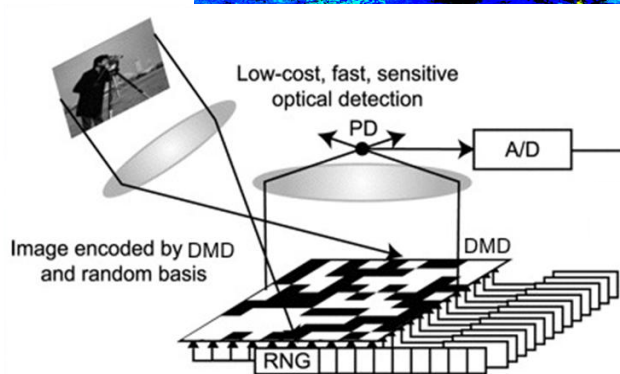
Challenge: Missing information

Complete measurements can be costly, time consuming and sometimes just impossible!

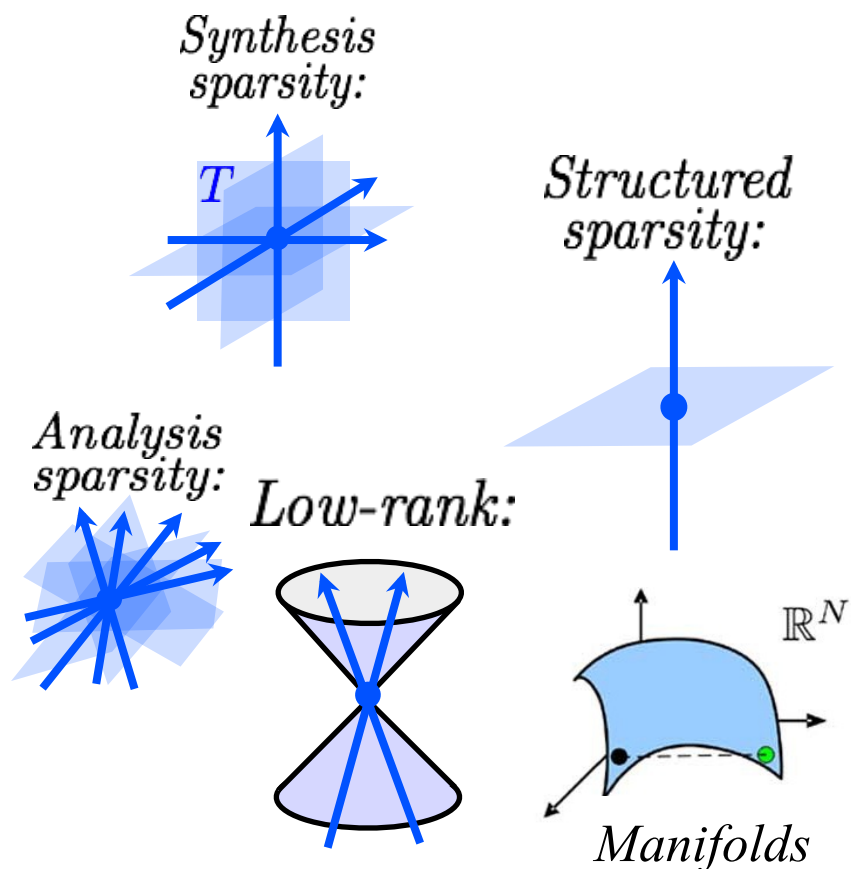
Compressed sensing to address the challenge



[Donoho'06; Candès, Romberg, and Tao, '06]

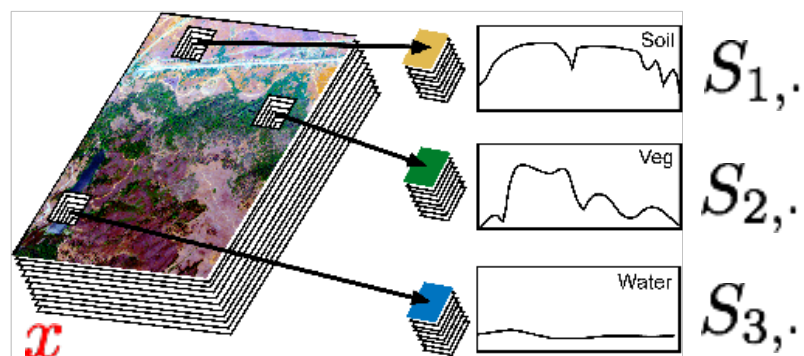
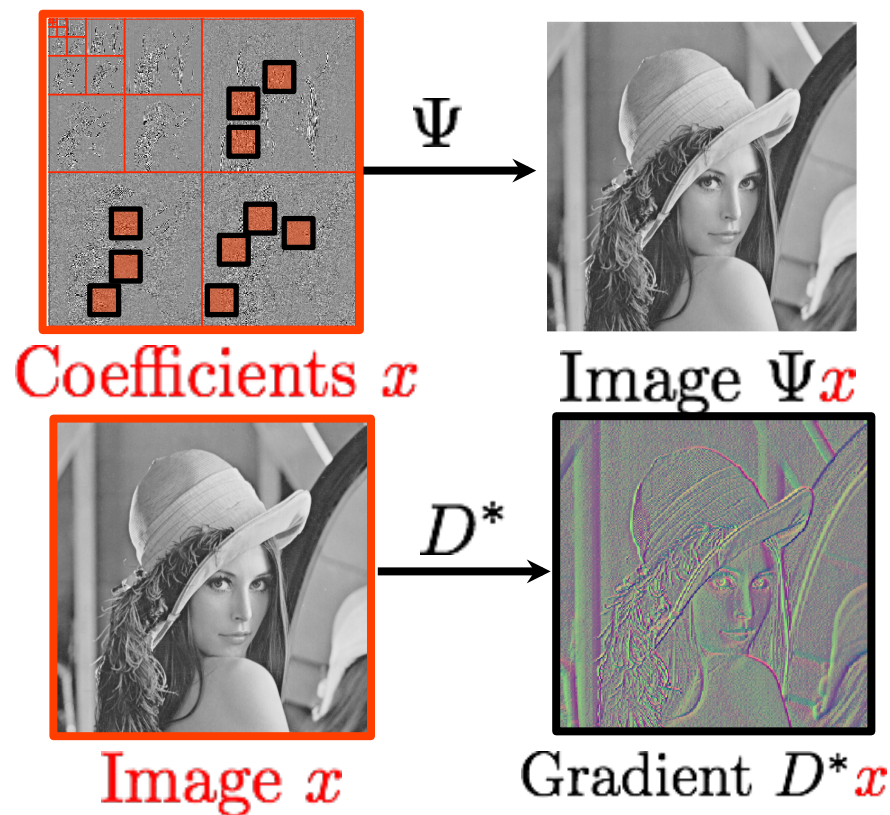


Data models/priors



Multi-spectral imaging:

$$x_{i,\cdot} = \sum_{j=1}^r A_{i,j} S_{j,\cdot}$$



Solving Compressed Sensing/Inv. Problems

Estimating by constrained least squares

$$x \in \arg\min \{ f(x) := 1/2 \|y - Ax\|_2^2 \} \quad s.t. \quad x \in C$$

!! NP-hard for most interesting models

(e.g. sparsity [Natarajan'95])

Iterative Gradient Projection (IPG)

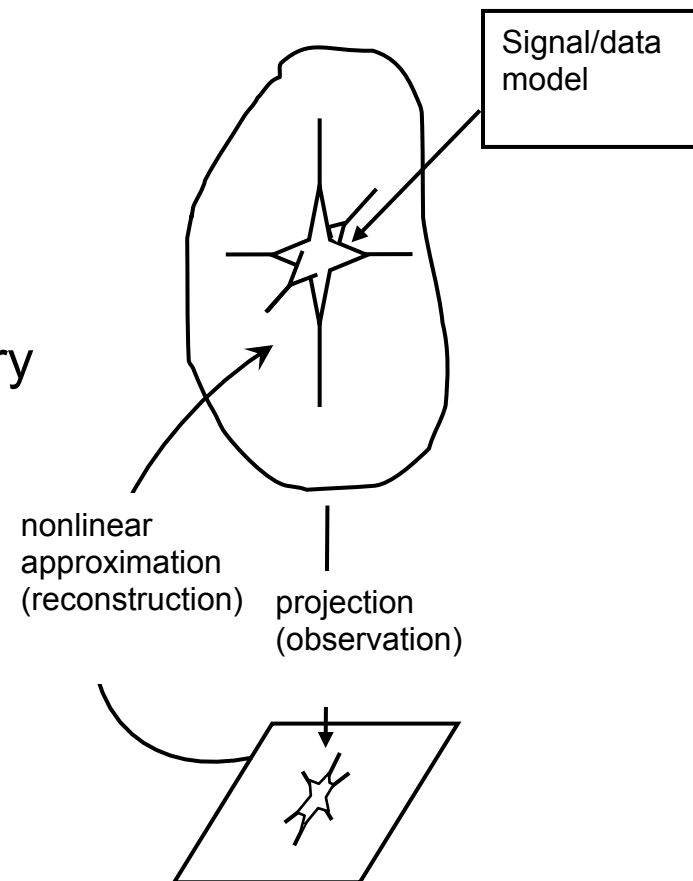
Generally proximal-gradient algorithms are very popular:

$$x^{\uparrow k} = P_C(x^{\uparrow k-1} - \mu A^{\uparrow T}(Ax^{\uparrow k-1} - y))$$

Gradient $A^{\uparrow T}(Ax^{\uparrow k-1} - y)$, step size μ ,

Euclidean projection

$$P_C(x) \in \arg\min \|x - x^{\uparrow}\|_2 \quad s.t. \quad x^{\uparrow} \in C$$



Embedding: key to CS/IPG stability

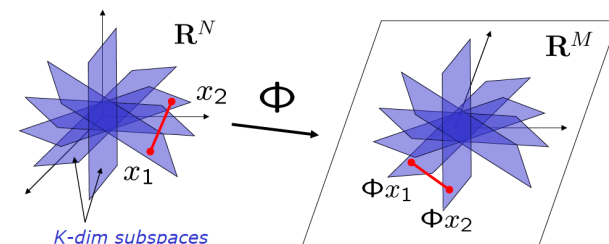
Bi-Lipschitz embeddable sets: $\forall x, x' \in \mathcal{C}$

$$\alpha \|x - x'\|_2 \leq \|A(x - x')\|_2 \leq \beta \|x - x'\|_2$$

Theorem. [Blumensath'11] For **any** $(\mathcal{C}, A)$ if holds $\beta \leq 1.5\alpha$,

IPG $\rightarrow$ stable & linear convergence: $\|x_{t+K} - x_0\| \rightarrow O(w) + \tau$, $K \sim (\log \tau^{-1} - 1)$

Global optimality even for **nonconvex** programs!



Model $\mathcal{C} \in \mathbb{R}^n$	$O(m)$	
p (unstructured) points	$\theta^{-2} \log(p)$	[Johnson, Lindenstrauss'89]
$U \in \mathbb{R}^{L \times K}$ -flats	$\theta^{-2} (K + \log(L))$	[Blumensath, Davies'09]
rank r ($\sqrt{n} \times \sqrt{n}$) matrices	$\theta^{-2} r \sqrt{n}$	[Candès, Recht; Ma et al.'09]
'smooth' d dim. manifold	$\theta^{-2} d$	[Wakin, Baraniuk'09]
K tree-sparse	$\theta^{-2} K$	[Baraniuk et al.'09]

Sample complexity e.g. $A \sim$ i.i.d. subgaussian

A limitation...

Exact oracles might be too expensive or even do not exist!

Gradient $\nabla f^T (Ax^{k-1} - y)$

- A too large to fully access or fully compute/update ∇f
- Noise in communication in distributed solvers

Projection $P_C(x) \in \arg\min_{x' \in C} \|x - x'\|_2 \text{ s.t. } x' \in C$

- P_C may not be analytic and requires solving an auxiliary optimization (e.g. inclusions $C = \bigcap_i C_i$, total variation ball, low-rank, tree-sparse,...)
- P_C might be NP hard! (e.g. analysis sparsity, low-rank tensor decomposition)

Is IPG robust against inexact/approximate oracles?

Inexact oracles I: Fixed Precision

$$x^{\uparrow k} = \mathbf{P} \downarrow \mathcal{C} (x^{\uparrow k-1} - \mu \nabla f(x^{\uparrow k-1}))$$

Fixed Precision (FP) approximate oracles:

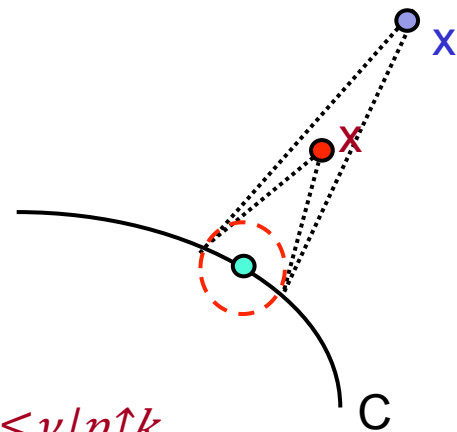
$$\|\nabla f(\cdot) - \nabla f(\cdot)\|_2 \leq \nu \downarrow g, \quad \|\mathbf{P} \downarrow \mathcal{C}(\cdot) - \mathbf{P} \downarrow \mathcal{C}(\cdot)\|_2 \leq \nu \downarrow p, \quad (\mathbf{P} \downarrow \mathcal{C}(\cdot) \in \mathcal{C})$$

Examples: TV ball, inclusions (e.g. Dijkstra alg.), and many more... (in convex settings, **Duality gap** $\rightarrow$ FP proj.)

Progressive Fixed Precision (PFP) oracles:

$$\|\nabla f(\cdot) - \nabla f(\cdot)\|_2 \leq \nu \downarrow g^{\uparrow k}, \quad \|\mathbf{P}(\cdot) - \mathbf{P}(\cdot)\|_2 \leq \nu \downarrow p^{\uparrow k}$$

Examples: Any FP oracle with progressive refinement of the approx. levels
e.g. convex sparse CUR factorization for $\nu \downarrow p^{\uparrow k} \sim O(1/k^3)$ [Schmidt et al.'11]



Inexact oracles II: $(1+\epsilon)$ -optimal

$(1+\epsilon)$ -approximate projections combined with FP/PFP gradient oracle

$$x^{k+1} = \mathbf{P}_{1+\epsilon}^{\downarrow C} (x^k - \mu \nabla f(x^k))$$

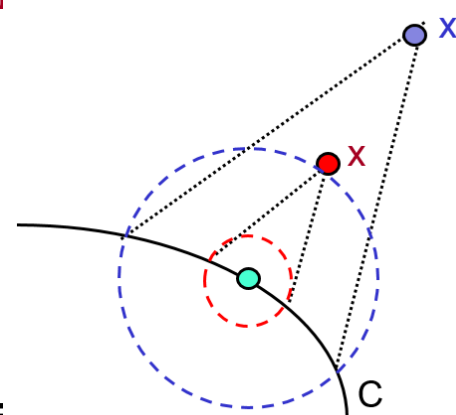
Gradient $\|\nabla f(x) - \nabla f(y)\|_2 \leq L \|x - y\|_2$

Projection $\|\mathbf{P}_{1+\epsilon}^{\downarrow C}(x) - x\|_2 \leq (1+\epsilon) \|\mathbf{P}^{\downarrow C}(x) - x\|_2$

Examples.

Many nonconvex constraints: Cheaper low-rank proximal
 randomized lin. algebra [Halko et al.'11], K-tree sparse signals [Hegde et al.'14],
 Tensor low-rank (Tucker) decomposition [Rauhut et al.'16],

... and shortly, ANN for data driven CS!



Robustness & linear convergence *of the* **inexact IPG**

IPG with (P)FP oracles

$$x^k = \mathbf{P} \downarrow C (x^{k-1} - \mu \nabla f(x^{k-1}))$$

Theorem. For any $(x^0 \in C, C, A)$ if $\beta \leq \mu^2 - 1 < 2\alpha^2$ then

$$\|x^k - x^0\| \leq \rho^k (\|x^0\| + \sum_{i=1}^k \rho^{1-i} e^i) + 2\sqrt{\beta} / (1-\rho) \alpha^2 \quad w$$

where, $\rho = \sqrt{1/\mu\alpha^2} - 1$ and $e^i = 2v \downarrow g^i / \alpha^2 + \sqrt{v \downarrow p^i} / \mu\alpha^2$

Remark:

ρ^{1-i} suppresses the early stages errors

$\Rightarrow$ use “progressive” approximations to get as good as exact!

IPG with (P)FP oracles

$$x^{\uparrow k} = \mathbf{P} \downarrow C(x^{\uparrow k-1} - \mu \nabla f(x^{\uparrow k-1}))$$

Corollary I. After $K = O(\log(\tau^{\uparrow-1}))$ iterations IPG-FP achieves

$$\|x^{\uparrow K} - x^{\downarrow 0}\| \leq O(w + v \downarrow g + \sqrt{v \downarrow p}) + \tau$$

linear convergence at rate $\rho = \sqrt{1/\mu \alpha \downarrow 0} - 1$.

Corollary II. Assume $\exists r < 1$ s.t. $e^{\uparrow i} = O(r^{\uparrow i})$, then after $K = O(\log(\tau^{\uparrow-1}))$ iterations IPG-PFP achieves

$$\|x^{\uparrow K} - x^{\downarrow 0}\| \leq O(w) + \tau$$

linear convergence at rate $\rho = \begin{cases} \max(\rho, r) & \rho \neq r \\ \rho + \xi & \rho = r \end{cases}$ (for any small $\xi > 0$)

IPG with $(1+\epsilon)$ -approximate projection

$$x^{\uparrow k} = \mathbf{P}_{C \uparrow \epsilon} (x^{\uparrow k-1} - \mu \nabla_{\downarrow k} f(x^{\uparrow k-1}))$$

Theorem. Assume for any $(x^{\downarrow 0} \in C, C, A)$ and an $\epsilon \geq 0$ it holds

$$\sqrt{2\epsilon + \epsilon^2} \leq \delta \sqrt{\alpha^{\downarrow 0}} / \|A\| \quad \text{and} \quad \beta \leq \mu^{\uparrow -1} < (2 - 2\delta + \delta^2) \alpha^{\downarrow} x^{\downarrow 0}$$

Then, $\|x^{\uparrow k} - x^{\downarrow 0}\| \leq \rho^{\uparrow k} (\|x^{\downarrow 0}\| + \kappa^{\downarrow g} \sum_{i=1}^{\uparrow k} \rho^{\uparrow - i} \nu^{\downarrow g} g^{\uparrow i}) + \kappa^{\downarrow z} / (1 - \rho)$ w

where, $\rho = \sqrt{1/\mu \alpha^{\downarrow 0}} - 1 + \delta$ $\kappa^{\downarrow g} = 2/\alpha^{\downarrow 0} + \sqrt{\mu}/\|A\| \delta$ and $\kappa^{\downarrow z} = 2\sqrt{\beta}/\alpha^{\downarrow 0} + \sqrt{\mu} \delta$

Remarks.

- Requires **stronger** embedding cond., **slower** convergence!
- Still linear conv. & $O(w) + \tau$ accuracy after $O(\log \tau^{\uparrow -1})$ iterations
- higher noise amplification

Application in data driven CS

Data driven CS

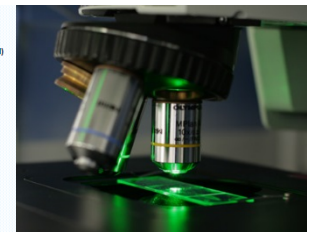
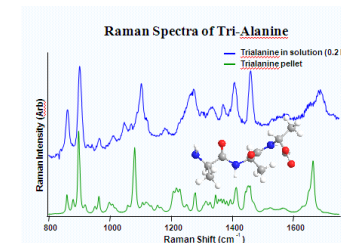
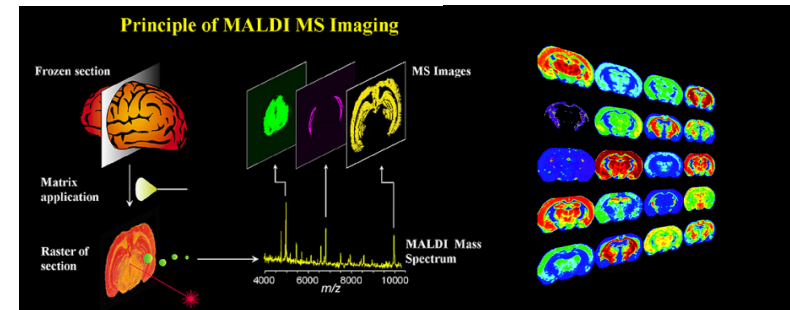
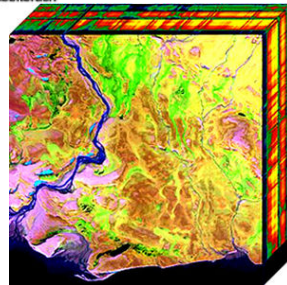
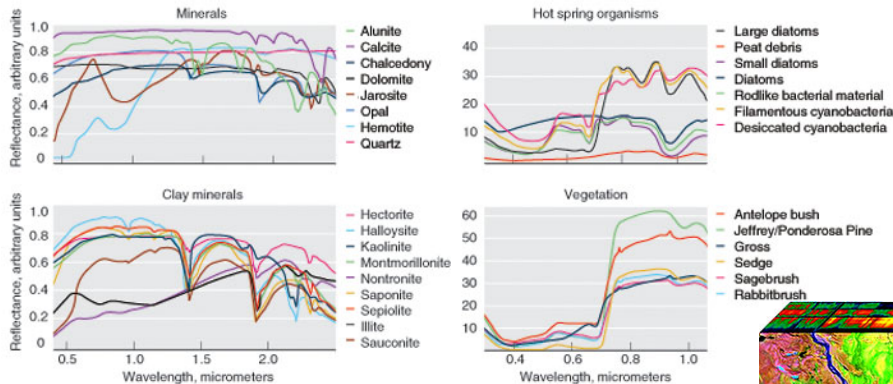
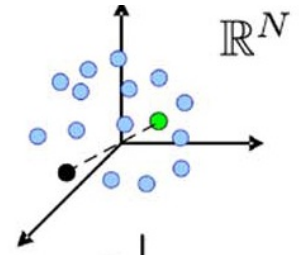
In the absence of (semi) algebraic physical models ($I_0, I_1, \text{rank}, \dots$)

Collect a possibly large dictionary (sample the model)

$$\mathcal{C} = \bigcup_{i=1, \dots, d} \{\psi_i\} \quad \psi_i \text{ atoms of } \Psi \in \mathbb{R}^{n \times d}$$

Examples in multi-dim. imaging:

Hyperspectral, Mass/Raman/MALDI, ... spectroscopy [Golbabaee et al '13; Duarte, Baraniuk '12; ...]



MR Fingerprinting: fast CQ-MRI

Goal: Measuring fast the NMR properties (**relaxation times: T1, T2**) [Ma et al.'13]

1. **Multiple** random/optimized excitations
(magnetic field rotation) [Cohen'15; Mahbub, Golbabaee, D., Marshall '16]

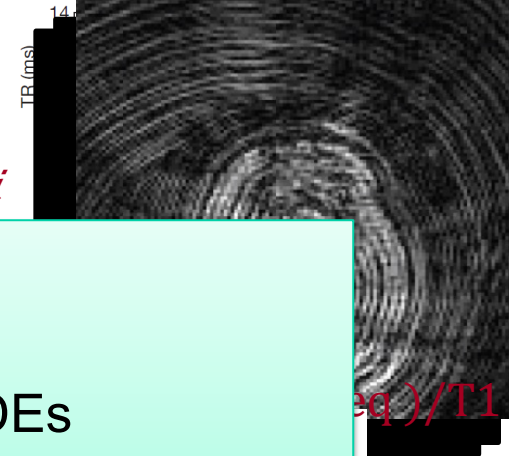
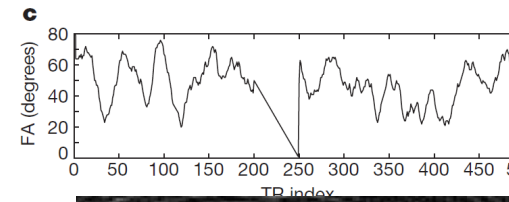
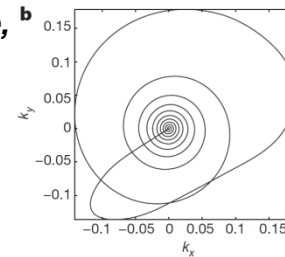
2. **Subsampling** the k-space (per excitation)

3. Construct a (huge) **dictionary** of "fingerprints"
i.e. par Continuous model (Bloch manifold) exists,

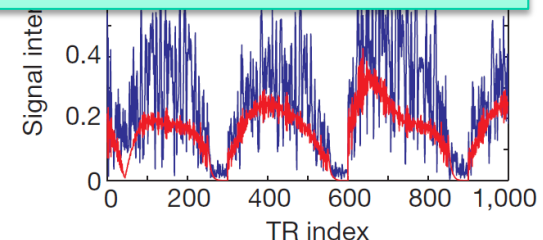
- but **not explicit** i.e. requires solving dynamic ODEs
- MRF sacrifices **memory** ~ enables **$P \downarrow C$ computation**

4. Use

- Exact IPG [D. et al '15]



$$\Psi = [B] \downarrow i$$



CS in product (tensor) space

Per pixel data driven model

$$\mathcal{C} = \bigcup_{i=1, \dots, d} \{\psi_i\} \quad \psi_i \text{ atoms of } \Psi \in \mathbb{R}^{n \times d}$$

Multi dim. image: $X \in \mathbb{R}^{n \times P}$ where $X_{\downarrow p} \in \mathcal{C} \quad \forall p=1, \dots, P$

Inverse problem: $\min_{X_{\downarrow p} \in \mathcal{C}} \|y - A(X)\|_2^2$

Direct recovery complexity $O(P \cdot d)$!

IPC (LNN Search C)

$X_{\downarrow p}$ Big datasets?

Con

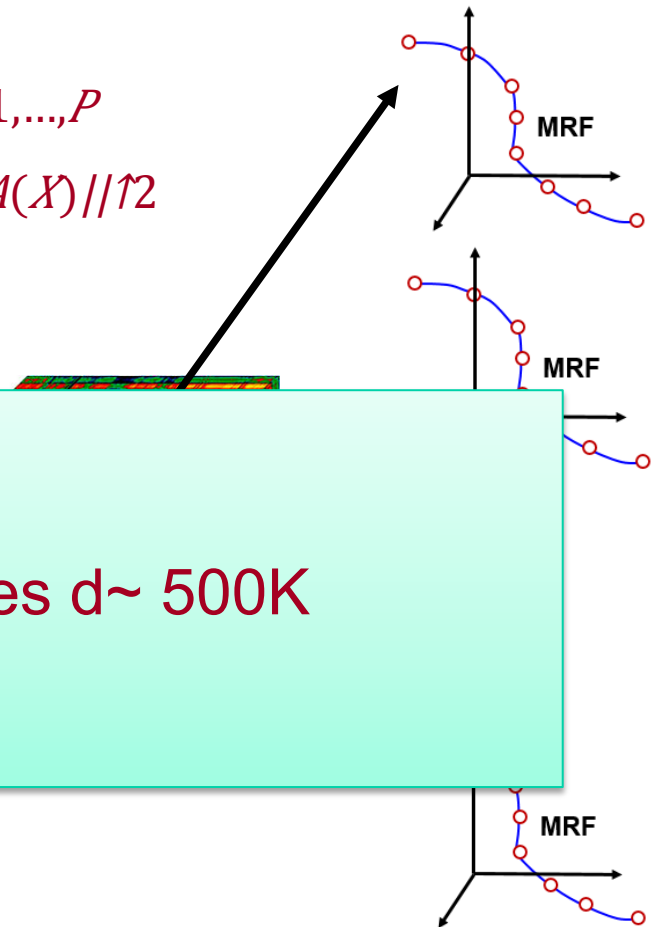
MRF uses large dictionaries $d \sim 500K$

Suff

Can we do better?

$M = O(P \cdot m)$, (ignoring inter channel structures)

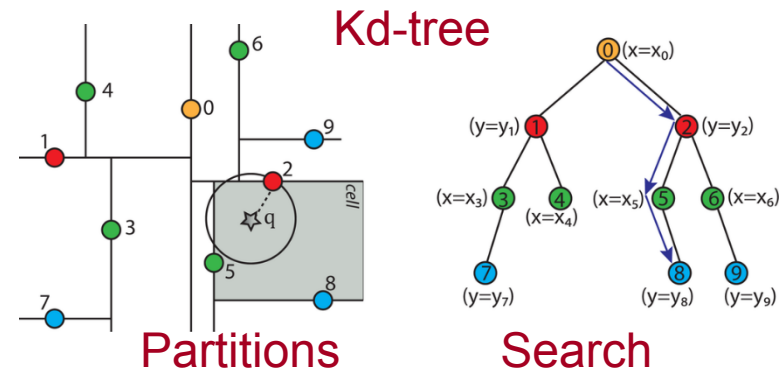
m = RIP complexity of each channel



Fast NN searches

Trees: (historical) approach to fast NN

- Hierarchical partitioning + Branch & bound search. e.g., *kd-trees*, *Metric/Ball-trees*, ...



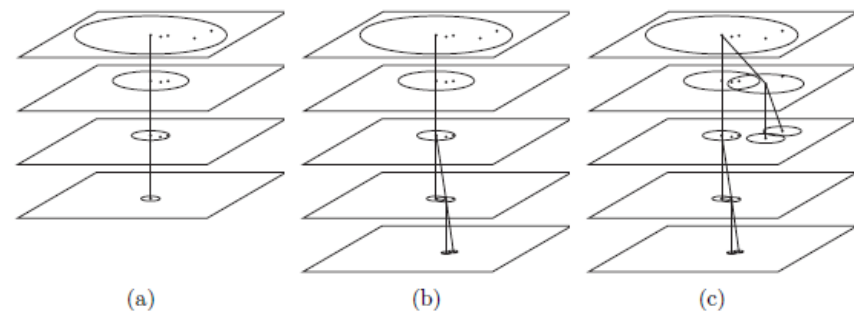
“Curse of dimensionality”: exact NN in $\mathbb{R}^n$ cannot achieve $o(nd)$ in time with a reasonable memory. Approximate NN can! If datasets $\sim$ low dim.

Navigating nets, Cover trees

[Krauthgamer, Lee'04; Beygelzimer'06]

At **scale** $l = 1, \dots, L$

- Covering (parent nodes) $\sigma 2^{l-1}$
- Separation (nodes appearing at scale l) $\sigma 2^{l-1}$



CT builds **multi-resolution** cover-nets
Brand & bound search

CT provably good for low dim data

Search options with cover trees

1. $(1+\epsilon)$ -ANN: as proposed by [Beygelzimer et al.'06]

Theorem. Cover tree $(1+\epsilon)$ -ANN complexity: [Krauthgamer + Lee 04]

$2^{\mathcal{O}(\dim(C)) \log \Delta + \epsilon^{\mathcal{O}(\dim(C))}}$ in time, $\mathcal{O}(d)$ space

(typically $\log \Delta = \mathcal{O}(\log(d))$)

2. FP-ANN: truncated tree $L = \lceil \log(\nu \downarrow p / \sigma) \rceil +$ exact NN

(complexity could be arbitrary large/theoretically)

3. PFP-ANN: progressing on truncation level $\nu \downarrow p \uparrow k = \mathcal{O}(r \uparrow k)$, $r < 1$

Inexact IPG

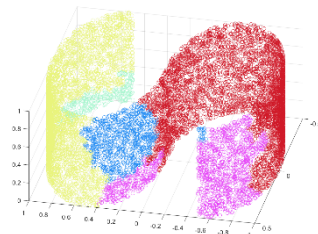
$$X \downarrow p \uparrow k = \text{ANN} \downarrow C ([X \uparrow k - 1 - \mu A \uparrow H(A(X \uparrow k - 1) - y)] \downarrow p), \forall p$$

Numerical experiments 1: Toy problem

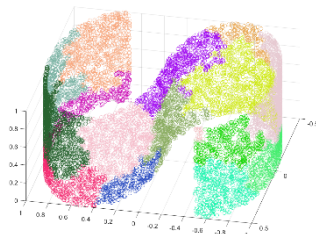
2D manifold data

$$\mathcal{C} = \bigcup_{i=1, \dots, d} \{\psi_i\} \quad \psi_i \text{ atoms } \Psi \in \mathbb{R}^{n \times d}$$

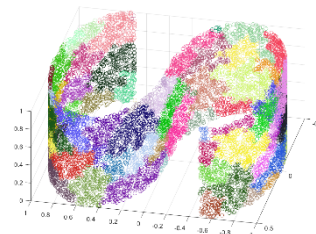
CT level 2



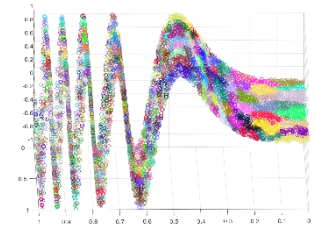
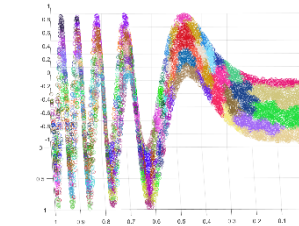
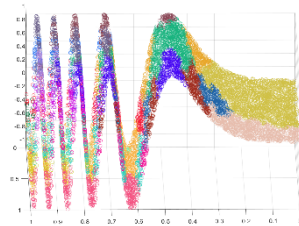
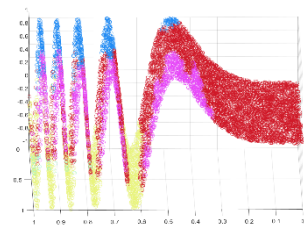
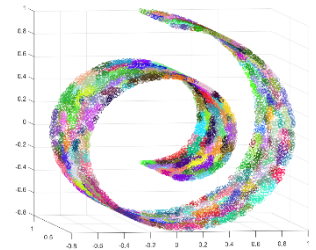
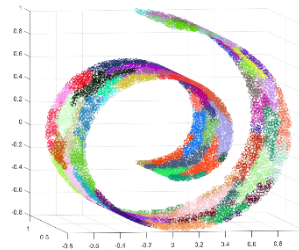
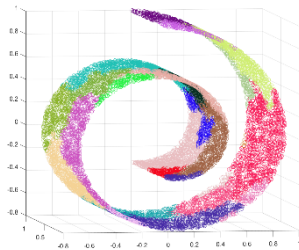
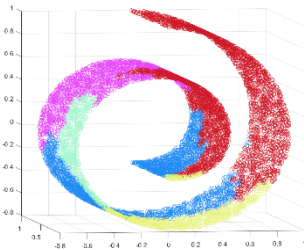
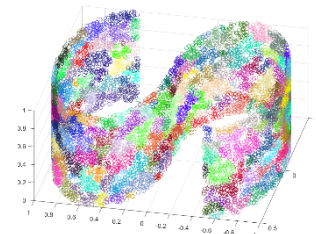
CT level 3



CT level 4



CT level 5



Dataset	Population	Ambient dim. (N)	CT depth	CT res.
S-Manifold	5'000	200	14	2.43E-4
Swissroll	5'000	200	14	1.70E-4
Oscillating wave	5'000	200	14	1.86E-4

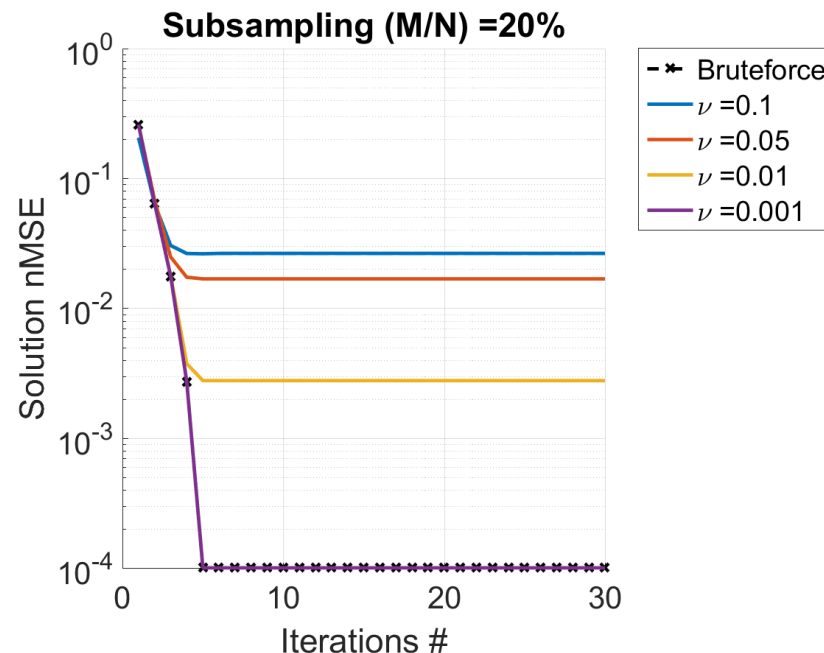
Solution accuracy vs. Iterations (FP)

Signal: $X \in \mathbb{R}^{n \times P}$ $n=200$, $P=50$ (randomly chosen $\in \mathbb{C}$)

$m \times P \times n \times P$ i.i.d. Normal A , CS ratio = m/n (noiseless)

$$\min_{X \downarrow \text{vec} \in \prod_{j=1}^P C} \|y - A(X)\|_2^2 \Leftrightarrow X \downarrow p \uparrow k = \text{ANN} \downarrow C ([X \uparrow k - 1 - \mu A \uparrow H(A(X \uparrow k - 1) - y)] \downarrow p), \forall p$$

Swiss roll



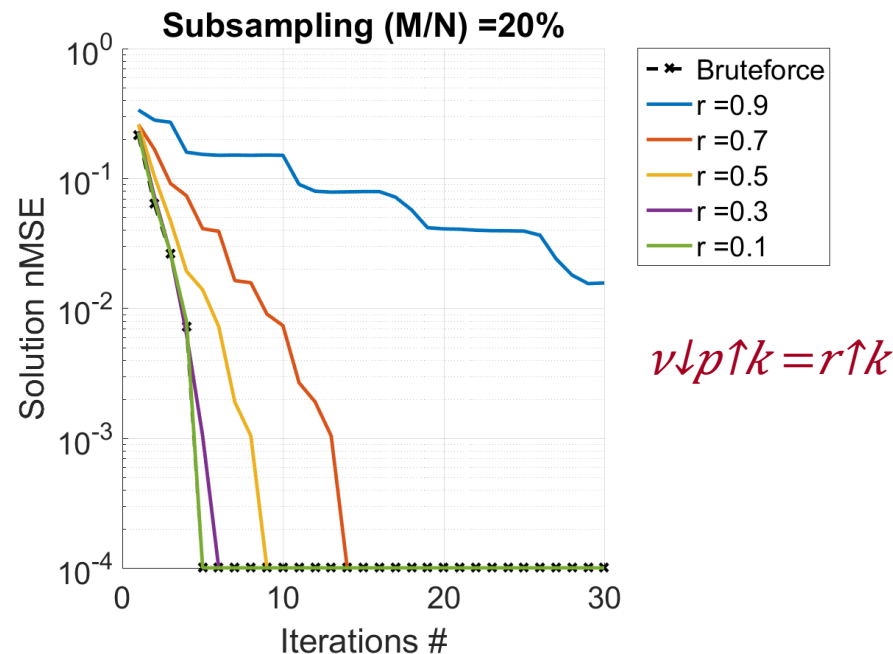
Solution accuracy vs. Iterations (PFP)

Signal: $X \in \mathbb{R}^{n \times P}$ $n=200$, $P=50$ (randomly chosen $\in \mathbb{C}$)

$m \times P \times n \times P$ i.i.d. Normal A , CS ratio $= m/n$ (noiseless)

$$\min_{\tau} X \downarrow_{\text{vec}} \in \prod_{j=1}^P C \quad ||y - A(X)||_2^2 \Leftrightarrow X \downarrow_{p \uparrow k} = \text{ANN} \downarrow_C ([X \uparrow_{k-1} - \mu A \uparrow H(A(X \uparrow_k - 1) - y)] \downarrow_p), \forall p$$

Swiss roll



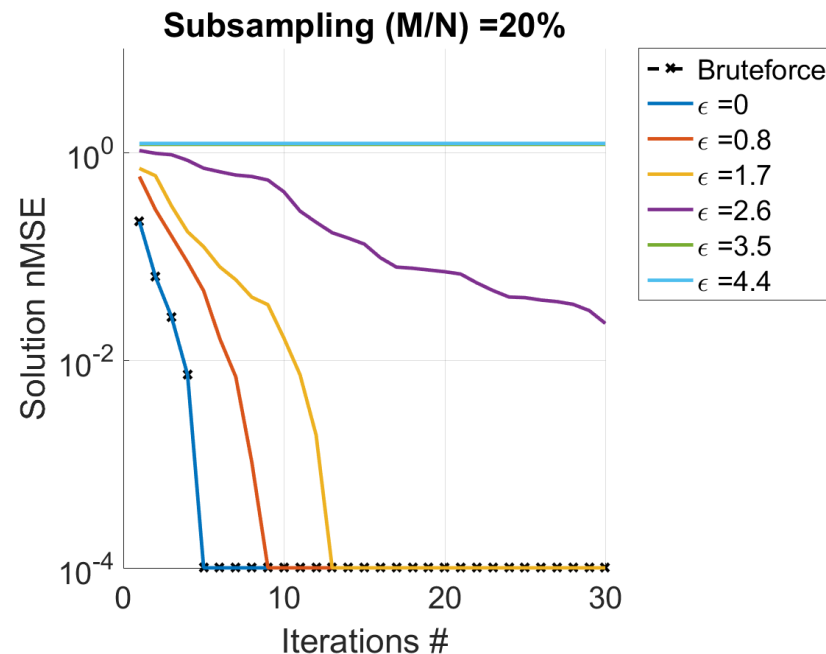
Solution accuracy vs. Iterations $(1+\epsilon)$ -ANN

Signal: $X \in \mathbb{R}^{n \times P}$ $n=200$, $P=50$ (randomly chosen $\in \mathbb{C}$)

$m \times P \times n \times P$ i.i.d. Normal A , CS ratio $= m/n$ (noiseless)

$$\min_{\tau} X \downarrow_{vec} \in \prod_{j=1}^P C \quad ||y - A(X)||^2 \Leftrightarrow X \downarrow_{p \uparrow k} = ANN \downarrow_C ([X \uparrow_{k-1} - \mu A \uparrow H(A(X \uparrow_k - 1) - y)] \downarrow_p), \forall p$$

Swiss roll

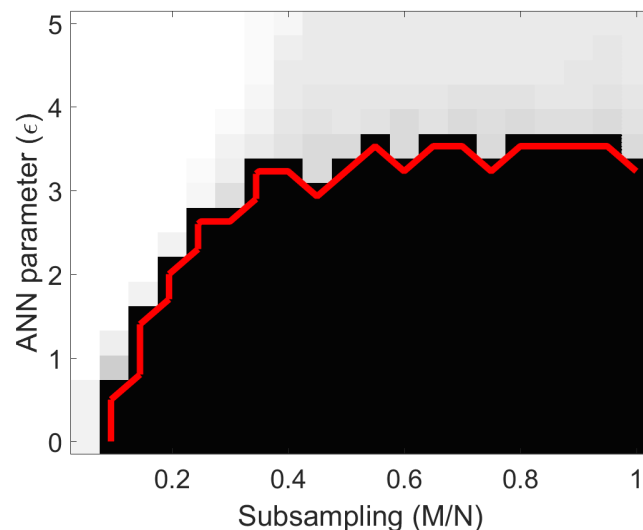


Phase transitions

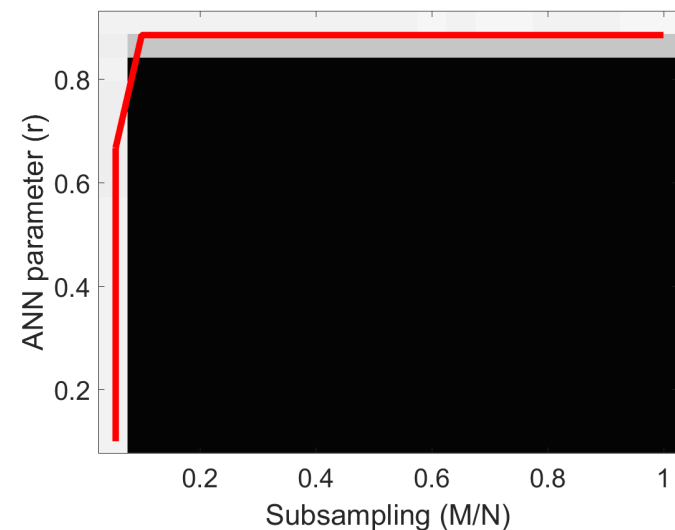
Signal: $X \in \mathbb{R}^{n \times P}$ $n=200$, $P=50$ (randomly chosen $\in \mathbb{C}$)

$m \times n$ i.i.d. Normal A , CS ratio = m/n (noiseless) ~ averaged 25 trials

Recovery PT: Black/white = low/high sol. nMSE, red curve = recovery region nMSE < $10e-4$)



$(1+\epsilon)$ -ANN IPG



PFP-ANN IPG $v \downarrow p \uparrow k = r \uparrow k$

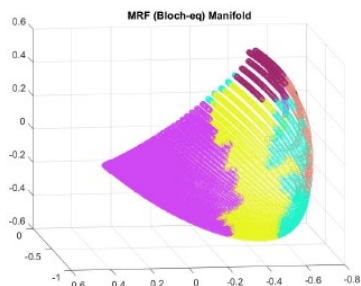
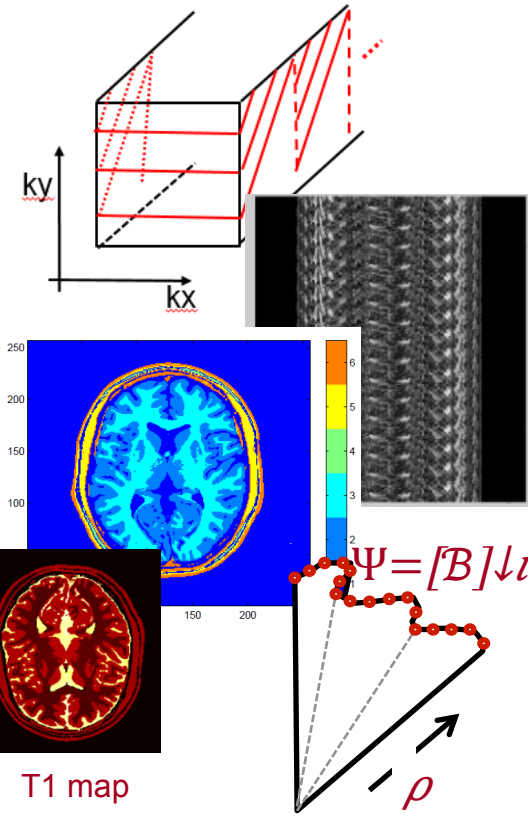
Numerical experiments 2: MRF

MRF cone & EPI acquisition

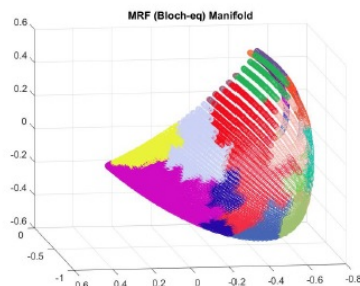
- K-space subsampling: Echo Planar Imaging (EPI)
- Anatomical phantom {Grey/White matters, CSF, muscle, skin}
- Bloch eq. dictionary $\Psi \in \mathbb{C}^{1512 \times \sim 50'000}$

$$\min_{\tau} \|y - A(X)\|_2^2 \quad s.t. \quad X \downarrow vec \in \prod_j \uparrow \text{cone}(\Psi)$$

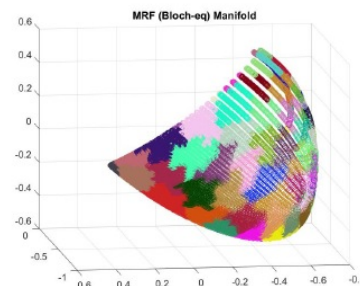
1. $a \downarrow p = X \downarrow p \uparrow k - 1 - \mu A \uparrow H (A(X \downarrow p \uparrow k - 1) - y)$
2. $\psi \downarrow p = ANN \downarrow normalised\{\Psi\} (a \downarrow p \uparrow k / \|a \downarrow p \uparrow k\|)$
3. $\rho \downarrow p \uparrow k = a \downarrow p / \psi \downarrow p$
4. $X \downarrow p \uparrow k = \rho \downarrow p \psi \downarrow p, \forall p$



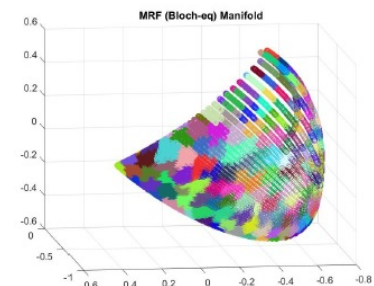
(a) Cover tree segments at scale 2



(b) Cover tree segments at scale 3



(c) Cover tree segments at scale 4

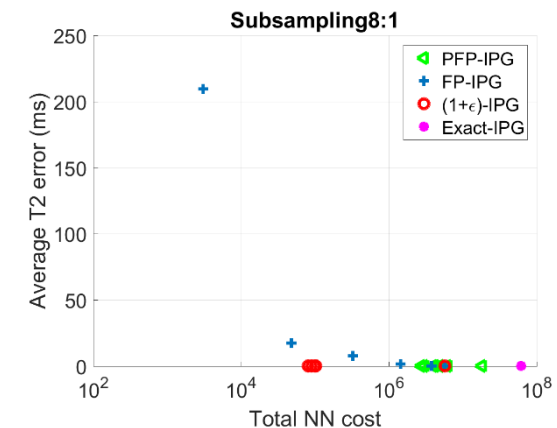
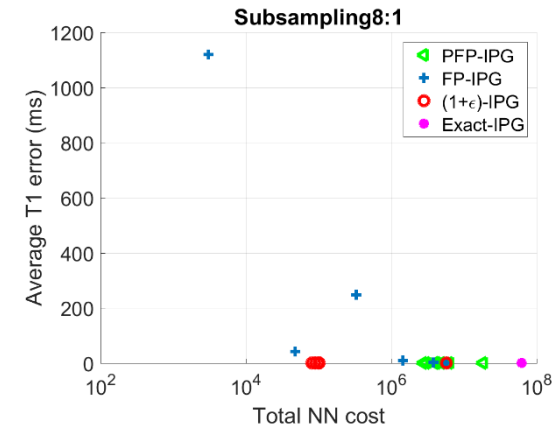
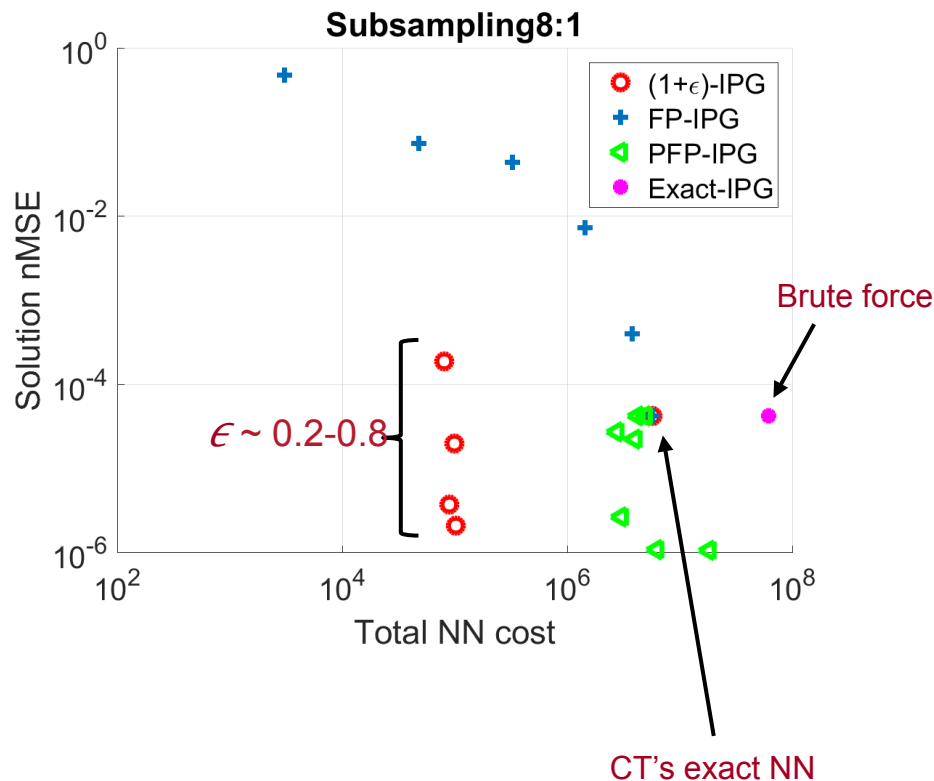


(d) Cover tree segments at scale 5

Accuracy vs. computation

Dominant cost $\rightarrow$ NN/ANN (since $\mathcal{A}$ is FFT)

Projection cost = # matches calculated (i.e. visited nodes on the tree)



Summary

- IPG robustness to inexact oracles (under embedding assumption)
- Linear convergence result:
 - PFP/ $(1+\epsilon)$ -oracles: same final accuracy vs. exact IPG
 - PFP: same convergence rate vs. exact IPG
 - $(1+\epsilon)$: stronger assumptions/sensitive to conditioning of A
- Implications in data driven CS (using ANN)
 - Cover trees for fast ANN: complexity $\sim$ intrinsic dim(data)
- $O(10e3)$ faster parameter estimation in MRF

Thnx!