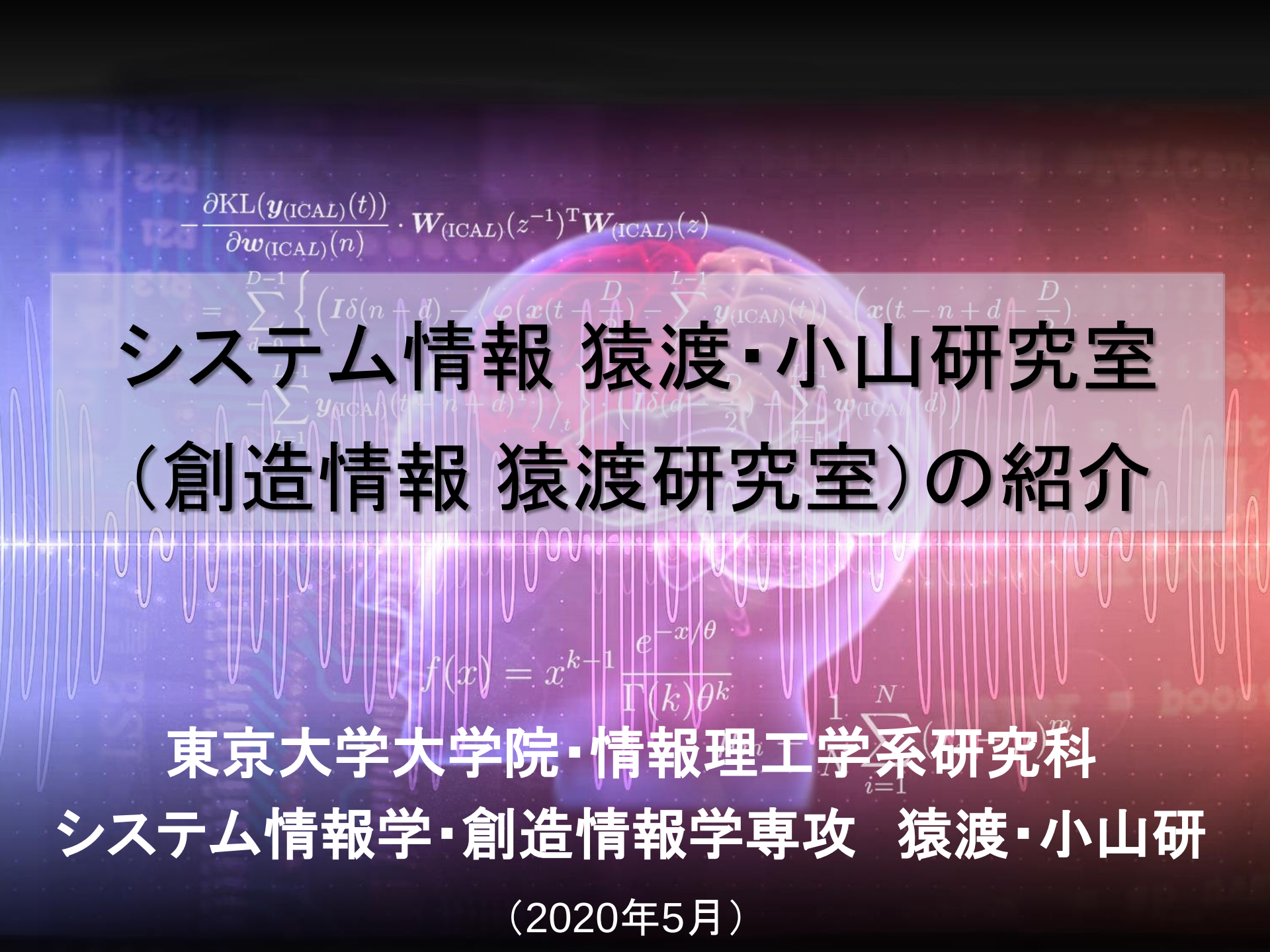




$$-\frac{\partial \text{KL}(\mathbf{y}_{(\text{ICAL})}(t))}{\partial \mathbf{w}_{(\text{ICAL})}(n)} \cdot \mathbf{W}_{(\text{ICAL})}(z^{-1})^T \mathbf{W}_{(\text{ICAL})}(z)$$

$$= \sum_{d=0}^{D-1} \left\{ \left(I\delta(n+d) - \left\langle \varphi\left(x(t-\frac{D}{2}) - \sum_{l=1}^{L-1} \mathbf{y}_{(\text{ICAL})}(t-l)\right) \cdot \left(x(t-n+d-\frac{D}{2}) - \sum_{l=1}^{L-1} \mathbf{y}_{(\text{ICAL})}(t-l-n+d)\right) \right\rangle_t \right) \cdot \left(I\delta(d-\frac{D}{2}) + \sum_{l=1}^{L-1} \mathbf{w}_{(\text{ICAL})}(d) \right) \right\}$$

システム情報 猿渡・小山研究室 (創造情報 猿渡研究室)の紹介

$$f(x) = x^{k-1} \frac{e^{-x/\theta}}{\Gamma(k)\theta^k}$$

東京大学大学院・情報理工学系研究科

システム情報学・創造情報学専攻 猿渡・小山研

(2020年5月)

猿渡・小山研究室(創造情報 猿渡研)

猿渡洋(教授)



専門分野

- ・教師無し最適化
- ・統計・機械学習
- 論的信号処理
- ・音場再生・伝送
(音ホログラフ)
- ・スパース最適化

小山翔一(講師)



専門分野

- ・音場再生・伝送
(音ホログラフ)
- ・スパース最適化

高道慎之介(助教)



専門分野

- ・統計的音声合成
- ・声質変換
- ・深層学習(DNN)

中村友彦(特任助教)



専門分野

- ・音楽情報処理
- ・多重解像度解析
- ・DNN音源分離

郡山知樹(助教)



専門分野

- ・統計的音声合成
- ・DNN
- ・深層ガウス過程

協力教員 北村大地先生
特任研究員 高宗さん
秘書 丹治さん

博士課程学生4名
修士課程学生9+9名



研究俯瞰図

- 音声・音響・音楽メディアに関する信号処理・情報処理
- ヒューマンインターフェイス・コミュニケーションシステムの構築
- 統計的・機械学習論的信号処理、数理最適化問題等を研究

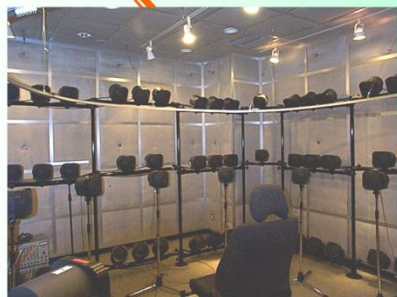
教師無し学習に基づく
ブラインド音源分離

多チャンネル信号処理



高次統計量
に基づく聴覚
印象定量化

逆フィルタ
音場再現



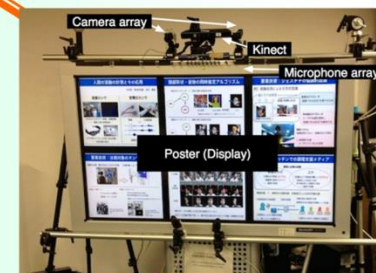
音拡張現実感
聴覚補助システム

音VRに基づく
バージンフリー
音声対話



実環境における
頑健な音声認識

ロボット
音声対話システム



マルチモーダル
ヒューマン
インターフェイス

スパース表現に
よる信号分解

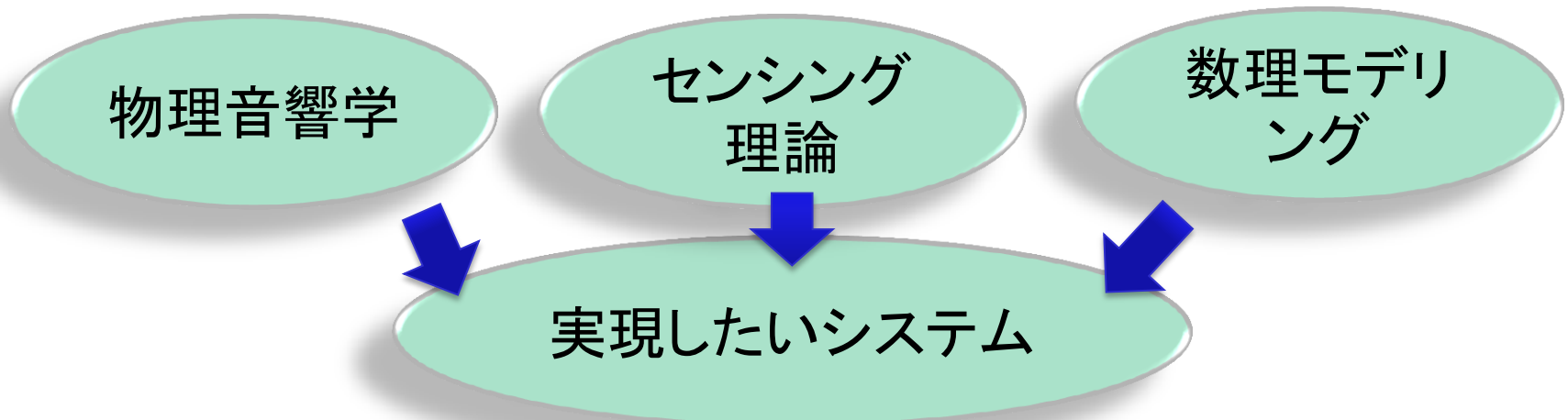
音像制御

音楽サムネイル
自動生成



音メディアに関する信号処理研究の魅力

- 音メディアに関する信号処理研究の魅力とは？
 - 自然界の音が持つ無限の**多様性**(cf. 無線通信信号)
 - 研究のアプローチに**多面性**あり(決定論的？統計的？)
 - 最後は聴かせてなんぼの評価 ⇒ **芸術性**も併せ持つ
- 「**物理世界(波動)と情報世界(抽象)をまたぐ学問**」であり、かつそれを「**統一的に取り扱うシステム工学**」である。
- 対象の多様性ゆえに「**なんでもあり**」の分野でもある。



音メディアに関する信号処理研究の魅力

波動方程式

室内音響

伝達関数

音生成過程 etc.

離散サンプリング

フーリエ解析

球面調和解析

圧縮センシング etc.

統計モデリング

最尤・ベイズ推定

機械学習

スパース最適化

物理音響学

センシング
理論

数理モデリ
ング

実現したいシステム



ビッグデータ vs. スモールデータ？

時代はビッグデータ！
しかし実際の音波動データは
簡単に集められるのか？



むしろ「**スモールデータ(一期一会データ)**」の研究が必要
ではないか？

全てのセンサはネットワーク
に繋がり「トリリオン(一兆個)
センシング」も登場。



しかし**ばらまかれた音センサ**
は位相情報を失っているので
活用できない！

ビッグデータ vs. スモールデータ？

時代はビッグデータ！
しかし実際の音波動データは

全てのセンサはネットワーク
に繋がリ「トリリオン（一兆個）

- 弊研ではこれら**両方のスケールを意識**した新しい情報処理・機械学習の枠組みを提案します。
- 単に「データ（観測）⇒抽象化（モデリング・認識）」の一方方向ではなく、「データ⇒モデリング・認識⇒物理波動の**芸術的**再構築」まで見据えたメディア処理を提案します。

むしろ「スモールデータ（一期一会データ）」の研究が必要ではないか？

しかし「はらまかれた音」センサは位相情報を失っているので活用できない！

研究紹介1. ブラインド音源分離

- 音の方向・声質・音量など、事前に何も分かっていなくても、瞬時に音を「聞き分ける」ことの出来るシステムを目指す。
- 独自に開発した高速独立成分分析(ICA)、独立低ランク行列分析(ILRMA)という教師無し数理最適化アルゴリズムに基づいて、音を統計的に独立な成分に分解することにより、別々の音声信号を見つける。

スモールデータ
教師無し最適化
低ランクモデル

音の教師無し分離(AI聖徳太子を造る)

■カクテルパーティー効果

- 人の聞き分け能力の模擬
- 補聴器への応用



■音声インターフェイス

- マイクロホンと口の間の距離が大きくなるにつれて増大してくる妨害音を抑圧・除去
- 対話IFやロボットへの応用
- 音監視システム・災害救助ロボット
(内閣府プロジェクト; プレスリリース済み)



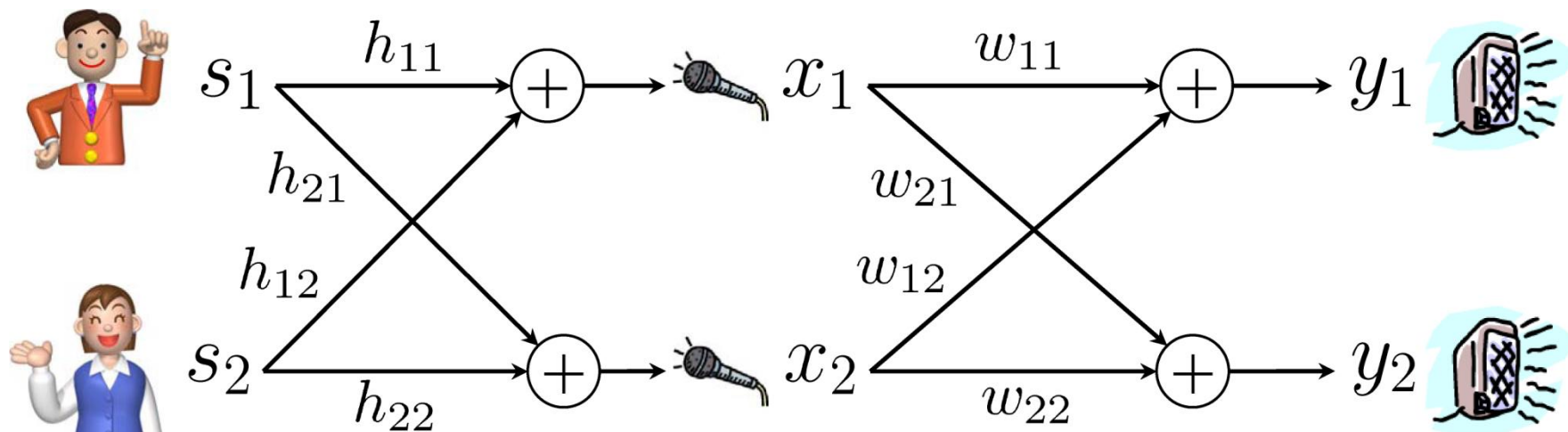
■音楽／楽器音分析

- ミックスダウン録音の解析
- 自動採譜、音楽情報処理

節々にセンサを配置⇒位置不定

ブラインド音源分離(B_lindS_ourceS_eparation)

- 混ざり合った信号 x_1, x_2 から元の信号を取り出す
- どのように混ざったかに関する情報 H は利用できない
- 事前トレーニング出来ない⇒ビッグデータではなく**スモールデータ**



実は上記は**2つのことを同時に推定**している

- [空間] 統計的に独立な音源の分類問題(分離行列 W の推定)
- [音源] 各音源が属する確率分布 $p(y)$ や構造の推定問題

上記を閉形式で解く方法は存在せず凸問題でもない⇒**大変困難!**

ILRMA: 音源の独立性と低ランク性に着目したBSS

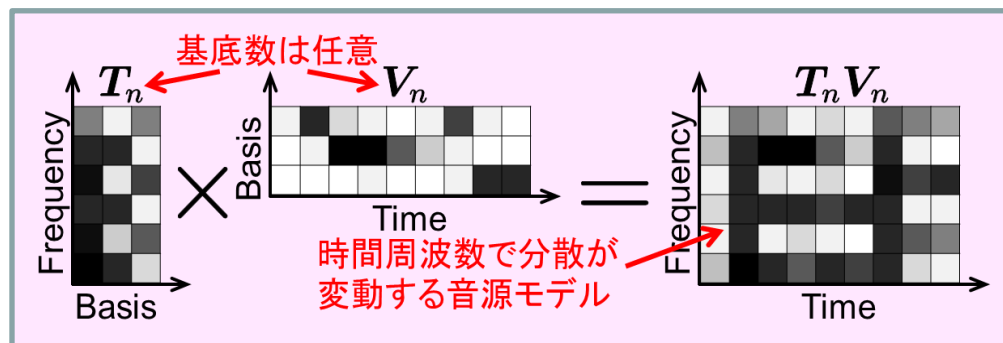
[IEEE Trans. ASLP 2016、IEEE SPJP論文賞・ASJ粟屋賞・JSPS育志賞]

- ILRMAのコスト(対数尤度)関数→これを最小化

$$\mathcal{J} = \sum_{i,j} \left[\underbrace{\sum_m \log \sum_k z_{mk} t_{ik} v_{kj}}_{\text{音源の低ランク性コスト関数}} + \underbrace{\sum_m \frac{|y_{ij,m}|^2}{\sum_k z_{mk} t_{ik} v_{kj}}}_{\text{音源の独立性コスト関数}} - 2 \log |\det \mathbf{W}_i| \right]$$

音源の低ランク性コスト関数
(音源NMFモデルの推定に寄与)

音源の独立性コスト関数
(空間モデル $\mathbf{W}$ の推定に寄与)



$$p(\mathbf{y}) = p(y_1) p(y_2) \dots p(y_m)$$

となる $\mathbf{W}$ を推定

両者を交互にMajorization-Minimizationアルゴリズムで反復最小化

- ✓ コスト値の単調減少性を保証(勾配法には無い特徴)
- ✓ 高速かつ安定な求解法を実現(従来の多入力NMFと比較して2ケタ速い)

モデルの多様化・数理解法の開拓

● 音源生成モデルの多様化 IEEE SPS Tokyo Joint Chapter 学生賞!

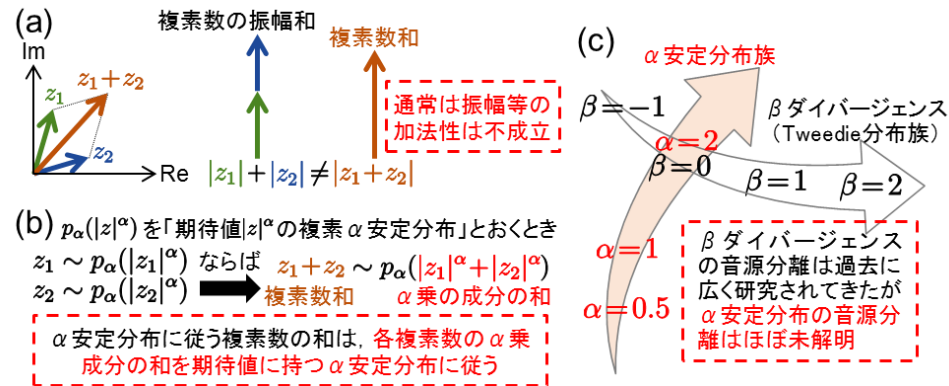
- 複素波形重ね合わせと整合する α 安定分布の導入

⇒ **t-ILRMA** [EURASIP-JASP2018]

- 複素球状ポアソン分布の導入

⇒ **$\beta=1$ -divergence最小化ILRMA**

[IEICE Trans. 2018]

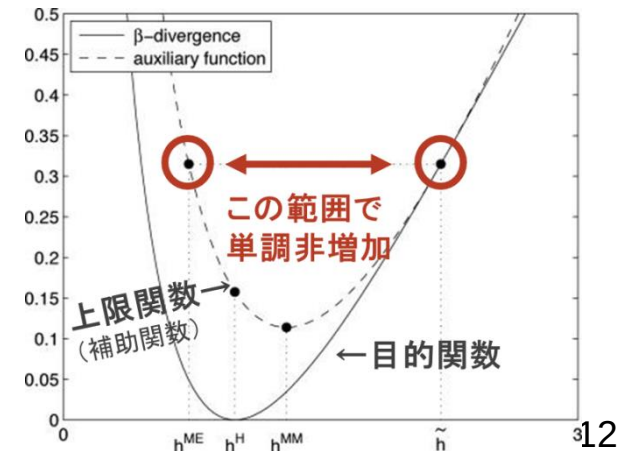


● 座標降下法におけるバリエーション

- **パラメトリックMajorization-Equalizationアルゴリズム**による音源・空間最適化の「バランス」化 [CAMSAP2017]

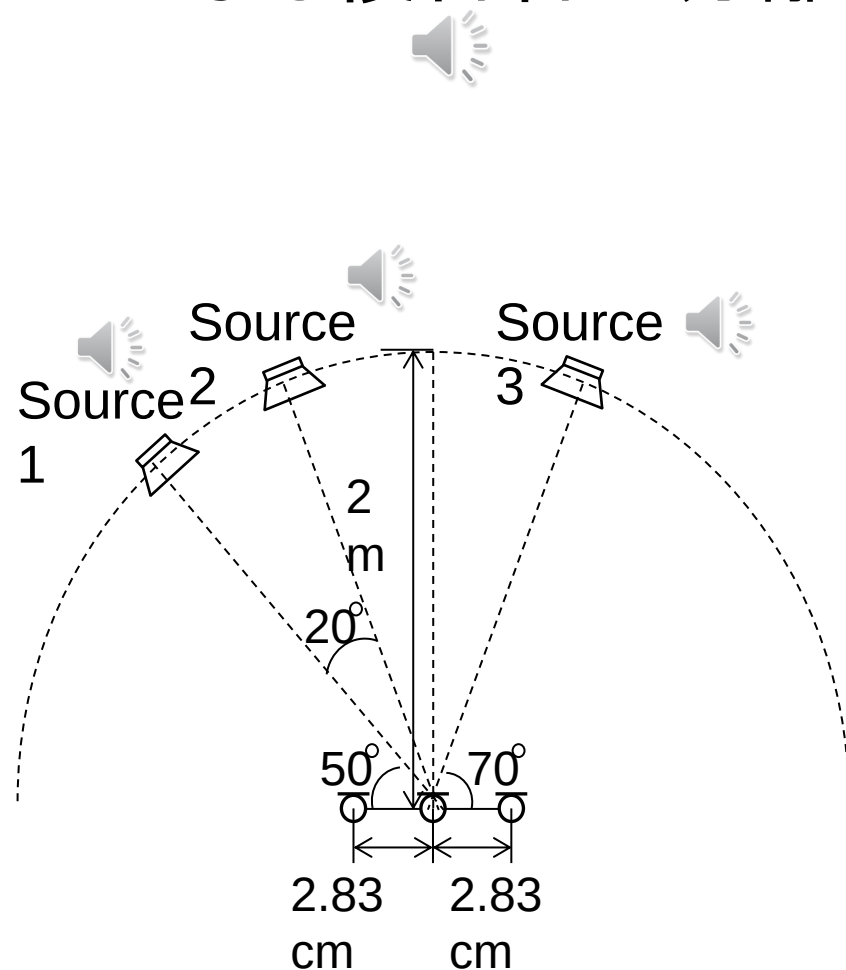
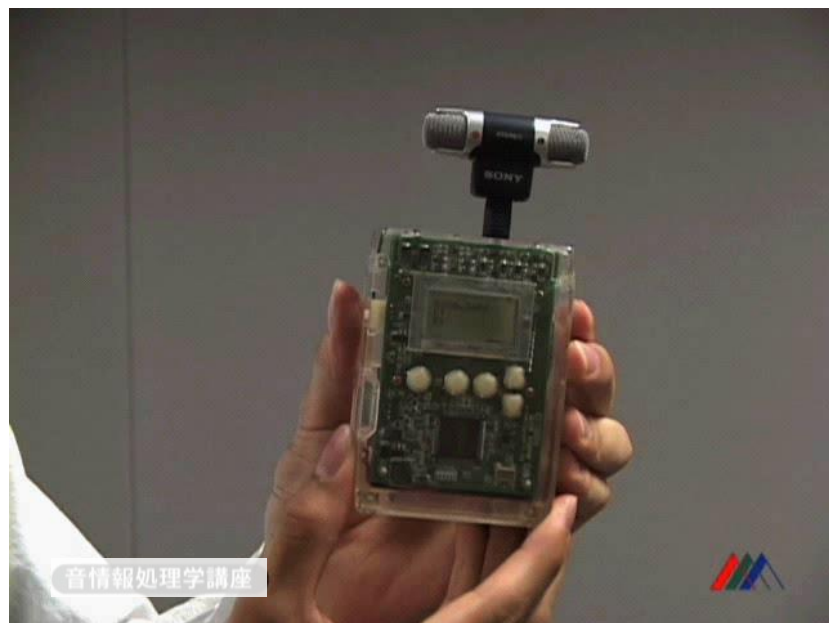
● 識別的な音源NMFモデルの追求

- 二段階最適化を直接解くことの出来るNMF [IEEE SP Letters 2019]



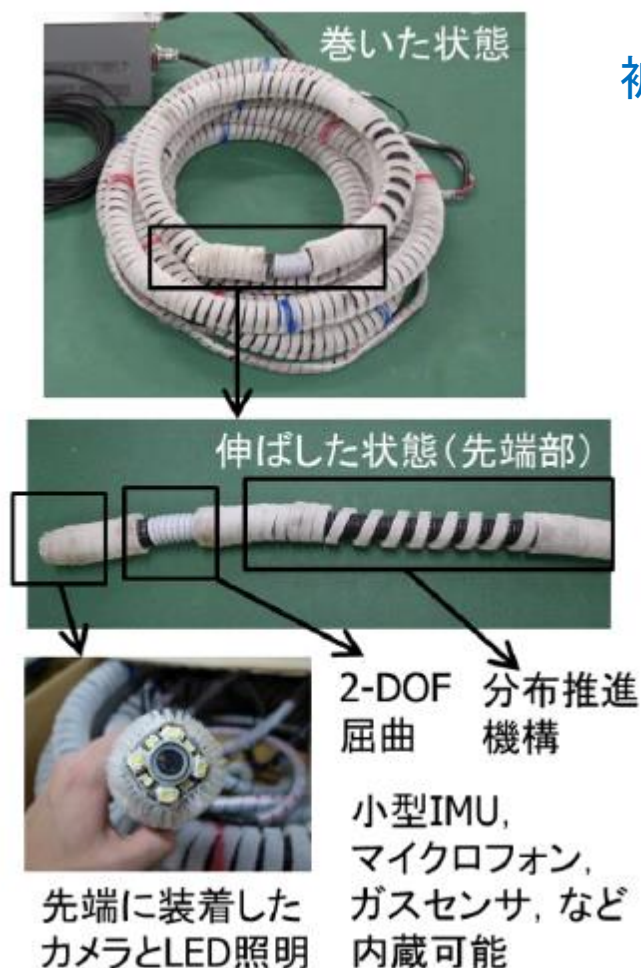
高速ICA、独立低ランク行列分析によるデモ

- リアルタイム音声聞き分け(警察備品に採用)
- ドラム、弦楽器、音声からなる複合音の分離

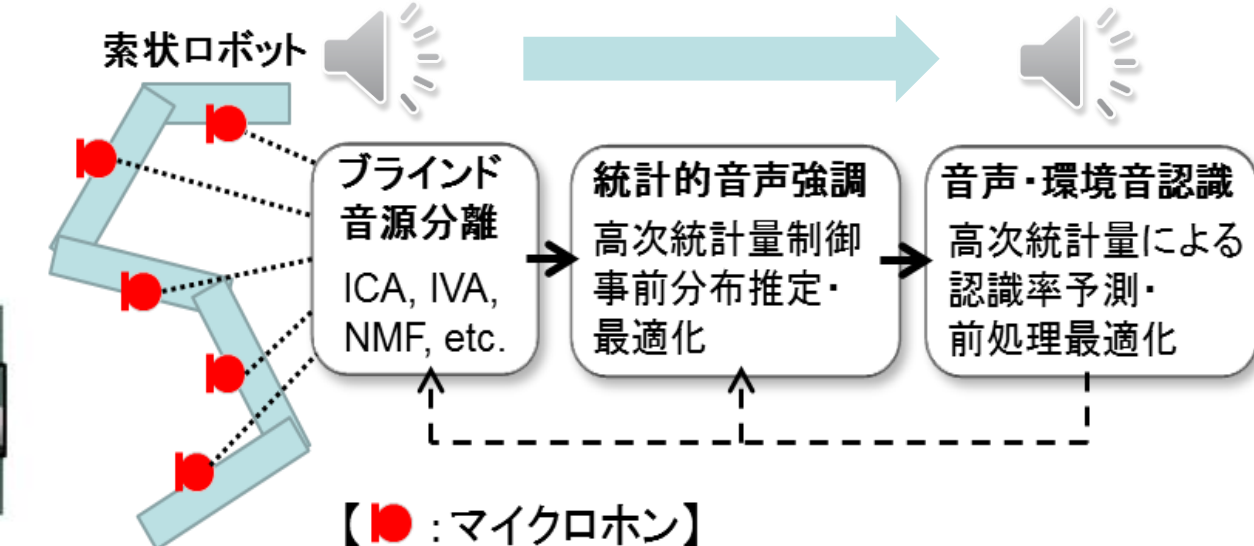


内閣府ImPACT災害対応タフロボット [2016年6月プレスリリース]

- 災害時の倒壊家屋に入り込んで被災者発見
- 環境音認識による状況把握・救助支援



被災者はいるか？



いかなる曲がりくねった形状においても
マイク同士が協調して騒音の中から被災者の声を見つけ出す

独立深層学習行列分析

Independent Deeply Learned Matrix Analysis
(IDLMA: 発音はアイドルエムエー)

ILRMAにおける問題点：音源の低ランク性？

$$\mathcal{J}_{\text{ILRMA}} = \frac{1}{J} \sum_{i,j,n} \left[\log \sum_l t_{il,n} v_{lj,n} + \frac{|w_{i,n}^H x_{ij}|^2}{\sum_l t_{il,n} v_{lj,n}} \right] - \sum_i \log |\det \mathbf{W}_i|^2$$

音源モデル（低ランク性）

空間モデル（音源間が独立）

音源によっては低ランク性が
成り立たない場合がある

音源・マイク位置，部屋の形状，
残響時間などの膨大な物理要因に依存

ならば！

事前に学習データを用いて音源モデル
の分散を推定する写像を作る

学習データの用意は非現実的
ブラインドに推定

ILRMAにおける問題点：音源の低ランク性？

$$\mathcal{J}_{\text{ILRMA}} = \frac{1}{J} \sum_{i,j,n} \left[\log \sum_l t_{il,n} v_{lj,n} + \frac{|w_{i,n}^H x_{ij}|^2}{\sum_l t_{il,n} v_{lj,n}} \right] - \sum_i \log |\det \mathbf{W}_i|^2$$

音源モデル（低ランク性）

空間モデル（音源間が独立）

- 深層学習（DNN）による強力なモデリング能力を活用する！
- 今まで培ってきた「教師あり音源分離（例：教師ありNMF）」の技術を昇華させる形で研究を発展できる。
- 急速に発展するDNN研究を我々ならではの視点で拡張する。

事前に学習データを用いて音源モデル
の分散を推定する写像を作る

学習データの用意は非現実的
ブラインドに推定

DNN音源モデルによる最尤推定

■ 独立深層学習行列分析 (IDLMA)

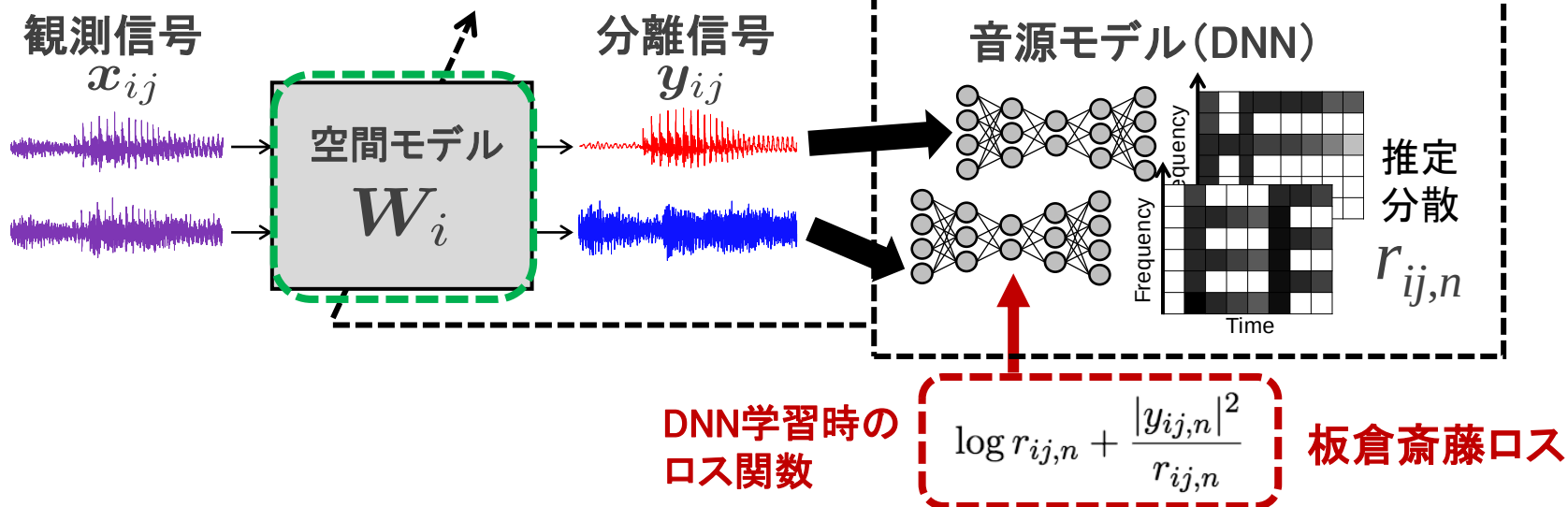
$$y_{ij,n} = \mathbf{w}_{i,n}^H \mathbf{x}_{ij}$$

$$\mathcal{J} = \frac{1}{J} \sum_{i,j,n} \left(\log r_{ij,n} + \frac{|\mathbf{w}_{i,n}^H \mathbf{x}_{ij}|^2}{r_{ij,n}} \right) - \sum_i \log |\det \mathbf{W}_i|^2$$

音源モデル (DNN)

交互に最適化

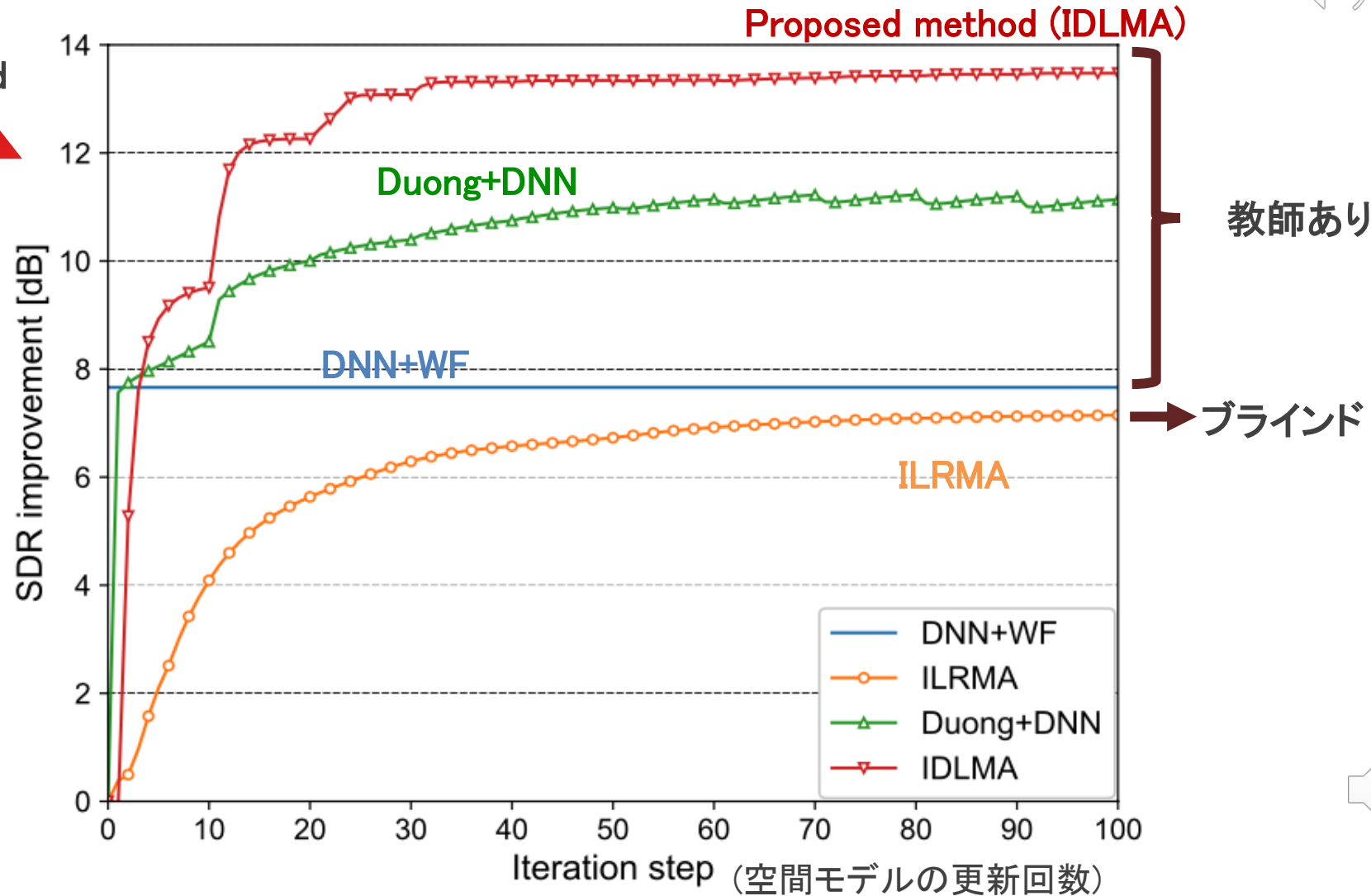
空間モデル (音源間が独立)



■ **空間モデル**: 各音源が統計的に独立となる分離行列を推定

■ **音源モデル**: J を最小化するような分散 $r_{ij,n}$ を推定するDNNを各音源ごとに構成

実験結果 (ASJ優秀学生発表賞・ICA2019 Young Scientist Grant Award)



■ 10回に1回 DNNで分散行列を更新

研究紹介2. 音楽信号解析

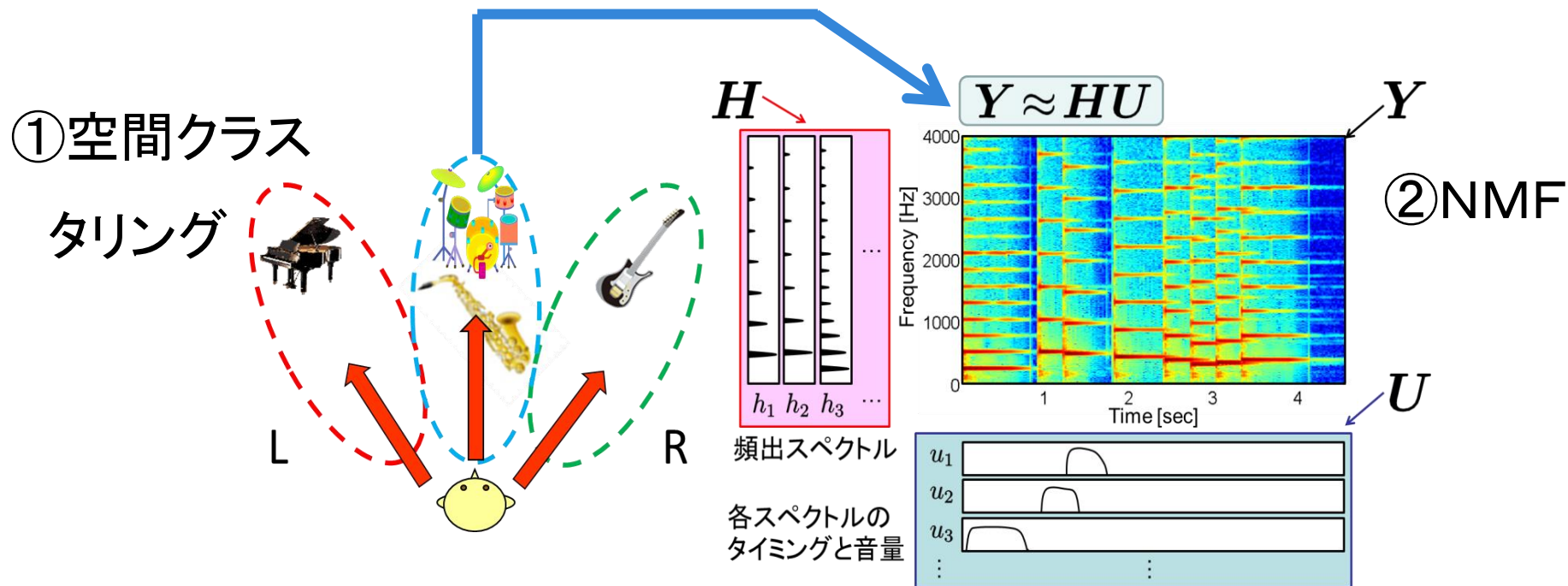
- 様々な楽器がまじりあった音楽信号の中から、自分の好きな楽器を見つけ出し、自分の好みのリミックス版を製作する。
- 非負値行列因子分解(NMF)という数理アルゴリズムに基づいて、音を事前に学習した「より簡単(低ランクかつ疎:スパース)な頻出パターン」に分解することにより、信号を解析する。

スパース分解
低ランクモデル
半教師有り学習

スパース・アンチスパース信号表現に基づく多チャネル音楽信号分離

[IEEE Trans. ASLP 2015, SIP若手奨励賞、日本音響学会学生優秀発表賞、TAF論文賞]

- 多チャネル音楽信号を効率よく分解するため、空間クラスタリングとスペクトル頻出パターン分解(非負値行列因子分解:NMF)を組み合わせた手法を提案している。
- 本手法では、空間クラスタリングでの分類エラーを学習基底で補修する機能が備わっているため、**分離に必要なスパース性**と**補修に必要なアンチスパース性**との間でトレードオフが生じる。
- そこで、各音源に合ったスパース性を選択する機能を導入した。



多チャンネル音楽信号分離デモ

📺 実際の演奏曲を教師有りNMFで分解してみた。

原曲



教師1



分離音1



教師2



分離音2



多チャンネル音楽信号分離デモ

🖥️ プロレコーディングに対応できる品質を目指して。

原曲（プロ演奏）



Saxのみを抜いた
伴奏部分



Copyright © 2014 Yamaha Corp.
All rights reserved.

サックス奏者が
消えた!?

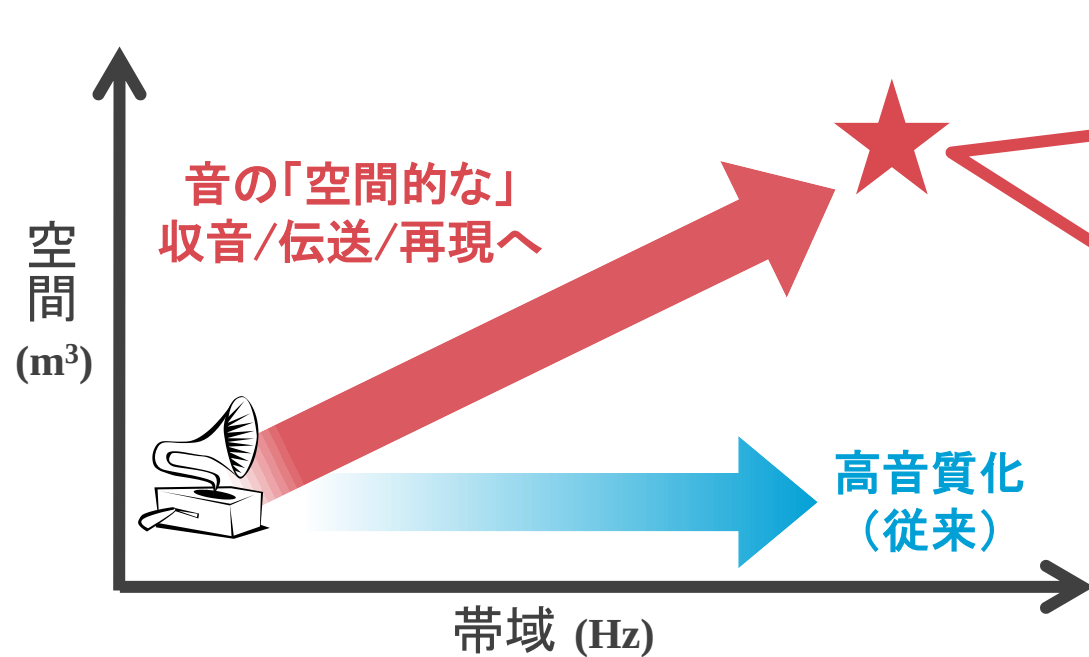


研究紹介3. 音場再現

- ステレオや5.1chサラウンドなどの従来の音の再生技術と異なり、広い領域で音の空間そのものを再現する。
- 音のホログラフを実装し、音バーチャルリアリティや音拡張現実感システムを実現する。
- 音場を「スパース(疎)」な性質に基づいて分解する。

球面調和解析
スパース最適化
音VR・AR

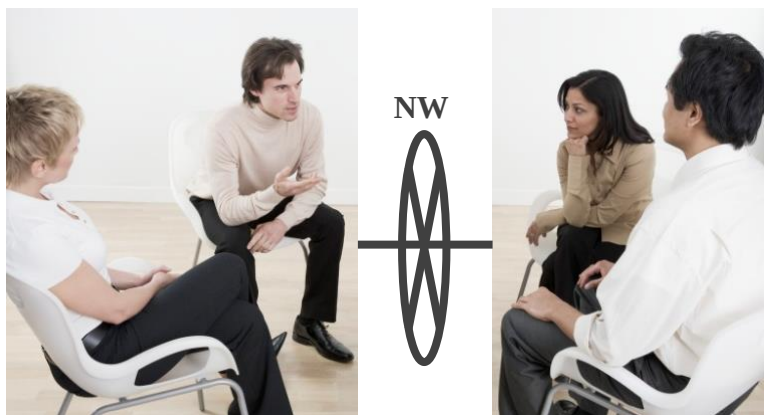
超高臨場感音響再生技術に向けて



あたかも同じ空間を
共有しているかのような
超高臨場感システム



テレプレゼンス



コンテンツ配信



家庭



劇場

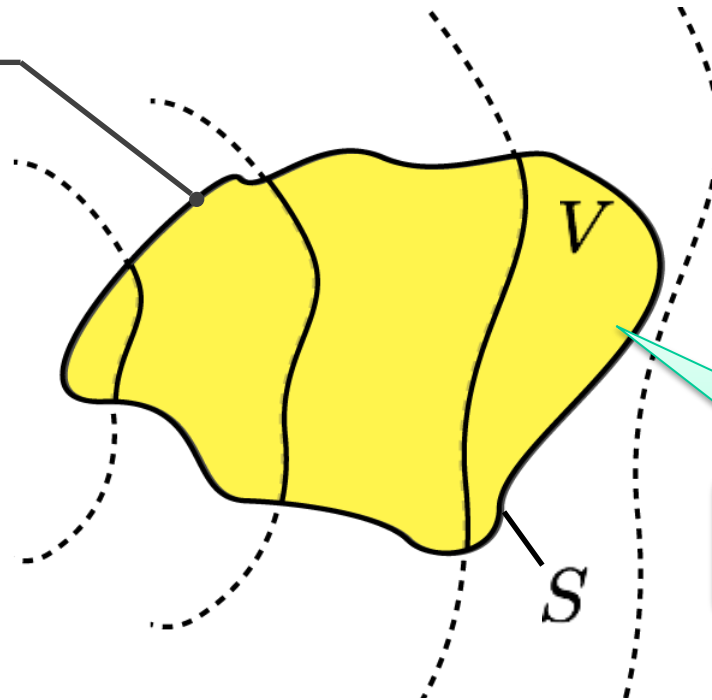


音場再現による高臨場音響再生

音場そのものを物理的に再現

Secondary source distribution

Virtual
primary sources



広い受聴領域を
実現できる可能性

- 対象領域 V 内の音場を, 境界面 S 上に配置した二次音源 (= スピーカ) を用いて, 所望の音場と一致させる

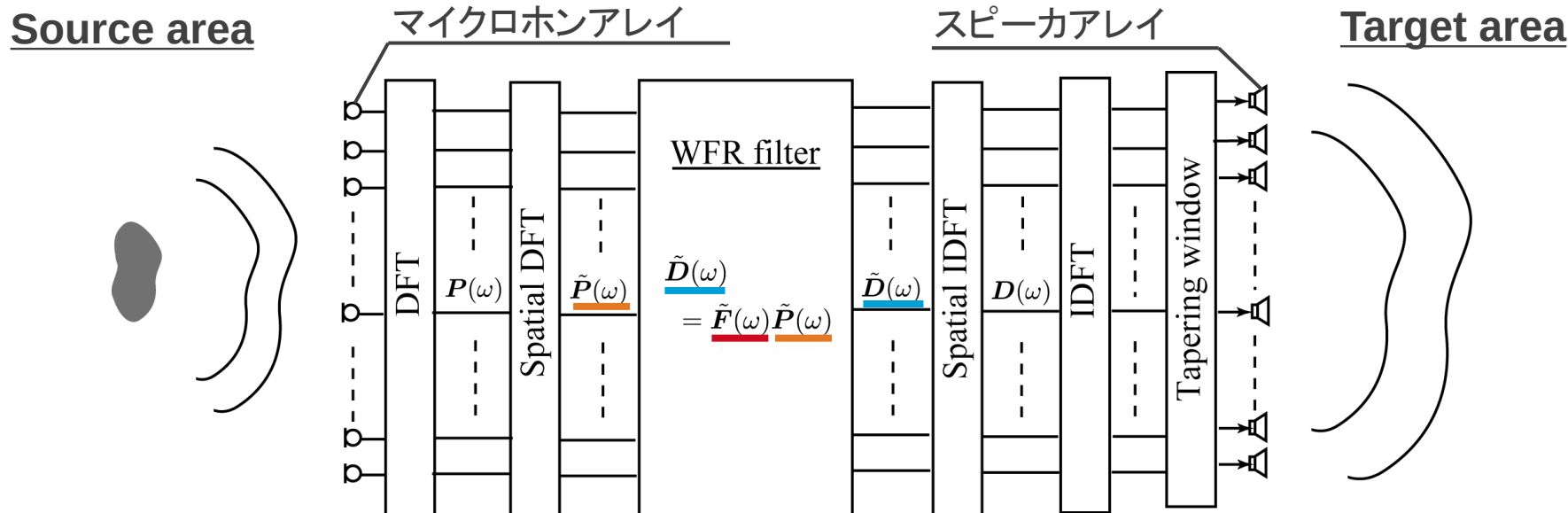
➡ 音場再現問題

直線状アレイのためのWFRフィルタ

[音響学会板倉賞]

[テレコムシステム技術賞]

実装上は時空間の2次元FIRフィルタ畳み込み: 波面再構成(WFR)フィルタ



$$\tilde{D}(\omega) = \tilde{F}(\omega) \tilde{P}(\omega)$$

スピーカ駆動信号の
時空間スペクトル

マイクロホン信号の
時空間スペクトル

WFRフィルタ: 各要素が以下で決まる対角行列

$$\tilde{F}_i(\omega) = -4j \frac{e^{j\sqrt{k^2 - k_{x,i}^2} y_{\text{ref}}}}{H_0^{(1)}\left(\sqrt{k^2 - k_{x,i}^2} y_{\text{ref}}\right)}$$

仮想再現音場例1（音源がスピーカの手前）

物理的なスピーカ列はここ



スパース表現に基づく音場分解

音場の生成モデル：モノポール音源と平面波の重ねあわせでモデル化

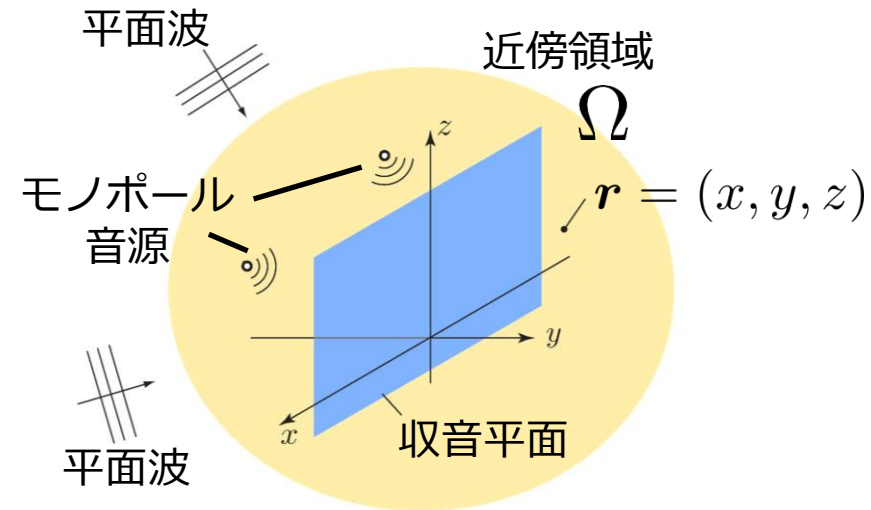
[Koyama, 2014]

非斉次・斉次

Helmholtz 方程式による表現

$$(\nabla^2 + k^2)p(\mathbf{r}, \omega) = \begin{cases} -Q(\mathbf{r}, \omega), & \mathbf{r} \in \Omega \\ 0, & \mathbf{r} \notin \Omega \end{cases}$$

$Q(\mathbf{r}, \omega)$: モノポール音源の分布



このHelmholtz方程式の解は、**非斉次項**と**斉次項**の和として書ける。

$$\begin{aligned} p(\mathbf{r}, \omega) &= p_i(\mathbf{r}, \omega) + p_h(\mathbf{r}, \omega) \\ &= \int_{\mathbf{r}' \in \Omega} Q(\mathbf{r}', \omega) G(\mathbf{r} | \mathbf{r}', \omega) d\mathbf{r}' + p_h(\mathbf{r}, \omega) \end{aligned}$$

Green関数

モノポール音源由来の
非斉次項

平面波由来の
斉次項

スパース表現に基づく音場分解

離散化し、スパース分解の枠組みで求解

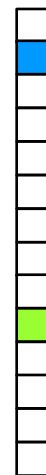
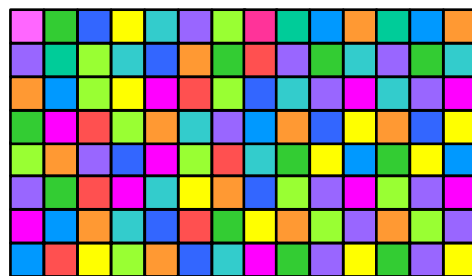
$$p(\mathbf{r}, \omega) = \int_{\mathbf{r}' \in \Omega} Q(\mathbf{r}', \omega) G(\mathbf{r} | \mathbf{r}', \omega) d\mathbf{r}' + p_h(\mathbf{r}, \omega)$$

$$\mathbf{y} = \mathbf{D}\mathbf{x} + \mathbf{z}$$

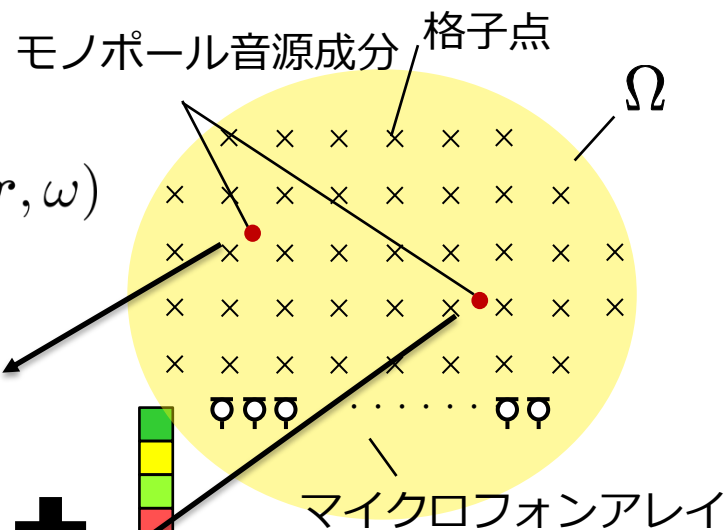
離散化



=



+



$$\mathbf{y} \in \mathbb{C}^M$$

$$\mathbf{D} \in \mathbb{C}^{M \times N}$$

$$\mathbf{x} \in \mathbb{C}^N$$

$$\mathbf{z} \in \mathbb{C}^M$$

マイクログフォンでの観測信号

Green関数を要素に持つ辞書行列

モノポール音源成分の分布

平面波由来の成分

スパース表現に基づく音場分解

モノポール音源成分が空間的にスパースであることを利用し、
スパース分解の枠組みで求解

$$y = D\mathbf{x} + \mathbf{z}$$

$y \in \mathbb{C}^M$ $D \in \mathbb{C}^{M \times N}$ $\mathbf{x} \in \mathbb{C}^N$ $\mathbf{z} \in \mathbb{C}^M$

少数の要素が非ゼロになる！

最適化問題としてのスパース分解問題

$$\begin{aligned} &\underset{\mathbf{x}}{\text{minimize}} && \|\mathbf{x}\|_p \\ &\text{subject to} && \|\mathbf{y} - \mathbf{D}\mathbf{x}\|_2^2 < \varepsilon \\ &&& 0 < p \leq 1 \end{aligned}$$

$\mathbf{X}$ のスパース性を誘導するノルム

[2015小山] $\mathbf{z}$ の項にも低ランク制約

[2015村田] 音源自体を低ランク表現

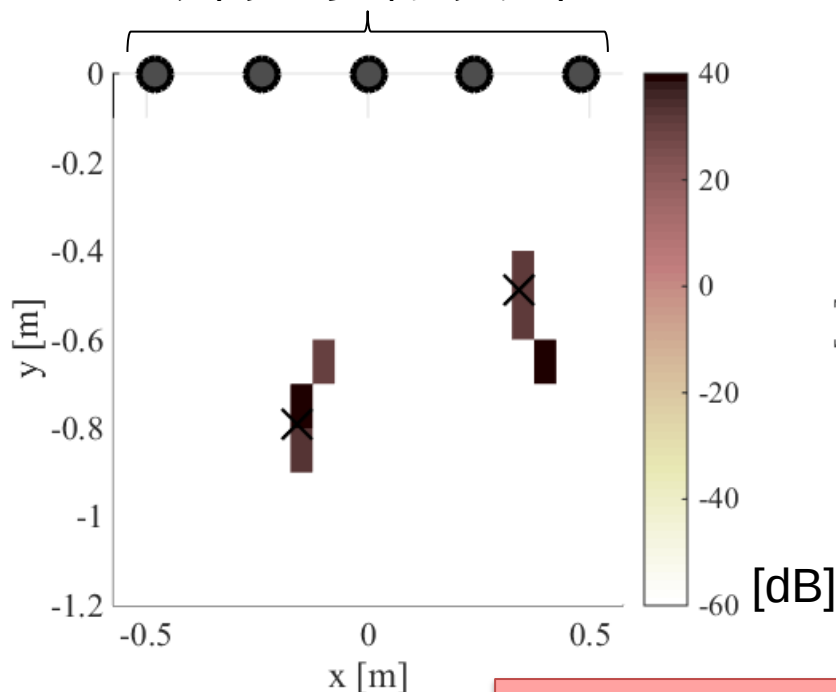
ASJ学生優秀発表賞受賞！

シミュレーション実験：音源位置推定結果

強い音源の相関，また強いノイズの環境下において頑健な性能を示した。

推定された音場のパワーの分布 白色雑音：SN比 10 dB

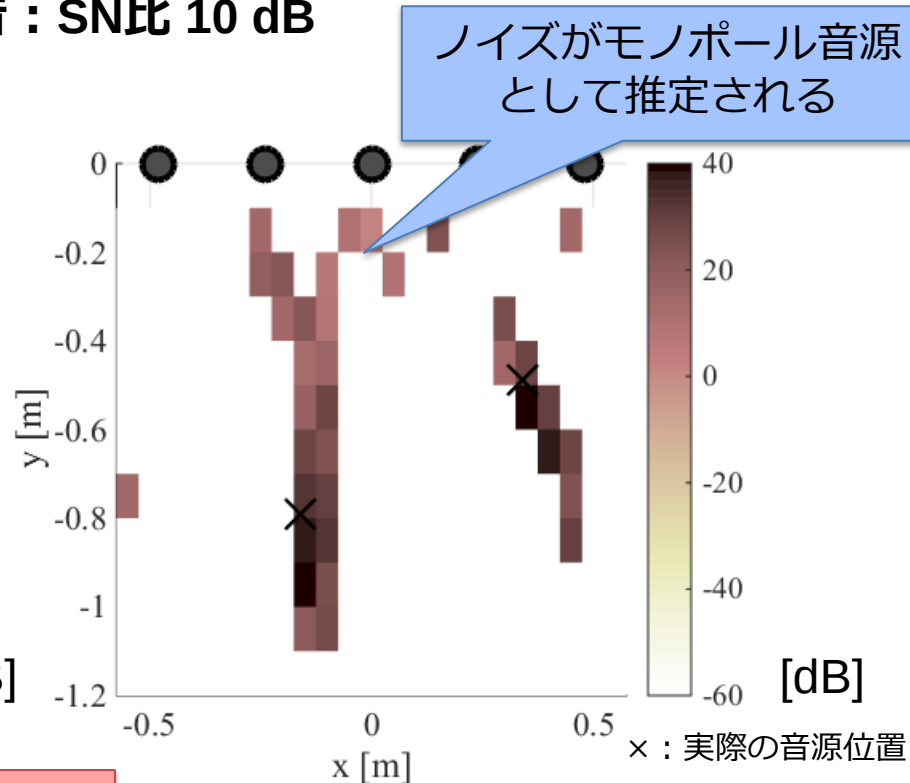
マイクロフォンアレイ



提案法

F-measure: 0.44

大きい方が高性能



従来手法 (M-FOCUSS)

F-measure: 0.25

研究紹介4. 統計的音声合成

- テキストもしくは自分の声を入れるだけで、誰の声でも、どんな訛りでも、何語でも、喋ることが出来るような統計的音声合成システムを実現する。
- 深層学習 (Deep Neural Net) の枠組みを活かし、「AIオレオレ詐欺師」と「AI防犯課刑事」を対決させて、お互いに精度を高めるAnti-Spoofing 敵対学習理論を独自に提唱

ビッグデータ
深層学習
敵対学習



音声合成の統計モデリング



入力 x と出力 y の関係をどう記述するか？ → **統計的逆問題**

声の**ゆらぎ**をどう扱うか？

人間らしい声とは何か？ 「人間らし

東京大学は世界一！

昼飯はとんこつラーメンに限る！

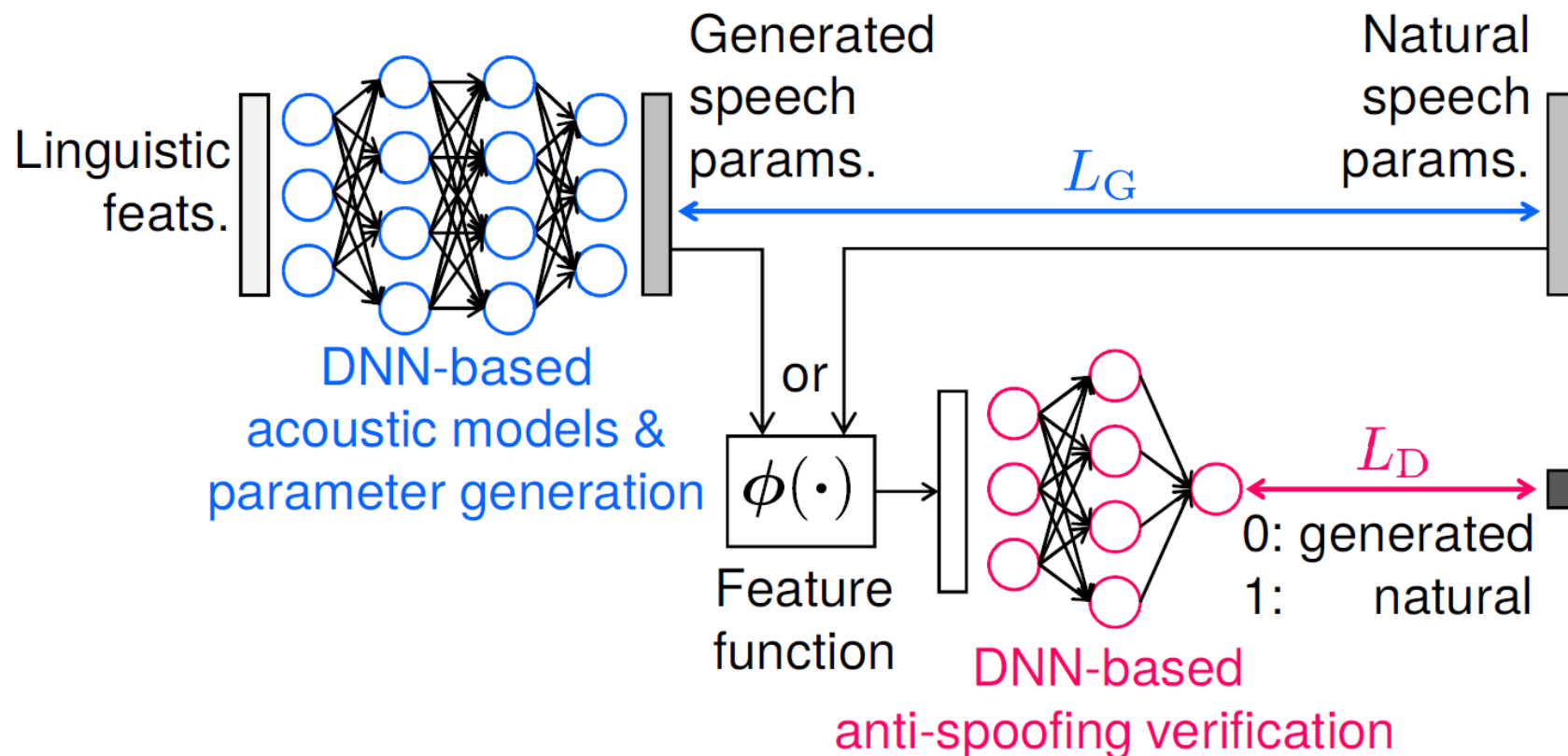
2015年
テキスト音声合成国際コンペ一位
2016年
音声変換国際コンペ一位





Anti-Spoofingと敵対する音響モデル学習理論

[ICASSP2017 SPL Student Grant賞・IEEE SPSJ 学生論文賞他]





リアルタイムDNN声質変換



2019年3月 日経xTECHプレスリリース



現在の研究課題（統計制約と最適化）



・敵対的DNN音声合成 → anti-spoofing と音響モデルの交互最適化

Anti-spoofing の学習: cross-entropy 最小化

$$L_D(\mathbf{y}, \hat{\mathbf{y}}) = -\frac{1}{T} \sum_{t=1}^T \log D(\mathbf{y}_t) - \frac{1}{T} \sum_{t=1}^T \log(1 - D(\mathbf{y}_t))$$

音響モデルの学習: MGE学習 + anti-spoofing との敵対的学習

$$L(\mathbf{y}, \hat{\mathbf{y}}) = L_G(\mathbf{y}, \hat{\mathbf{y}}) + \omega_D \frac{E_{L_G}}{E_{L_D}} \times \left(-\frac{1}{T} \sum_{t=1}^T \log D(\hat{\mathbf{y}}_t) \right)$$

コンテキストの
条件付き敵対
学習（音では
世界初！）

敵対的学習: $\mathbf{y}$ と $\hat{\mathbf{y}}$ が従う確率分布間のJSダイバージェンス最小化

[Goodfellow et al., 2014]

・異なる学習規範を用いた Generative Adversarial Network (GAN)

f -GAN [Nowozin et al., 2016]

KL, JSダイバージェンスの拡張

Wasserstein GAN (W-GAN) [Arjovsky et al., 2017]

Wasserstein 距離 (earth mover's distance) の最小化

分布に対して
異なる感度の
制約が加わる

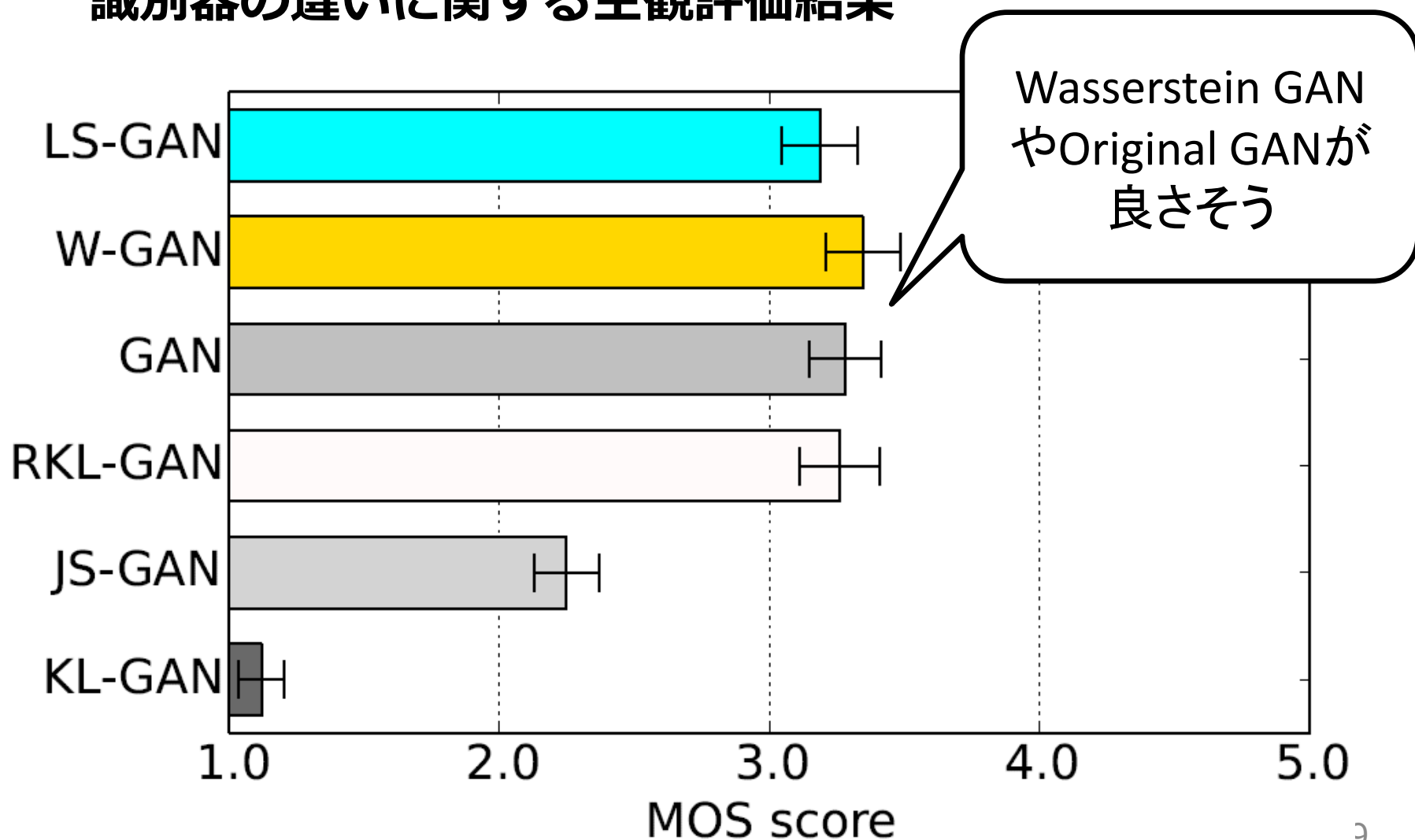
・積極的に自然音の分布を取り込むmoment-matching-DNN



様々なGANによる声質の違い



識別器の違いに関する主観評価結果



第一研への進学を志望する方へ

研究プロジェクト・国内外研究者コラボ

積極的な外部交流を推奨しています！

[2020年～]

- ・スモールデータ音響AR(科研費基盤A) with 筑波大牧野先生
- ・音響AR(二国間共同研究) with 中国科学院
- ・極限音響センシング(二国間共同研究) with オークランド大学(NZ)
- ・楽音信号分解(ヤマハ研究開発センター) with 高橋さん・近藤さん
- ・独立半正定値テンソル分析(NTT-CS研) with 池下さん
- ・分散配置アレイ音空間記録・再構成(JSTさきがけ)
- ・音場の計測と合成(東大卓越)
- ・時空間音響メディア解析制御(DMM.comラボ)
- ・次世代音声翻訳(科研基盤S分担) with NAIST中村先生・名大戸田先生
- ・高品質音声合成(科研費若手)
- ・なりすまし識別(セコム財団チャレンジ)
- ・リアルタイム広帯域声質変換(総務省SCOPE) with 電気系齋藤先生
- ・多重解像度深層分析による楽曲解析(立石財団)
- ... (他多数?)

とにかく音が好き人は集まれ！

2014年以降の研究教育実績

- ・原著論文：Top論文誌IEEE, ASA, EURASIP, ISCA 23編, 他19編
- ・国際・国内会議発表：**数えきれないので略**
- ・教員・指導学生が2014年以降に**60件の学術賞**を受賞



音の情報処理に興味がある人
統計的信号処理の数理に興味がある人
機械学習を使って生のデータを扱いたい人



一期一会なスモールデータの研究をしたい人
波動現象を数理的に考えたい人 etc.
そんなあなたにお勧めです。

