

JOINT TRACKING *and* ENHANCEMENT of AUDIO SIGNALS in REVERBERANT ENVIRONMENTS

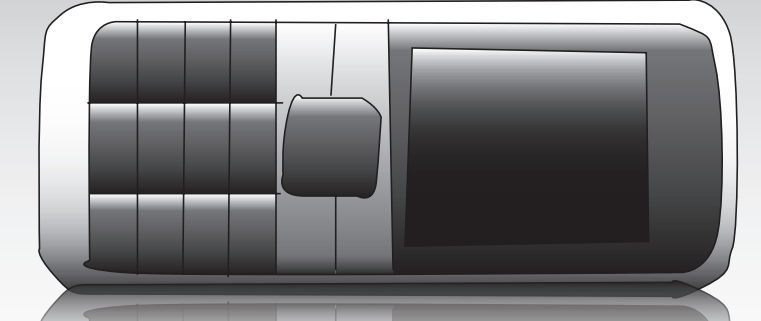
Overview

Environmental noise, music, and acoustic reverberation often lead to **unintelligibility** in speech in applications such as mobile phones or hearing aids. **Enhancement** of the degraded noisy signal is an extremely important Engineering problem for many real world audio communication systems. Enhancement relies on **prior information** about the speech, noise, and room acoustics. As exact prior knowledge is often not available, we **model belief** about the system to find the most probable restored signal.

Applications

Speech processing

Hearing aids, teleconferencing, speech recognition systems, handsfree telephones

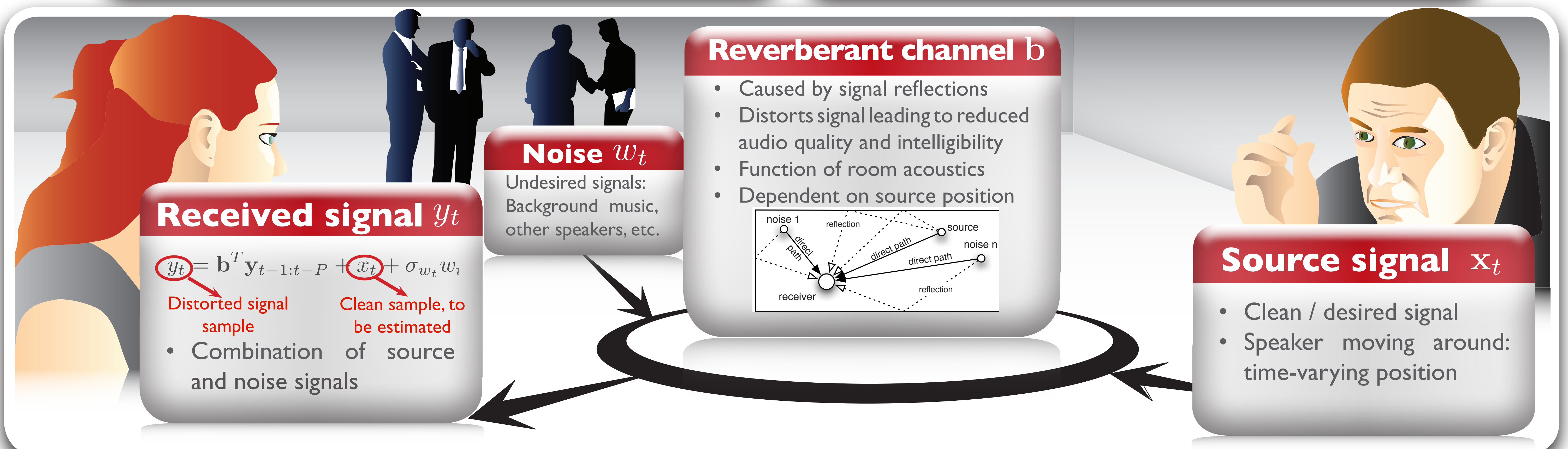


Underwater acoustics & scene analysis

Target localization, tracking, identification, classification, detection

Military & security applications

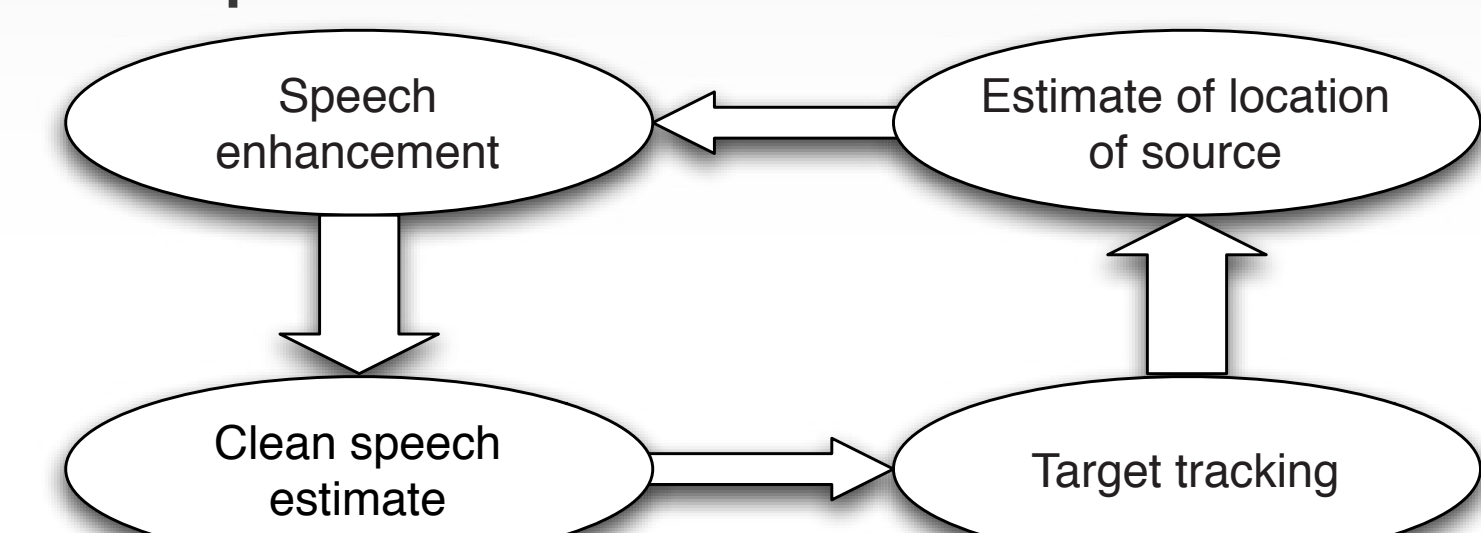
Object guidance, air traffic control, surveillance



Problem: Joint Enhancement & Tracking

Enhancement of source signal

- Removes noise and reverberation from received signal
- Dereverberation: Estimate of the room acoustics is required
- Improve estimate of room acoustics using source location



Tracking of audio source

- Estimates source location
- Source signal estimate necessary to identify true source from reflections and noise

Bayesian Blind Dereverberation

Bayesian probabilistic inference allows us to introduce uncertainty in estimation methods of the unknown parameters and to update our belief in the estimates as new data becomes available.

$$p(\mathbf{x}_{0:t}, \boldsymbol{\theta}_{0:t} | \mathbf{y}_{1:t}) = \frac{\overbrace{p(\mathbf{y}_{1:t} | \mathbf{x}_{0:t}, \boldsymbol{\theta}_{0:t})}^{\text{Likelihood}} \overbrace{p(\mathbf{x}_{0:t}, \boldsymbol{\theta}_{0:t})}^{\text{Joint prior}}}{\underbrace{p(\mathbf{y}_{1:t})}_{\text{Evidence}}}$$

Marginalization of channel parameters

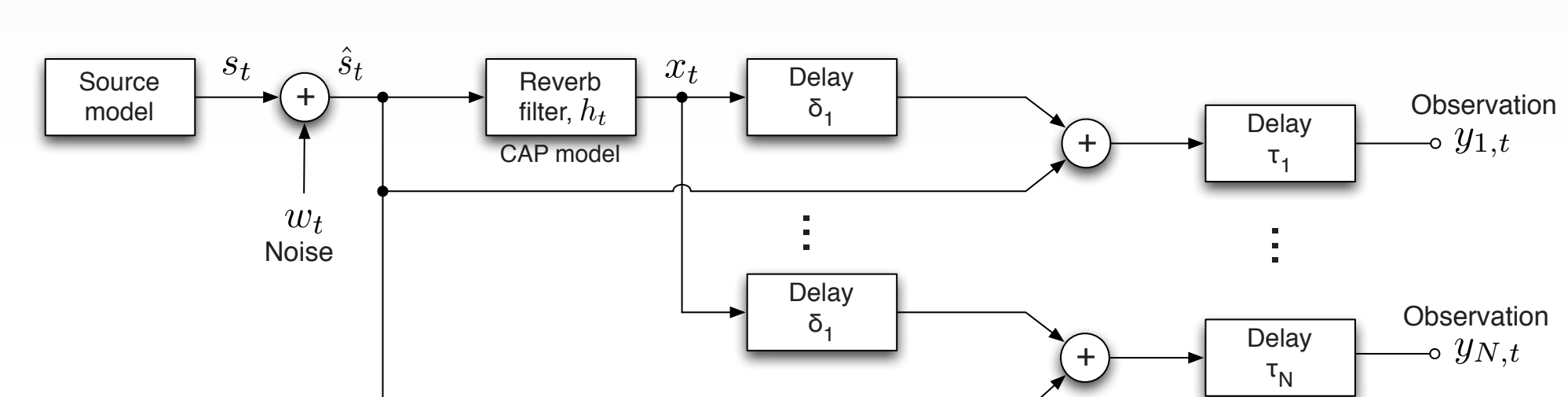
The distorting channel effects can be removed from the observed signal by **analytically marginalizing** the channel parameters.

$$p(\mathbf{x}_{0:t} | \mathbf{y}_{1:t}, \boldsymbol{\theta}_{0:t}) = p(\mathbf{x}_0 | \boldsymbol{\theta}_0) \prod_{k \in T} \left[\int p(\mathbf{x}_t | \mathbf{y}_{1:t}, \boldsymbol{\theta}_{0:t}, \mathbf{b}) p(\mathbf{b} | \mathbf{y}_{1:t}, \boldsymbol{\theta}_{0:t}) d\mathbf{b} \right]$$

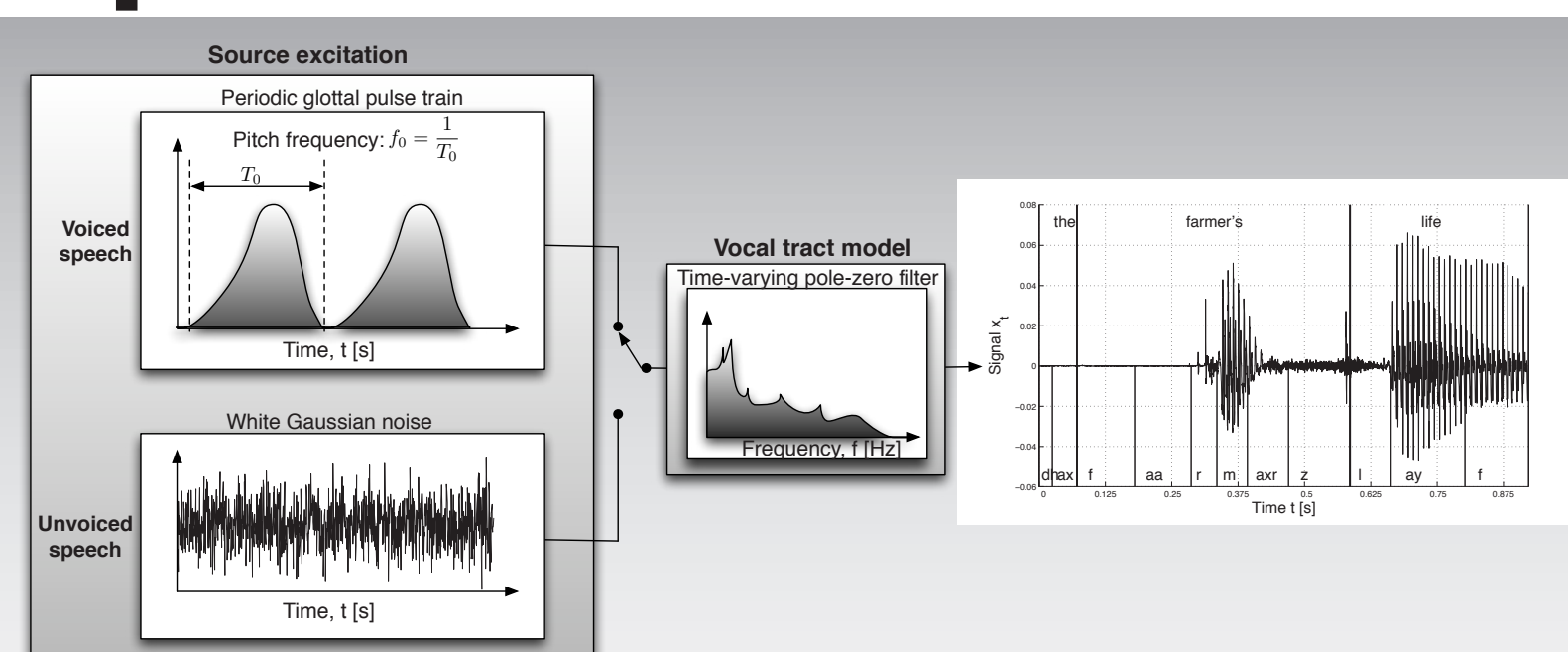
As the channel parameters are usually **unknown**, they need to be estimated like all other unknown variables. Thus, an ensemble of Kalman filters for stochastically sampled channel parameters is run simultaneously to marginalize over all the hypothesized channels.

Tracking using a time-varying channel

The source position is implicitly tracked as side-product of the channel parameter estimation by modelling the room impulse response as a function of the location of the source position relative to the sensor. The channel model is extended to the multi-sensor case.



Speech Models

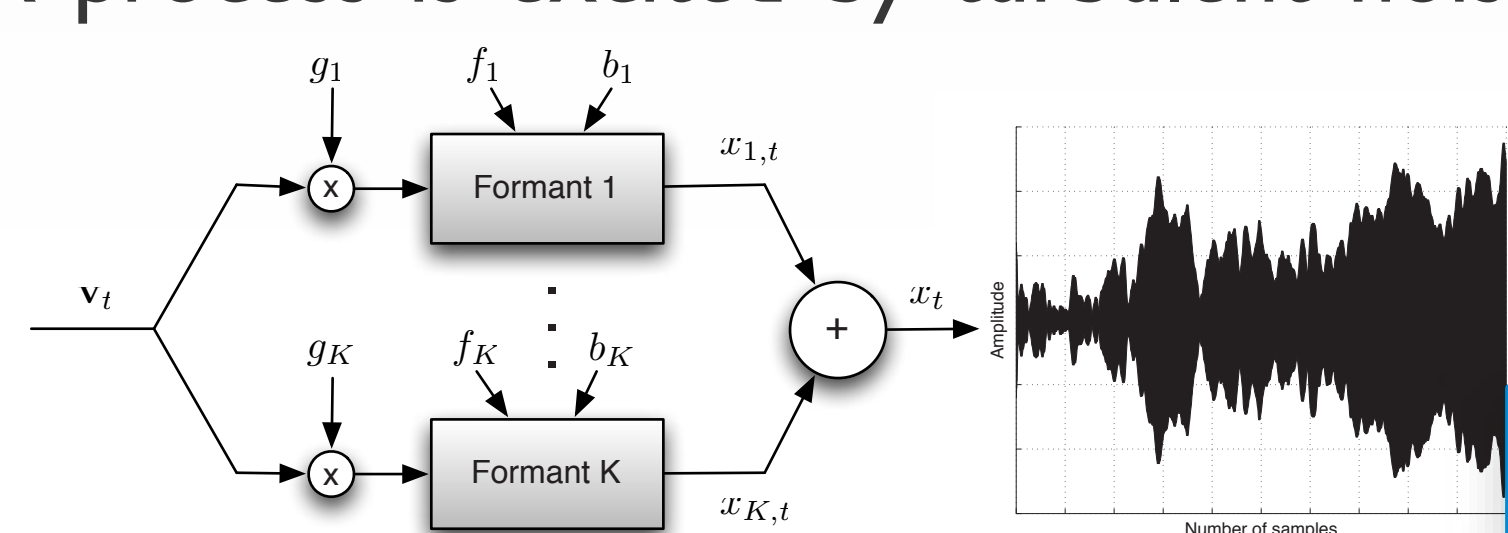


Speech can be modeled by autoregressive (AR) processes due to their accurate capturing of the short-term spectrum of speech. Time-varying AR (TVAR) parameters are used

in order to capture the continually changing properties of the vocal tract. The TVAR parameters evolve according to a first-order Markov chain to allow for a flexible representation. Voice activity detectors distinguish between voiced and unvoiced speech sequences.

For **unvoiced speech** the AR process is excited by turbulent noise, i.e., a white Gaussian noise sequence.

The vibrations of the vocal chords occurring for **voiced speech** are represented by a glottal pulse excitation of the AR process. Parallel formant synthesizers (PFSs) can be used to approximate the harmonic behavior of voiced speech. PFSs generate speech by means of several digital resonators.



[1] C. Evers and J. R. Hopgood, "Blind speech dereverberation using particle filters by marginalization of channel parameters", in process, *IEEE J. Trans. Speech Audio Proc.*, 2009.

[2] C. Evers and J. R. Hopgood, "Parametric modelling for single-channel blind dereverberation of speech from a moving speaker", *IET J. Sig. Proc.*, pp. 59–74, vol. 2, June, 2008.

[3] C. Evers, J. R. Hopgood, and J. Bell, "Acoustic models for blind source dereverberation using sequential Monte Carlo methods", *Proc. IEEE ICASSP*, Las Vegas, Apr. 2008.

[4] C. Evers, J. R. Hopgood, and J. Bell, "Blind speech dereverberation using batch and sequential Monte Carlo methods", *Proc. IEEE ISCAS*, Seattle, May 2008.

[5] J. R. Hopgood, C. Evers and S. Fortune, "Bayesian Single Channel Blind Dereverberation of Speech from a Moving Speaker", to appear in *Speech Dereverberation*, P. Naylor and N. Gaubitch (Eds.), Springer, 2008.