

REPORT DOCUMENTATION PAGE

AFRL-SR-AR-TR-05-

0135

The public reporting burden for this collection of information is estimated to average 1 hour per response, including gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments information, including suggestions for reducing the burden, to Department of Defense, Washington Headquarters Service, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any form of collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.

1. REPORT DATE (DD-MM-YYYY) 31-01-2005		2. REPORT TYPE Final Report		3. DATES COVERED (From - To) Aug 2002 - Jan 2005	
4. TITLE AND SUBTITLE Integrated Adaptive Compression				5a. CONTRACT NUMBER F49620-02-C-0047	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER 62702E	
				5d. PROJECT NUMBER	
6. AUTHOR(S) Goldstein, J. Scott Witzgall, Hanna Greene, Robert R. Zoltowski, Michael D.				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Science Applications International Corporation 10260 Campus Point Drive, San Diego, CA 92121				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) AF Office of Scientific Research 4015 Wilson Blvd. Room 713 Arlington, VA 22203-1954 NM				10. SPONSOR/MONITOR'S ACRONYM(S) USAF, AFOSR	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release, distribution is unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT The objective of the Integrated Sensing and Processing (ISP) program was a holistic approach to the design of systems. 1) SAIC designed a Joint Source-Channel Coding algorithm, designed to integrate the source data with channel coding. The work resulted in an algorithm that both incorporated a feedback of sensor information into the processing chain and broke through the traditional processing flow of optimizing the compression algorithms based on separate black boxes of source compression and channel coding and modulation. 2) SAIC also designed the Fast Adaptive Modulation algorithm that demonstrated a novel method of nonlinear optimization applied to end-to-end optimization of a radio communications system. Optimization is in "least-logs" rather than the more familiar "least-squares" optimization. Least-logs has the advantage that it is less sensitive to outliers than least-squares optimization. As such, it can be applied to curve-fitting and curve-finding applications in noisy images and displays. This methodology measures features, not pixels, with the resulting SNR gain due to integration over the feature.					
15. SUBJECT TERMS nonlinear optimization, least-logs, low probability of intercept, blind Rake receiver, source compression algorithm, reconstruction error					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT None	18. NUMBER OF PAGES 36	19a. NAME OF RESPONSIBLE PERSON Dr. Belinda King
a. REPORT U	b. ABSTRACT U	c. THIS PAGE U			19b. TELEPHONE NUMBER (Include area code) 703/696-8409

Integrated Adaptive Compression (Final Report)

J. Scott Goldstein, Project Manager, jay.s.goldstein@saic.com, 703-814-7750

Hanna Witzgall, Principal Investigator, witzgallh@saic.com, 703-814-8246

Robert R. Greene, Principal Investigator, greener@saic.com, 703-814-7700

Michael D. Zoltowski, Consultant, Purdue University

*Science Applications International Corporation
10260 Campus Point Drive, San Diego, CA 92121*

Contract Number: F49620-02-C-0047

Start Date: August 1, 2002

End Date: January 31, 2005

Contract Amount: \$733,538

<u>Contract Issued by:</u>	<u>Government Program Managers:</u>
USAF, AFRL	Dr. Douglas Cochran, DARPA/DSO
AF Office of Scientific Research	Dr. Carey Schwartz, DARPA/DSO
4015 Wilson Blvd. Room 713	Defense Advanced Research Projects Agency
Arlington, VA 22203-1954	3701 North Fairfax Drive
Exiquio R. Balboa (703)696-5956	Arlington, VA 22203-1714

January 31, 2005

"This report was prepared as an account of work sponsored by an agency of the United States Government. Neither the United States Government, nor any agency thereof, nor any of their employees, makes any warranty, express or implied, or assumes any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represents that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government or any agency thereof. The views and opinions of authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof."

Table of Contents

Chapter 1 Problem Statement/ Research Objective	1
Chapter 2 Integrated Sensing and Processing for Joint Source and Channel Coding	2
How the New Technique relates to the ISP Objective.....	2
Summary	3
Background	4
DPCM Compression Architecture	5
Mitigating Channel Noise	6
Performance Results and Analysis.....	9
Conclusions.....	15
Publications.....	15
Chapter 3 Integrated Sensing and Processing for Fast Adaptive Modulation	16
A Novel Non-linear Optimization Approach: Least-Logs.....	16
Background: Multipath Spreading	18
Mitigation of Inter-Symbol Interference due to Multipath	18
Least Logs Optimization: Mitigation of Timing Errors.....	22
Coherent Rake Receiver Implementation	24
FAM Concepts for Low Probability of Intercept Applications	26
Results and Estimate of Technical Feasibility	32

Table of Figures

Figure 2-1: Plots the Reconstruction Error for various loading levels and noise levels for DPCM, PCM and the new L-DPCM compression methods.....	3
Figure 2-2: Compression Architecture.....	5
Figure 2-3: The effect of loading the auto-correlation vector on the impulse response of the prediction filter: (a) Impulse response with no loading, (b) Impulse response with loading.	8
Figure 2-4: Times series of the speech data on which the performance analysis is based. (a) Short section of speech data, (b) Long section of speech data.....	9
Figure 2-5: (a) Reconstruction Error vs. loading level for 2 bit quantization over short voiced sequence in three AWGN channels, (b) Quantization Noise vs. Loading, (c) Impulse Response Energy vs. Loading.....	10
Figure 2-6: (a) Reconstruction Error vs. loading level for 4 bit quantization over short voiced sequence in three AWGN channels, (b) Quantization Noise vs. Loading, (c) Impulse Response Energy vs. Loading.....	12
Figure 2-7: (a) Reconstruction Error vs. loading level for 4 bit quantization over long voiced sequence in three AWGN channels; (b) Quantization Noise vs. Loading; (c) Impulse Response Energy vs. Loading.....	13
Figure 2-8: Comparison of original and reconstructed signal with a low loading level of -28 [dB] (left) with a high loading level 8[dB] (right) in a noise environment of $\sigma_c^2/\sigma_s^2 = .564$	14
Figure 3-1: Feedback of a limited amount of information from the receiver to the transmitter allows the transmitter to modify the codes used to represent the data.	16
Figure 3-2: (a) The horizontal partial derivative of $\log(r)$, red is negative, blue is positive, green is near zero. (b) A white vertical line of ones against a black background of zeros. (c) The convolution of (a) with (b), note that the positive and negative regions in (c) have expanded not only along the line, but away from the line as well, a special property of the logarithm.	17
Figure 3-3: Multipath can create several correlation peaks when the signal is despread by correlation.....	18
Figure 3-4: The correlation peaks from consecutive symbols can overlap, when the total multipath spread is larger than the symbol period.	19
Figure 3-5: Using different spreading codes for each symbol period will de-interleave the multipath arrivals from each symbol.	19
Figure 3-6: A lag display is formed by time shifting the output from successive correlators by the length of one symbol. Each arrival path will form a curve in the lag display.....	20
Figure 3-7: A simulated version of the display using three pairs of spreading codes.....	20
Figure 3-8: A simulated version of the lag display using only one pair of spreading codes. Note the threefold increase in correlations, increased side lobe background level, and Inter-Symbol Interference at the top where the lines cross.....	21
Figure 3-9: Doppler was simulated by calculating the variation in delay due to the time varying change in distance from a source (T) to a receiver (R) via three point reflectors, as the receiver moves between the reflectors.	22
Figure 3-10: The gradient of the $\log(r)$ potential formed by convolving the function in Figure 2(a) with the lag display.....	23
Figure 3-11: Linear templates translate and rotate in the gradient field. Points on the template in the blue tend to move to the left, in the red to the right. Translational and rotational moments on the templates can be calculated by integrating the gradient field along the templates.	23
Figure 3-12: At least one of the templates converges to each of the lines in the lag display, located at the positive-negative (red-blue) boundaries in the gradient field.....	24
Figure 3-13: The real and imaginary parts of the sampled correlator outputs along the template trajectories. Note the sinusoidal envelope.	25
Figure 3-14: Segments of a set of six decoded output sequences from six templates are shown in Figure 14. It is easy to see that the second, third, and fourth output sequences are very similar. They would be coherently combined.....	26

Figure 3-15: A lag display of the absolute value squared of the complex correlator output data. A very slight expression of several sinusoidal curves can be seen running vertically down the display.	27
Figure 3-16: The sinusoidal curves in the lag display are enhanced by the convolution of the derivative of $\log(r)$ with the real (a) and imaginary (b) parts of the complex version of the lag display.	28
Figure 3-17: Six multipath arrivals are now clearly seen in the enhanced lag display. The second, third, and fourth arrivals form the cluster of interest, as before.	28
Figure 3-18: Templates with three parameters, delay with the amplitude and phase of the artificially induced cosine Doppler delay, easily converge to an accurate estimate of the timing.	29
Figure 3-19: An additional LPI application is the coherent combination of transmissions from several tight beam transmitters whose signals converge on a central location.	29
Figure 3-20: The lag display for a scenario with seven multipath arrivals, received at -12 dB, and appearing at 3 dB at the correlator output.	30
Figure 3-21: The algorithm successfully enhances the seven multipaths in the lag display.	31
Figure 3-22: At least one template successfully fits each of the curves in the enhanced lag display.	31

Chapter 1 Problem Statement/ Research Objective

The stated objective of the Integrated Sensing and Processing (ISP) program was to consider a more holistic approach to the design of systems so that relevant sensor information is astutely incorporated into the processing chain to improve user objectives. Feeding back sensor information into the processing chain to hone the output performance was considered a critical element of the integrated design as was breaking out of traditional black box processing elements to optimize over the entire system objective. To validate ISP's thesis, SAIC explored several different variations on the theme of integrating sensors with their communication systems.

The Joint Source-Channel Coding effort focused on algorithms designed to integrate the source with channel coding. The work resulted in an algorithm that both incorporated a feedback of sensor information into the processing chain and broke through the traditional processing flow of optimizing the compression algorithms based on separate black boxes of source compression and channel coding and modulation. It achieved this by feedback of the sensor's channel information to the source compression algorithm, which enabled the algorithm to adaptively optimize itself to minimize the systems overall reconstruction error.

The Fast Adaptive Modulation task has demonstrated a novel method of nonlinear optimization that can be applied to end-to-end optimization of a radio communications system. The optimization methodology applied here is optimization in "least-logs" rather than the more familiar "least-squares" optimization. Least-logs optimization has the advantage that it is less sensitive to outliers than least-squares optimization. As such, it can be applied to curve-fitting and curve-finding applications in noisy images and displays. It follows that this methodology measures features, not pixels, with the resulting SNR gain due to integration over the feature.

In this application, the optimization procedure allows the estimation of a small number of parameters that characterize the combined effects of variations in clock rates that depend on components in both radios as well as the motion of the platforms and number and types of scatterers in the propagation environment. A single successful optimization process enables several different functionalities:

- 1) A classifier at the receiver that clusters the signal output sequences into related multipath arrivals, then selects the modulation parameters used by the transmitter,
- 2) A pilot-signal-free (blind) Rake receiver,
- 3) Coherent combination of signals from multiple transmitters for low probability of interception applications,
- 4) Recovery of hidden timing information for low probability of exploitation applications.
- 5) Enhancement of linear features in displays based on a log transform is achieved as an intermediate product in the optimization process.

Chapter 2 Integrated Sensing and Processing for Joint Source and Channel Coding

This research effort focused on algorithms to optimize the source compression for minimum reconstruction error in various channel noise environments. Specifically SAIC proposed to “feedback” knowledge of the communication environment into improving upon a new data compression design. Ideally the new compression algorithms would incorporate some novel reduced rank signal processing techniques, an expertise of this group.

The work resulted in a fundamentally new approach to lower the reconstruction error over noisy channels with no additional error correction coding or bit redundancy required. Results showed significant ($> 10\text{dB}$ on speech data) reductions in the reconstruction error power when estimates of the channel noise were dynamically incorporated into the source compression algorithm. In this way the source compression filter could be optimized to minimize the over-all reconstruction error as opposed to the source residual.

How the New Technique relates to the ISP Objective

The research validated ISP’s thesis that integrated sensing and processing can improve system performance. The work resulted in an algorithm that both incorporated a feedback of sensor information into the processing chain and broke through the traditional processing flow of optimizing the compression algorithms based on separate black boxes of source compression and channel coding and modulation. It achieved this by feedback the sensor’s channel information to the source compression algorithm which enabled the algorithm to adaptively optimize itself to minimize the systems overall reconstruction error. Previous research has focused on incorporating channel information into channel modulation designs. This is the first application that the authors are aware of where knowledge of the channel noise is feed into the source compression algorithm. This architecture enables the reduction of the reconstruction error over certain noise environments with no added redundancy for error control coding. In this sense the research achieved the ISP objective by improving performance through the incorporation of a more holistic system design and sensor feedback mechanism.

Summary

This section describes a computationally efficient technique to improve the reconstruction quality of differential pulse coding modulation (DPCM) source compression in AWGN channels. DPCM is often used to compress a correlated signal prior to transmission. However, it is well known that DPCM systems are more sensitive to channel noise than the simple pulse code modulation (PCM) technique, which simply quantizes the source signal. The proposed technique integrates estimates of the channel noise directly into the source compression algorithm to optimally load the auto-correlation vector used in DPCM. Analysis shows that loading this vector modifies the prediction filter's impulse response in such a way that the impact of channel noise is reduced, albeit at the expense of increased quantization noise. This tradeoff allows the new loaded (L-DPCM) to minimize the reconstruction error, the user metric of interest, rather than the source residual. The technique is demonstrated on speech data for several noise environments, quantization levels and loading levels and shows significant performance improvement over the standard, non-integrated DPCM implementation.

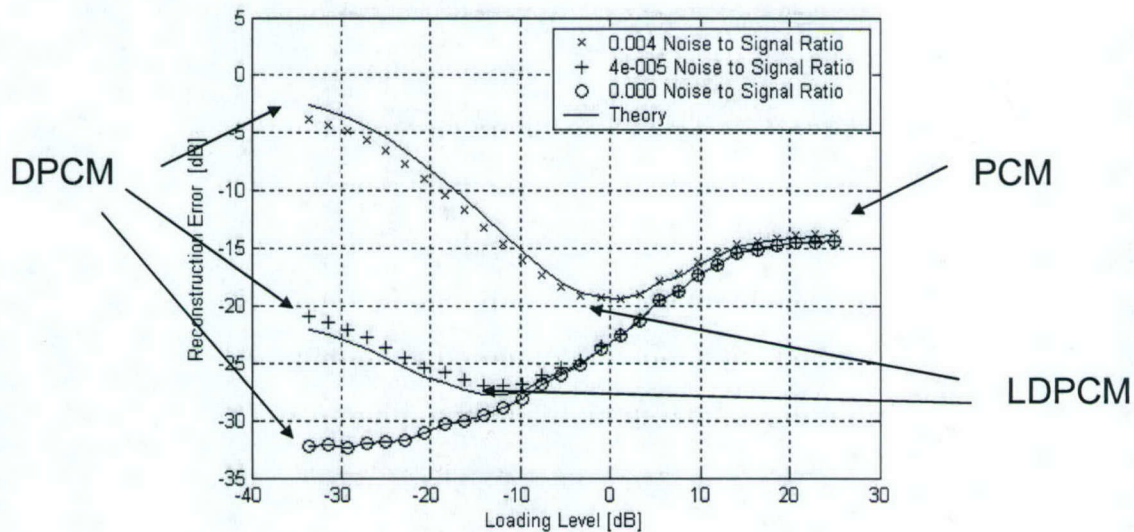


Figure 2-1: Plots the Reconstruction Error for various loading levels and noise levels for DPCM, PCM and the new L-DPCM compression methods.

Figure 2-1 illustrates the potential of the new technique and compares its performance to the commonly implemented unloaded DPCM and PCM compression methods. Unloaded DPCM performance in three different noise environments is shown by the reconstructed values on the far left of Figure 2-1. PCM's performance is indicated by the right most reconstruction values. The loaded (L-DPCM) reconstruction values are illustrated by the points in between these two methods. Note how the L-DPCM can optimize itself for a any noise environment to improve performance upon either the classical DPCM or PCM techniques.

Background

Differential pulse code modulation (DPCM) is often used to compress correlated signals prior to transmission. A well compressed signal can minimize a system's power and bandwidth requirements. The basic premise of DPCM source compression is to model the source signal with a prediction filter and encode the remaining information in a quantized error residual. The modeling reduces the dynamic range of the information needed to be transmitted. This in turn minimizes the quantization error which contributes to the overall reconstruction error for a fixed number of bits.

A major problem with this type of compression is its sensitivity to channel noise. In general DPCM systems achieve higher source compression ratios at the expense of increased sensitivity to channel noise. In extreme cases, the quality of a decompressed DPCM signal can become much worse than using no prediction and simply quantizing and sending the original signal, as is done with the standard pulse code modulation (PCM) implementation. The reason DPCM is so sensitive to channel noise is that the prediction filter is designed to minimize the residual error, causing the filter's impulse response to be powerful. This results in maximizing the spurious noise response in the reconstruction process. If the channel noise corrupting the transmitted residual is relatively large, then the reconstructed signal will contain a large noise component. Therefore, in noisy channels, it may be better not to minimize the residual error power to reduce the contribution of channel noise at reconstruction. Thus the traditional method of designing prediction filters to minimize the residual error often results in poor performance in noisy channels.

Several methods have been proposed to reduce DPCM sensitivity to channel noise. Early attempts to address this problem suggested using estimates of the channel noise to decide how many bits to allocate to the quantization of the error residual versus how many bits to allocate to error correction. The study determined that it is ineffective to finely quantize the error residual in relatively high channel noise. Therefore, in noisy channels the approach allocates more bits to error correction at the expense of increased quantization noise. The method showed an 11 dB improvement on 8 bits/sample speech data in noisy channel conditions. One drawback with this approach is that it does not improve performance for low bit rates because it uses bits to both quantize the residual and for error correction. Later attempts to address the DPCM sensitivity problem suggested much more elaborate solutions based, for example, on nonlinear prediction filters or Markov models combined with Viterbi decoding. None of these methods fundamentally change the optimization metric of the prediction filter from an intermediate measure of residual error to the user metric of overall reconstruction error.

To address these issues, SAIC developed a new technique that integrates knowledge of the channel conditions directly into the source compression algorithm so that the prediction filter can be optimized based on the reconstruction error. The filter modification takes the form of a scalar addition, effectively loading the diagonal of the Toeplitz autocorrelation matrix. This causes the adaptive prediction filter to change its impulse response function in such a way as to make the reconstruction process less sensitive to channel noise, albeit at the expense of a decrease in prediction accuracy, and

an increase in quantization noise. This results in a performance tradeoff similar to the one described above which trades increased quantization noise for reduced sensitivity to channel noise. The technique differs from the previous approach in that it acts directly on the prediction filter itself to increase its robustness to channel noise, instead of allocating more bits to error correction. In this way the new technique can show significant improvement even at low bit rates.

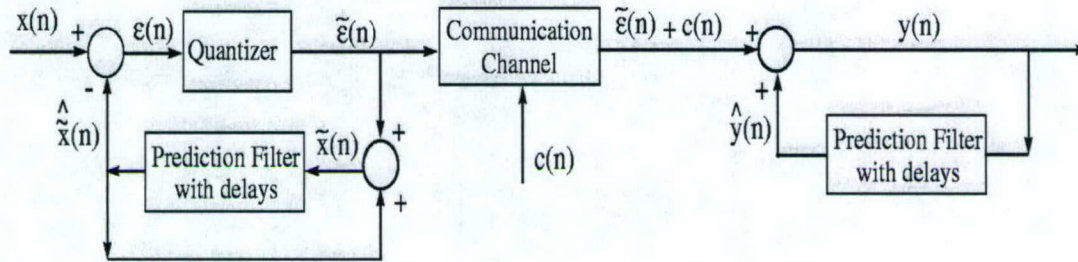


Figure 2-2: Compression Architecture

DPCM Compression Architecture

The DPCM source compression architecture is shown in Figure 2-2. A temporal signal $x[n]$ enters the encoder. A residual error $\varepsilon[n]$ is formed from the difference between the signal $x[n]$ and the encoder's signal estimate $\hat{x}[n]$:

Equation 2-1: Residual Error

$$\varepsilon[n] = x[n] - \hat{x}[n]$$

The quantized residual $\tilde{\varepsilon}[n]$ and the filter model are transmitted across the channel. Note that the predicted value $\hat{x}[n]$ is estimated from the previously reconstructed values $\tilde{x}[n]$ which are based on the quantized error residual $\tilde{\varepsilon}[n]$. This is an important because it has a stabilizing effect which prevents a large accumulation of the quantization noise in the reconstructed signal.

At the receiver, the signal $y[n]$ is reconstructed from the quantized residual $\tilde{\varepsilon}[n]$ corrupted by the channel noise $c[n]$, and the receiver's predictive estimate $\hat{y}[n]$ which is contaminated by the response of the prediction filter to previous channel noise:

Equation 2-2: Reconstructed Value

$$y[n] = \hat{y}[n] + \tilde{\varepsilon}[n] + c[n]$$

If there is no channel noise, $y[n] = \tilde{x}[n]$ and the error is entirely due to the quantization of the residual.

The prediction filter $\mathbf{w}$ is designed to minimize the mean squared error between the signal input $x[n]$ and its predicted value $\hat{x}[n]$

Equation 2-3: Filter minimizes difference between input and prediction.

$$\mathbf{w} = \arg \min_{\mathbf{w}} \langle \|x[n] - \hat{x}[n]\|^2 \rangle$$

Its solution is

Equation 2-4: Wiener-Hoff Solution

$$\mathbf{w} = (\mathbf{X}_{M-1} \mathbf{X}_{M-1}^H)^{-1} \mathbf{X}_{M-1} \mathbf{x}^H$$

where $\mathbf{X}_{M-1}$ is a $M-1 \times K$ data matrix, whose K columns consist of snapshots containing the $M-1$ past input values for the k -th input of the $1 \times K$ temporal signal vector $\mathbf{x}$. By zero padding the original data sequence $\mathbf{x}$, a covariance matrix with Toeplitz symmetry can be formed. This symmetry allows a computationally efficient solution to $\mathbf{w}$ using the popular Levinson-Durbin recursion. Use of this Toeplitz symmetric matrix also implies that the data statistics are completely described by signal's auto-correlation vector $\mathbf{r}$ with a zero lag as its first element. This will be a key point to the compression method described in this paper.

Mitigating Channel Noise

This paper describes a technique which improves the performance of DPCM systems in channel noise. The performance metric examined is the reconstruction error σ_R^2 . The empirical calculation of the reconstruction error is defined as the difference between the unquantized input $x[n]$ and the reconstructed output $y[n]$:

Equation 2-5: Empirical Reconstruction Error

$$\sigma_R^2 = \sum_{n=1}^N |x[n] - y[n]|^2$$

where N is the length of the data under consideration. Previous studies have shown that the reconstruction error can also be analyzed in terms of the quantization noise power σ_q^2 and the filter's cumulative response to the channel noise σ_c^2 :

Equation 2-6: Analytical Reconstruction Error

$$\sigma_R^2 = \sigma_q^2 + \sigma_c^2 \frac{1}{N} \sum_{n=1}^N |h_w[n]|^2$$

where $\mathbf{h}_w[n]$ represents the impulse response of the filter. Equation 2-6 shows that the impact of the channel noise on the reconstruction error is proportional to the prediction

filter's impulse response energy. This is intuitively appealing since the impulse response exactly describes how a noise impulse will distort the output signal. If the prediction filter's impulse response energy is large, then the channel noise corrupting the residual will have a large contribution. The cumulative effect that white Gaussian noise σ_c^2 will have on the reconstructed signal is therefore proportional to the energy in the filter's impulse response $\mathbf{h}_w[n]$ over the time interval N . Equation 2-6 allows the calculation of the reconstruction error without explicit knowledge of the reconstructed signal $y[n]$. This allows the transmitter to optimize the prediction filter based on minimizing the reconstruction error.

The quantization noise σ_q^2 is also indirectly associated with the impulse response energy of the prediction filter. The quantization noise can be calculated directly using:

Equation 2-7: Quantization Noise

$$\sigma_q^2 = \sum_{n=1}^N |\varepsilon[n] - \tilde{\varepsilon}[n]|^2$$

where $\varepsilon[n]$ and $\tilde{\varepsilon}[n]$ represent the residual and quantized residual error respectively. In general the better the filter is at predicting the signal, the more it can compress the dynamic range of the residual, which results in less quantization noise. However better compression results in higher impulse response energy, since a filter's impulse response will have to be greater to adequately represent the energy contained in the captured modes of the original signal. Lower impulse response energy therefore translates into an increased quantization error. Note that quantization noise is not multiplied by the impulse response energy in the reconstruction process in Equation 2-6 since the DPCM compression architecture incorporates these known quantization effects into the prediction estimate. As noted before, this prevents the quantization noise from accumulating over the entire response of the filter.

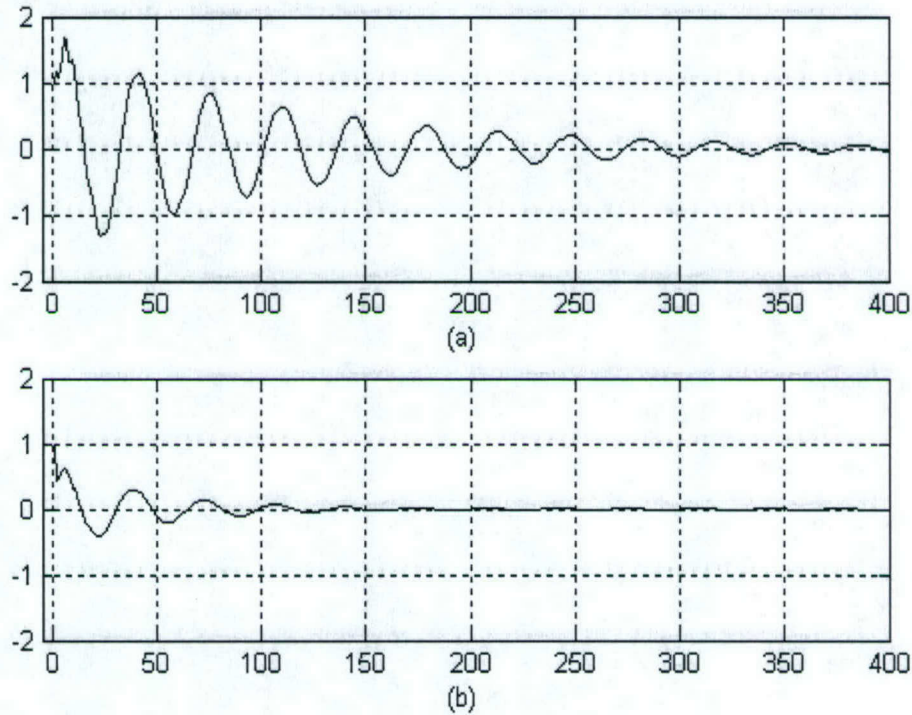


Figure 2-3: The effect of loading the auto-correlation vector on the impulse response of the prediction filter: (a) Impulse response with no loading, (b) Impulse response with loading.

The new approach, called loaded DPCM (L-DPCM), resolves these conflicting design goals for the prediction filter by monitoring the channel noise power to ensure that its impact on the reconstruction does not overpower the quantization noise. This is done by controlling the impulse response energy with a simple scalar addition σ_L^2 that loads the first term in the autocorrelation vector

Equation 2-8: Loaded Cross-Correlation Vector

$$\mathbf{r}_L[1] = \mathbf{r}[1] + \sigma_L^2$$

This loaded auto-correlation vector can then be used in the Levinson Durbin algorithm to calculate the prediction filter $\mathbf{w}$. Figure 2-3 (a) shows the impulse response of the prediction filter over one block of data. Figure 2-3 (b) shows the impulse response after loading. Clearly the loaded impulse response has less energy. Loading effectively eliminates the energy contribution due to the weaker correlations by making them even weaker relative to the zero-lag correlation. In this way the technique provides a simple adaptive solution to modify the prediction filter's impulse response energy, so that it minimizes the reconstruction error.

The authors believe this is the first technique to use loading to improve the compression performance of DPCM. Historically, loading has been applied to regularize the data covariance matrix in cases of low sample support. In these cases, loading the covariance matrix makes the adaptive filters less likely to over-adapt to spurious noise peaks that were not sufficiently averaged out due to low sample support. By limiting adaptation to noise peaks, the filter's performance is improved on independent, but statistically similar data. However, loading the DPCM prediction filter has a different purpose in the compression application. Here the filter is applied to the same data that was used for training, so there are no spurious noise peaks. In this case, the purpose of loading is not to obtain a better statistical estimate of the true covariance by covering up the noise correlations. Rather, the purpose is to limit the contribution of the channel noise and ensure that it is not the dominant factor in the reconstruction error.

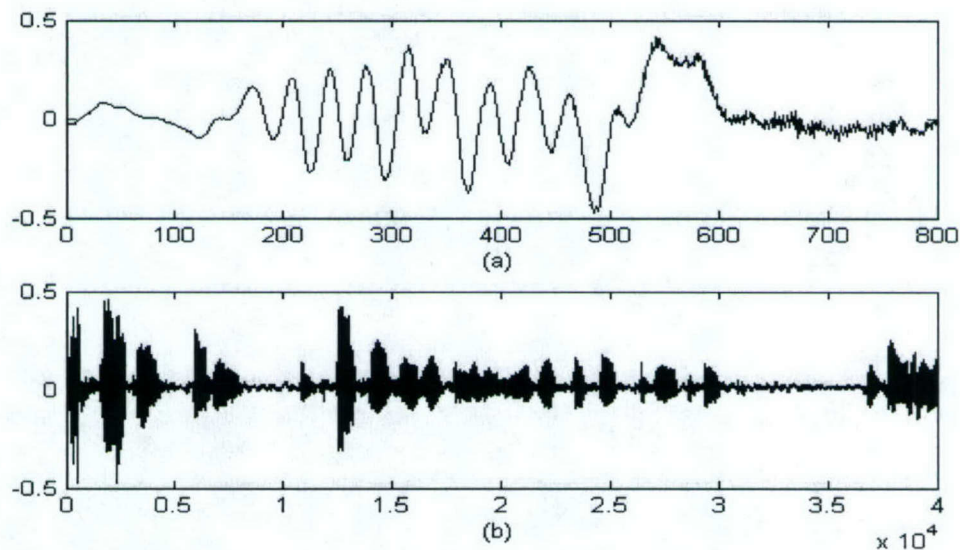


Figure 2-4: Times series of the speech data on which the performance analysis is based. (a) Short section of speech data, (b) Long section of speech data.

Performance Results and Analysis

L-DPCM's performance analysis is assessed on recorded speech sampled at 8 kHz, and analyzed for different channel conditions and quantization levels. Varying channel noise conditions are simulated by adding white Gaussian noise of different variances σ_c^2 to the data. The different channel conditions are expressed as the ratio of the channel noise power σ_c^2 to the signal power σ_s^2 . Data blocks of 400 samples are modeled with a 10th order prediction filter. The reconstruction error powers are plotted as a function of loading level and averaged over these blocks. Both the reconstruction error and the

loading level are also normalized by the signal variance σ_s^2 and expressed in dB as $10 \times \log_{10}(\sigma_R^2 / \sigma_s^2)$ and $10 \times \log_{10}(\sigma_R^2 / \sigma_L^2)$. To demonstrate the potential of the new technique, the results are first compared over a short 2-block section of voiced speech in which there is a significant advantage to performing DPCM in a noise free environment. Later the technique will be applied to a longer data stream consisting of both voiced and non-voiced segments. The respective data sequences are shown in Figure 2-4 (a) and (b).

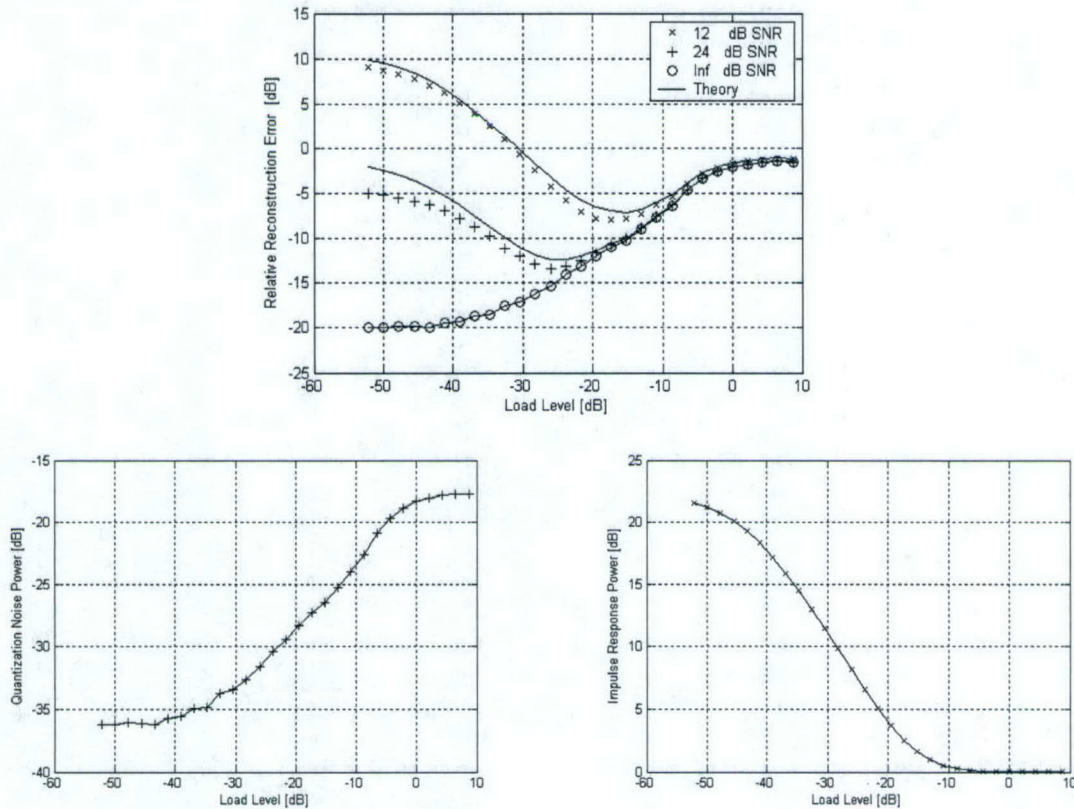


Figure 2-5: (a) Reconstruction Error vs. loading level for 2 bit quantization over short voiced sequence in three AWGN channels, (b) Quantization Noise vs. Loading, (c) Impulse Response Energy vs. Loading.

Figure 2-5 (a) plots the reconstruction error vs. loading level in three different channel environments over the short speech segment. Two bits are used to quantize the residual. The theoretically predicted values, calculated using Equation 2-6, are indicated by the solid lines. The results show close agreement between the theoretically predicted and empirically observed reconstruction errors. This allows the transmitter to accurately select the loading level to minimize the reconstruction error.

Selecting the appropriate loading level also enables the new L-DPCM technique to achieve the same performance as both DPCM in low noise channels and PCM in noisy channel conditions. Indeed, it is interesting to note that the L-DPCM algorithm with a

loading level of zero has the same performance as the conventional implementation of DPCM, in which no modifications are made to the prediction filter. The performance of standard DPCM is therefore approximated by the left-most reconstruction points in Figure 2-5 (a), which correspond to very low loading levels. For sufficiently large loading levels, the performance of L-DPCM approaches that of PCM. This is because the loading effectively modifies the autocorrelation vector in such a way that there appears to be no correlation between consecutive data samples. In this case the predicted value $\hat{x}[n]$ is zero, causing the residual to become the full input data, and the signal is passed unaltered to the quantizer. The points on the far right of the curves, which correspond to high loading levels, therefore approach the reconstruction performance obtained if PCM is used to quantize the source.

The curves in Figure 2-5 (a) indicate both the performance incentives and drawbacks to using DPCM in different noise environments. For example the lower curve of Figure 2-5 (a) demonstrates the incentive to use DPCM to compress data in a low noise channel. The plot indicates that DPCM prediction decreases the reconstruction error by 18 dB over PCM's simple quantization performance. Figure 2-5 (a) also shows the dangers of using DPCM in noisy channels. A reconstruction error of -20 dB in the noise free channel increases to 10 dB in a channel with $\sigma_c^2 / \sigma_s^2 = 0.068$.

The reason for this dramatic drop in performance can be explained in terms of the contributions of the quantization noise and the channel noise as a function of the loading. Figure 2-5 (b) plots the quantization noise as a function of the loading level and demonstrates how increased loading increases the quantization noise. This observation is indicative of the fact that loading the auto-correlation vector decreases its sensitivity to weaker correlations which results in a loss of the prediction filter's accuracy. This loss in turn results in a larger residual with increased dynamic range and causes an associated increase in the quantization noise. Thus the increase in quantization noise as a function of loading that is shown in Figure 2-5 (b) suggests and the noise free lower curve of Figure 2-5 (a) confirms that increased loading in a noise free environment will degrade the reconstructed signal.

Conversely, increased loading in a noisy environment improves the reconstruction performance. This can be explained by how the impulse response energy of the prediction filter is controlled by loading. Recall that the impact of channel noise is proportional to the impulse response energy of the prediction filter. Any noise corrupting the residual also gets amplified, proportional to the energy of the impulse response. A large signal compression is associated with a large impulse response energy resulting in a large reconstruction error, if noise corrupts the residual. To reduce the prediction filter's sensitivity to channel noise, its impulse response energy can be decreased by loading. Figure 2-5 (c) plots the impulse response energy of the prediction filter as a function of loading. It shows how loading causes the impulse response energy of the filter to decrease and thereby reduce its sensitivity to noise. By monitoring the noise environment a loading level can be selected so that the impact of the channel noise will never dominate the quantization noise. Thus the technique ensures that the performance in

noisy channels will never be worse than the noise-robust PCM encoding, and that the performance in quiet channels can be as good as DPCM.

Additionally Figure 2-5 (a) shows that selecting the optimum loading level can offer an improvement over either the standard DPCM or the PCM implementations in moderate noise conditions. Specifically, it shows that for $\sigma_c^2 / \sigma_s^2 = .004$ neither the standard implementations of DPCM or PCM are optimum and L-DPCM can offer a 10 dB improvement. In this region, the noise contributions of both the quantization noise and the channel noise are roughly equivalent and the channel noise is no longer the dominating contributor to the reconstruction error.

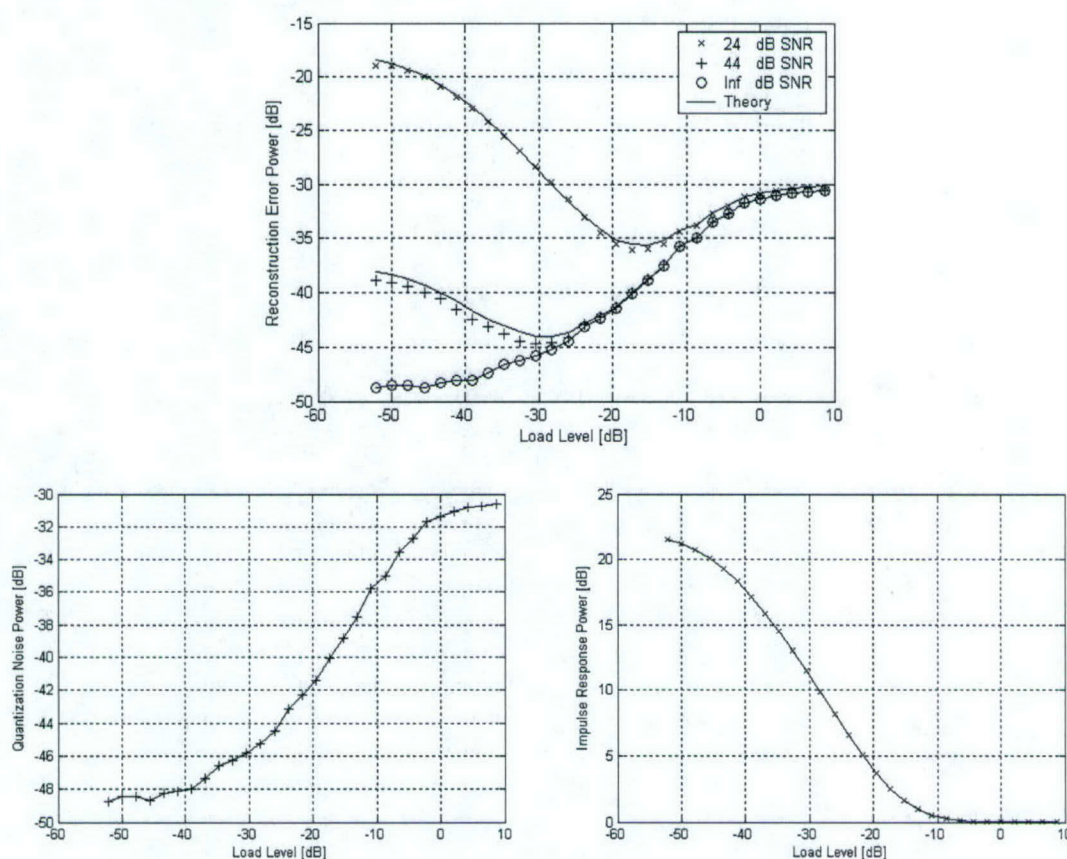


Figure 2-6: (a) Reconstruction Error vs. loading level for 4 bit quantization over short voiced sequence in three AWGN channels, (b) Quantization Noise vs. Loading, (c) Impulse Response Energy vs. Loading.

These same trends can be observed in smaller noise environments when more bits are used to quantize the residual. Figure 2-6 (a) shows a similar plot of the reconstruction error with 4 bit uniform quantization when AWGN sequences with relative noise ratios of

$\sigma_c^2/\sigma_s^2 = [0.004, 0.00004 \text{ and } 0.00]$ are added to the short voiced segment shown in Figure 2-4 (a). In this example the use of 4 bits decreases the quantization noise but also makes the compression that much more sensitive to any channel errors. Here even a noise ratio of $\sigma_c^2/\sigma_s^2 = 0.00004$ which corresponds to a -44 dB noise-to-signal ratio can have a significant impact on 4 bit DPCM reconstruction performance. This is because the impulse response again integrates the -44 dB channel noise up to a -20 dB reconstruction error. By loading the auto-correlation vector one limits the reconstruction amplification to ensure that the channel noise is not dominating reconstruction performance. The top curve of Figure 2-6(a) indicates that when $\sigma_c^2/\sigma_s^2 = 0.004$ corresponding to a -24 dB channel noise, and the conventional implementation of 4 bit DPCM fails, L-DPCM

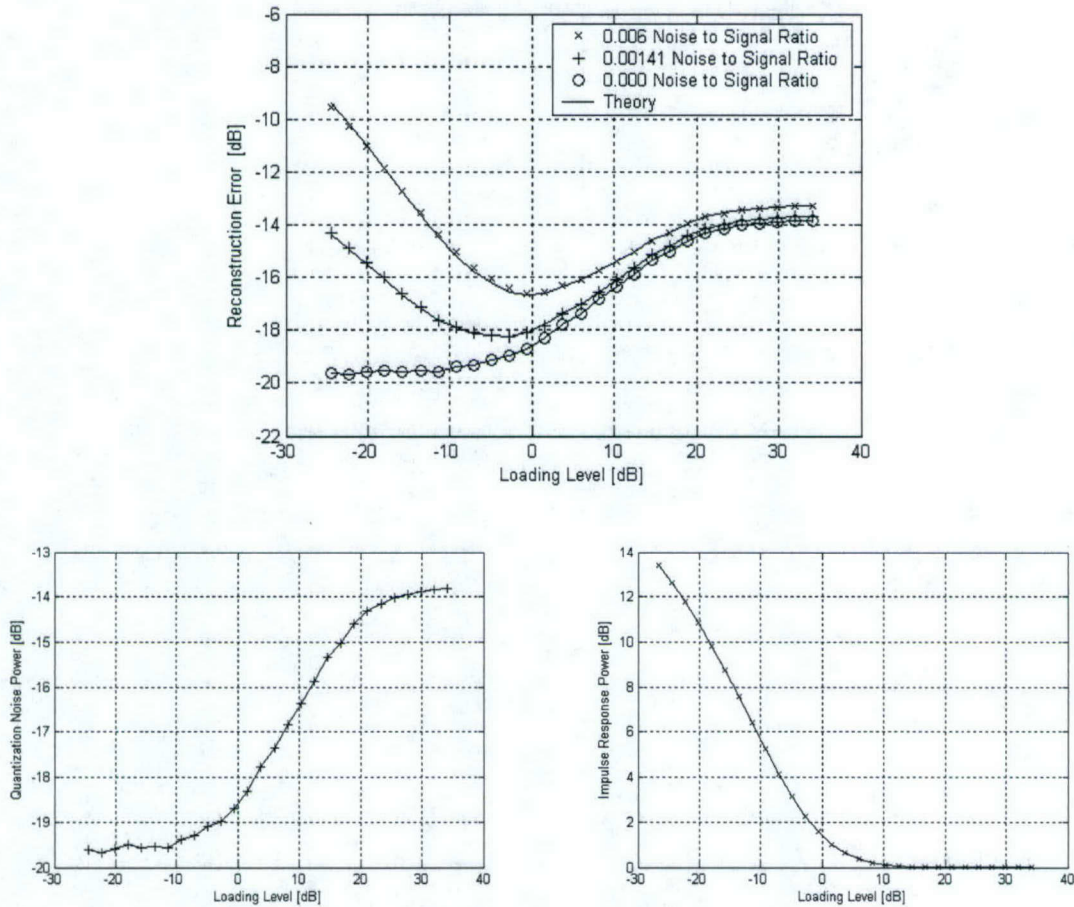


Figure 2-7: (a) Reconstruction Error vs. loading level for 4 bit quantization over long voiced sequence in three AWGN channels; (b) Quantization Noise vs. Loading; (c) Impulse Response Energy vs. Loading.

can offer some improvement over PCM encoding. These results show that the L-DPCM technique is applicable for a wide variety of channel conditions and quantization levels.

Finally Figure 2-7 plots the reconstruction error as a function of the channel noise over a the long series of speech containing 100 blocks of data. The residual is quantized to 4 bits and sent through the noise channels of $\sigma_c^2/\sigma_s^2 = [0.006, 0.00141, \text{ and } 0.00]$. The overall improvement of DPCM offers over PCM on the long blocks of speech is reduced to 6 dB since portions of unvoiced segments where linear prediction does not offer a significant improvement are averaged into the results. Still 6 dB represents on average 1 bit per sample less for the same reconstruction quality. Within this range the plot shows that selecting an optimum loading level improves the performance of L-DPCM over both DPCM and PCM in the two noise channels shown. In moderate noise conditions the plot shows that finding the optimum loading level can result in approximately a 4 dB improvement over either the standard DPCM or the standard PCM implementations. This represents a significant fraction of the 6 dB improvement offered by DPCM in a noise free channel. However if DPCM was used without this loading its performance would be no better than the simply quantized reconstruction quality of PCM.

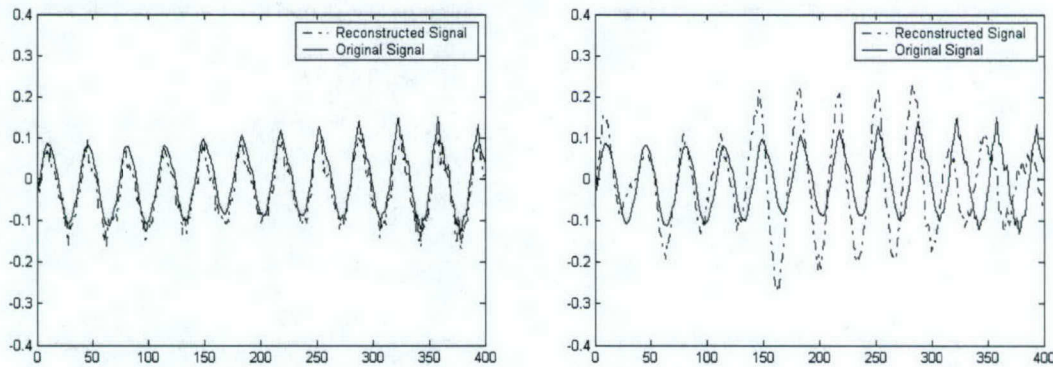


Figure 2-8: Comparison of original and reconstructed signal with a low loading level of -28 [dB] (left) with a high loading level 8[dB] (right) in a noise environment of $\sigma_c^2/\sigma_s^2 = .564$.

Some additional insight into the performance afforded by the technique can be gained by observing the reconstruction over a single block of data. Figure 2-8 (b) plots the original and reconstructed signal for a single block of data in a relatively high noise environment. No loading is applied. It visually demonstrates DPCM's poor reconstruction performance in the presence of channel noise. Figure 2-8 (a) plots the original and reconstructed signals for the same block of data in the same noise environment when loading is applied. The results show the improvement that loading can have on a DPCM reconstructed signal.

Overall these results illustrate that the performance gains of DPCM can be maintained by integrating knowledge of the channel noise directly into the source compression algorithm so that the prediction filter can be optimized for directly minimizing the reconstruction error instead of the residual error. These results also show that a scalar loading of the auto-correlation vector allows for a computationally simple way to get the best reconstruction error over a wide range of noise environments.

Conclusions

This paper demonstrates a simple technique that combines knowledge of the channel conditions into the source compression algorithm to minimize the overall reconstruction error. The analysis characterizes the reconstruction error in terms of the impulse response of the predictive filter and the quantization noise. Loading the prediction filter is shown to lower its impulse response energy and make it less sensitive to channel noise while increasing its quantization noise. The technique balances these conflicting noise sources by loading the prediction filter to minimize the overall reconstruction error.

Publications

- 1) H. Witzgall, W. Ogle and J. Goldstein, "Mitigating the effects of Channel Noise in Source Compression with Reduced Rank Processing," in *Proceedings 37th Asilomar Conference Signals Systems Computing.*, Pacific Grove, CA. November 2003.
- 2) H. Witzgall, and J. Goldstein, "Reduced Rank Predictive Source Coding," in *Proceedings 2003 IEEE Workshop on Statistical Signal Processing.*, St. Louis, MO September 2003.
- 3) H. Witzgall, and J. Goldstein, "On Optimizing the DPCM Prediction Filter for Minimum Reconstruction Error in Noisy Channels," *submitted to IEEE Transactions on Signal Processing*, November 2004.
- 4) H. Witzgall, W. Ogle and J. Goldstein, "Loading Levinson-Durbin to mitigate channel noise in DPCM Source Sompresion," *SAIC Internal Tech Paper*, September 2004.

Chapter 3 Integrated Sensing and Processing for Fast Adaptive Modulation

The Fast Adaptive Modulation task will demonstrate a novel method of nonlinear optimization that can be applied to end-to-end optimization of a radio communications system. This technique will optimize the end-to-end performance of a pair of CDMA radio transceivers on rapidly moving platforms in a dynamic multipath and Doppler spreading environment. Fast Adaptive Modulation (FAM) dynamically changes the modulation of a radio link to provide better performance in changing environments and loads, as a mission evolves.



Figure 3-1: Feedback of a limited amount of information from the receiver to the transmitter allows the transmitter to modify the codes used to represent the data.

Feedback of a limited amount of information from the receiver to the transmitter allows the transmitter to modify the codes used to represent the data, so as to present an easily decoded signal to the receiver. For this process to be efficient, only a limited amount of data can be fed back. The optimization procedure allows the estimation of a small number of parameters that characterize the combined effects of variations in clock rates that depend on components in both radios as well as the motion of the platforms and number and types of scatterers in the propagation environment.

One important potential application of this technique is to defeat algorithms that attempt to recover the spreading code of a signal. Such algorithms generally require the length of the code in time to be nearly constant. These algorithms could be defeated by smoothly varying the chip length during the frame, effectively applying a warbling Doppler shift to the signal. The intended receiver, knowing the spreading code can estimate the warbling Doppler shift and compensate, an adversary would not be able to do so.

A Novel Non-linear Optimization Approach: Least-Logs

The optimization methodology applied here is optimization in “least-logs” rather than the more familiar “least-squares” optimization. Least-logs optimization has the advantage that it is less sensitive to outliers than least-squares optimization. As such, it can be applied to curve-fitting and curve-finding applications in noisy images and displays. In principal, it can be applied to fitting surfaces in high dimensional optimization problems,

but as this is difficult to visualize, the application demonstrated here is restricted to curve fitting in two dimensions. The application developed here builds on the author's "Process and Apparatus for the Automatic Detection and Extraction of Features in Images and Displays," Patent #4,906,940.

This methodology measures features, not pixels, with the resulting SNR gain due to integration over the feature in question, a curve in this case. This gain results from a curious property of the logarithm function. For taking the x, or horizontal, partial derivative of $\log(r)$:

$$\frac{\partial}{\partial x} \ln(r) = \frac{\partial}{\partial x} \ln(\sqrt{x^2 + y^2}) = \frac{x}{x^2 + y^2},$$

the resulting function differs markedly from zero only over a limited region, the red and blue dots in figure 2(a). But when convolved with a vertical line shown as $x=50$ in figure 2(b), the resulting integral:

$$\int_{-\infty}^{\infty} \frac{x}{x^2 + y^2} dy = \arctan(y) \Big|_{-\infty}^{\infty} = \text{sign}(x) \cdot \pi$$

is, except for the sign, independent of x . It is in fact a step function, with a jump at the position of the line. The result of the (finite) convolution of the derivative of the logarithm in Figure 3-2(a) and the vertical line in Figure 3-2(b) clearly shows this effect in Figure 3-2(c). The negative region to the left of the line (red) is extended not only along the line, but also blooms for an extended region to the left of the line, and similarly for the positive (blue) region to the right.

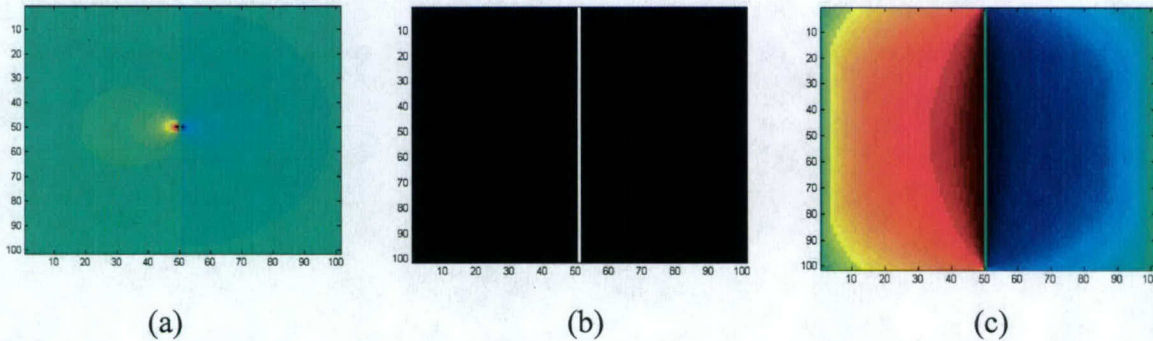


Figure 3-2: (a) The horizontal partial derivative of $\log(r)$, red is negative, blue is positive, green is near zero. (b) A white vertical line of ones against a black background of zeros. (c) The convolution of (a) with (b), note that the positive and negative regions in (c) have expanded not only along the line, but away from the line as well, a special property of the logarithm.

Following the negative gradient (left or right), from any point will necessarily lead to the line. Linear features can also be markedly enhanced relative to a random noise background, as will be shown subsequently.

Background: Multipath Spreading

One major issue that a tactical radio system has to deal with, as a mission evolves, is a radically changing multipath spreading environment. The major effect of spreading is that different versions of the same symbol can show up at different, apparently random, times due to propagation over several unrelated paths. One practical effect of this is that adjacent symbols will interfere with each other, since the receiver responds to these symbols will overlap. When the signals lie in a narrow frequency band, adjacent symbols may cancel each other entirely, the so-called “flat-fading” channel. This study will explore spread spectrum systems, which artificially spread the signal over such a wide bandwidth that only a portion of the signal can be cancelled at any given time, the so-called “frequency-selective-fading” channel.

Mitigation of Inter-Symbol Interference due to Multipath

In the modulation scheme proposed here, a different spreading code is associated with each symbol to be transmitted. In the simplest case of transmitting zeros and ones, one spreading code would be transmitted for a zero and a second code would be transmitted for a one. The spread spectrum receiver decodes a transmission by correlating the spreading codes with the signal and selecting the symbol corresponding to the spreading code with the strongest correlation. However, a single transmitted code, propagating over several paths, may produce several strong correlations, as shown in Figure 3-3. One major source of errors in a spread spectrum system is timing errors, due to confusion over the exact arrival time of one or more of the correlation peaks.

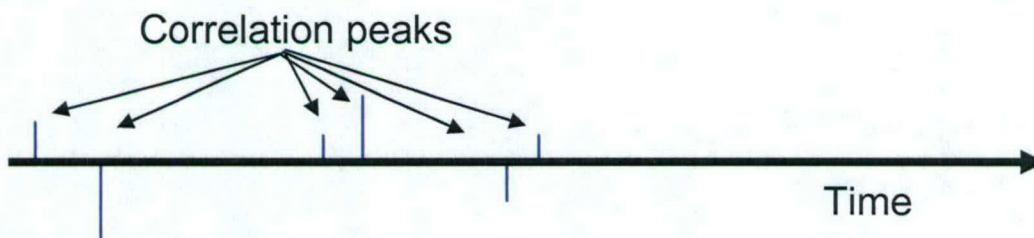
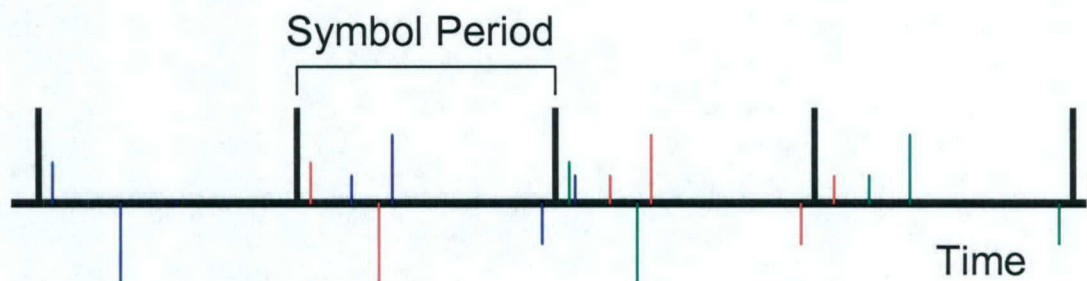


Figure 3-3: Multipath can create several correlation peaks when the signal is despread by correlation.

When the spreading of the signal is more than the length of the spreading code (or corresponding symbol), then correlations from one symbol will appear later, in the time interval associated with another symbol. Correlations from one symbol may even obscure or cancel correlations from another symbol. This produces the second major source of error in spread spectrum systems, Inter-Symbol Interference (ISI). This phenomenon is illustrated in Figure 3-4, in which correlations from the first (blue) symbol, run over into the time interval of the second (red) symbol, and the third (green) symbol.



Interleaved correlation peaks from multiple symbols

Figure 3-4: The correlation peaks from consecutive symbols can overlap, when the total multipath spread is larger than the symbol period.

One way to mitigate Inter-Symbol Interference is to determine over how many symbol periods that the signal is spread, and rotate through a corresponding number of sets of spreading codes. This will de-interleave the symbols, at the cost of requiring an increase in the number of correlators in the receiver. For the current example, if the propagation channel spreads the symbol response over three symbol periods then the transmitter rotates through three mutually independent pairs of spreading codes. For the first symbol, the appropriate code is chosen from the first pair, for the second symbol, a code from the second pair, for the third symbol, a code from the third pair. But at the fourth symbol, the multipath from the first symbol will have died out, so that the fourth symbol may be coded with a code from the first pair. Figure 3-5 shows the output of three correlators, one for each of the three successive spreading codes. Note that the outputs are similar but lag each other by one symbol period. If the correlator outputs for the three codes are shifted by a lag of one symbol period, so that the beginnings of successive symbol periods line up, the multipath arrivals line up in a “lag display” as shown in Figure 3-6.

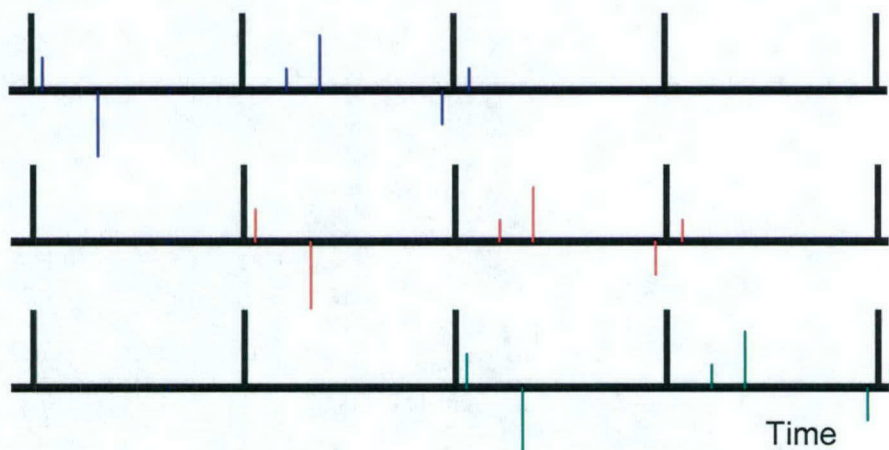


Figure 3-5: Using different spreading codes for each symbol period will de-interleave the multipath arrivals from each symbol.

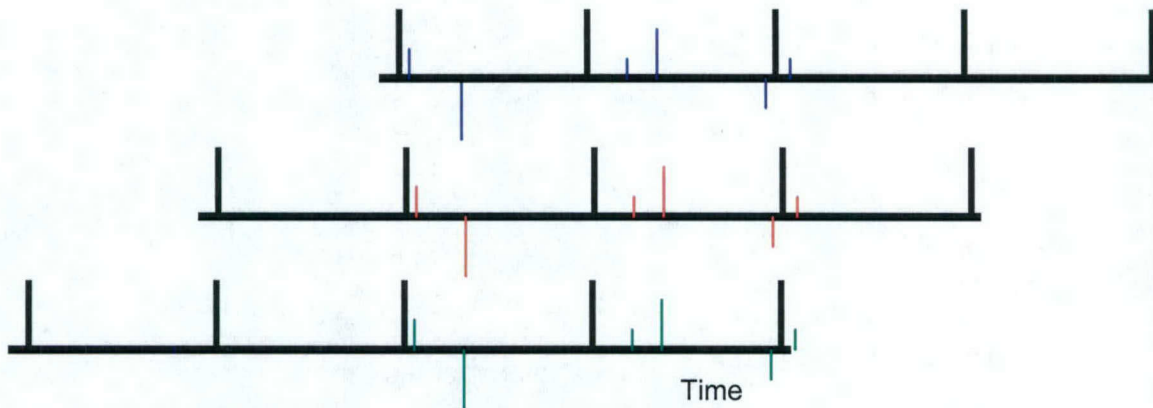


Figure 3-6: A lag display is formed by time shifting the output from successive correlators by the length of one symbol. Each arrival path will form a curve in the lag display.

A simulated lag display, which rotates through three spreading code pairs, is shown in Figure 3-7. The symbol length is .1 microsecond, corresponding to a bit rate of 10 MHz. The display is six symbol periods wide. The first two rows of the lag display are outputs from the zero and one correlators of the first pair. The next two rows are from the second pair, left shifted by one symbol period. The next two rows are from the third pair, left shifted by one additional symbol period. The seventh and eighth rows return to the first pair again.

The second, third, and fourth light-blue vertical lines in the lag display correspond to three multipath arrivals, bit aligned in time (vertically) and spread over three symbol periods (horizontally). The first, fifth, and sixth lines correspond to time-shifted versions (vertically) of the middle three.

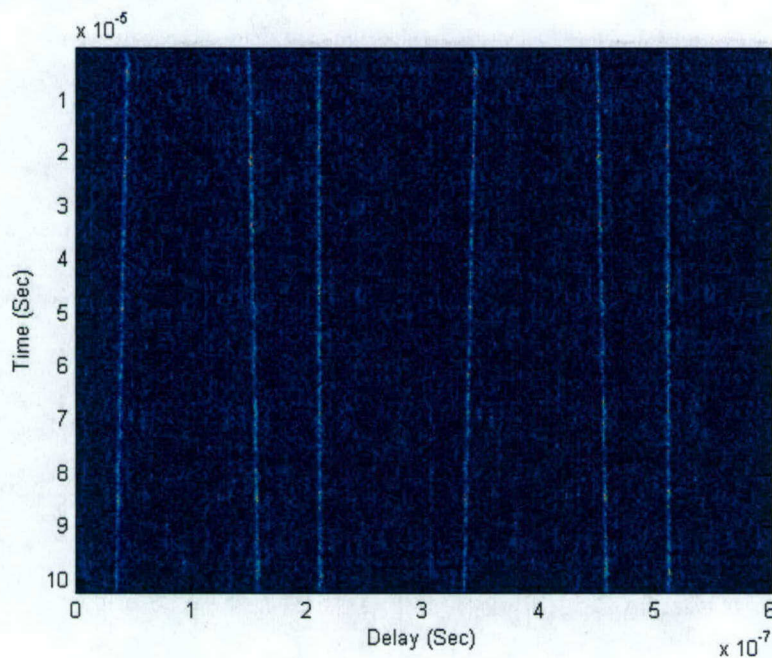


Figure 3-7: A simulated version of the display using three pairs of spreading codes.

The same simulation shown in Figure 3-7, but using only one pair of spreading codes, is shown in Figure 3-8. There are now three times as many lines since the correlator in each row produces a response not only for the current symbol, but also for the two previous symbols. This is because the multipath spreading spreads each symbol over three symbol periods. There are also increased side lobe levels as well as non-recoverable Inter-Symbol Interference near the top of the display, where the lines cross.

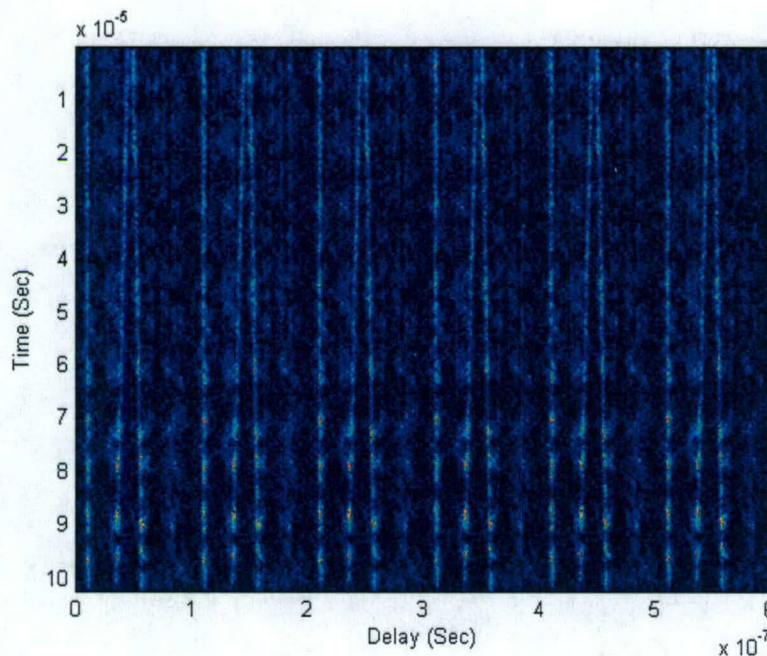


Figure 3-8: A simulated version of the lag display using only one pair of spreading codes. Note the threefold increase in correlations, increased side lobe background level, and Inter-Symbol Interference at the top where the lines cross.

The data for the simulations was generated using three pairs of 31-chip Gold codes. A frame consisted of 1023 bits, with a bit rate of 10 MHz. The bits were encoded by rotating through the three pairs of Gold codes, selecting the first code from the pair for a zero, and the second code from the pair for a one. Doppler was simulated by calculating the variation in delay due to the time varying change in distance from a source (T) to a receiver (R) via three point-reflectors, as the receiver moved between the reflectors, as shown in Figure 3-9. The Doppler correction was applied to the signal codes by resampling and to a 1 GHz carrier by direct calculation. In addition, independent clock rate errors for the bit rate and carrier frequency were applied, 5 parts per million for the receiver and .5 parts per million for the transmitter. The slope of the lines in the display is due to a combination of all these effects. The slope is exaggerated for illustration by using a receiver speed of 60,000 knots. Of course, the slopes of the lines are different depending on the rate at which the receiver is approaching or moving away from each reflector.

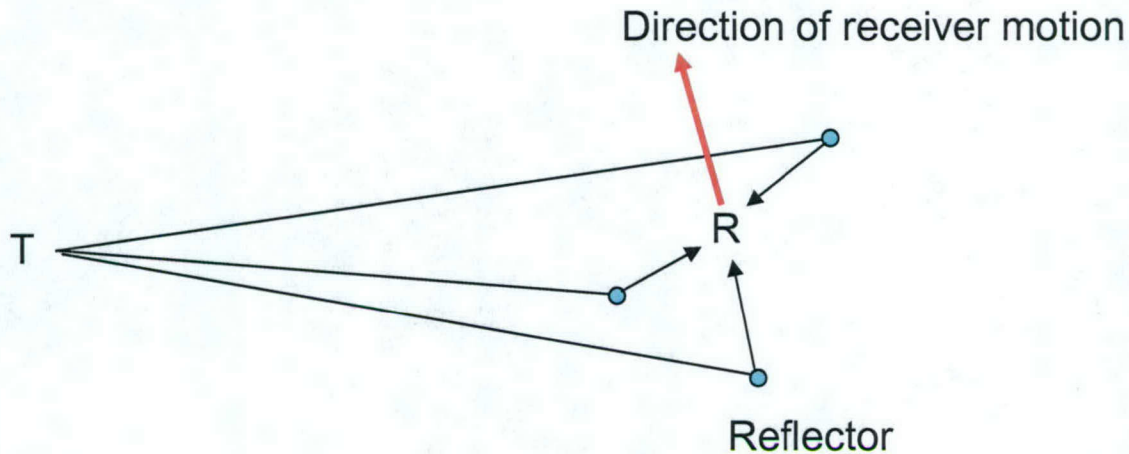


Figure 3-9: Doppler was simulated by calculating the variation in delay due to the time varying change in distance from a source (T) to a receiver (R) via three point reflectors, as the receiver moves between the reflectors.

Least Logs Optimization: Mitigation of Timing Errors

In order to build a coherent Rake Receiver, a very accurate estimate of the time delay properties of each multipath arrival is required. This can be accomplished by fitting a straight line to each curve in the lag display. This curve-fitting uses least-logs optimization. It is accomplished by convolving the lag display with the x-derivative of the logarithm, shown in Figure 3-2(a). The result is a gradient field with a step function, switching from positive to negative at the lines in the display, shown as jumps between red and blue in Figure 3-10.

The curve-fitting is carried out using linear templates that move in the gradient field using, for instance, equations for the motion of a rigid rod in a force field. In this case, the rod is allowed to translate and rotate. Qualitatively, if the template is mostly in the blue, it will move to the left, mostly red to the right. Similarly, if the top of the template is in the blue and the bottom is in the red, the template will rotate counterclockwise, and vice versa. The initial positions of the templates are shown in Figure 3-11. The converged positions of the templates are shown in Figure 3-12. At least one template has converged to each of the lines.

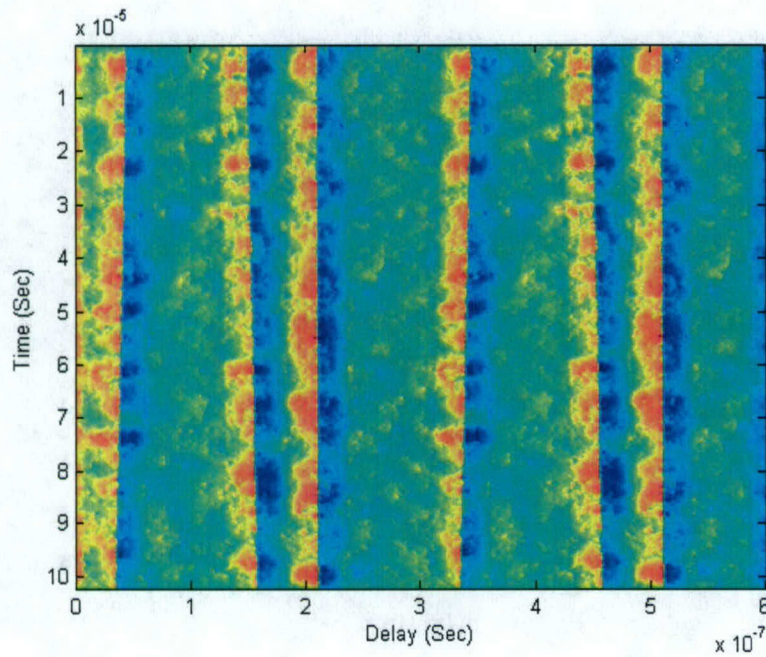


Figure 3-10: The gradient of the $\log(r)$ potential formed by convolving the function in Figure 2(a) with the lag display.

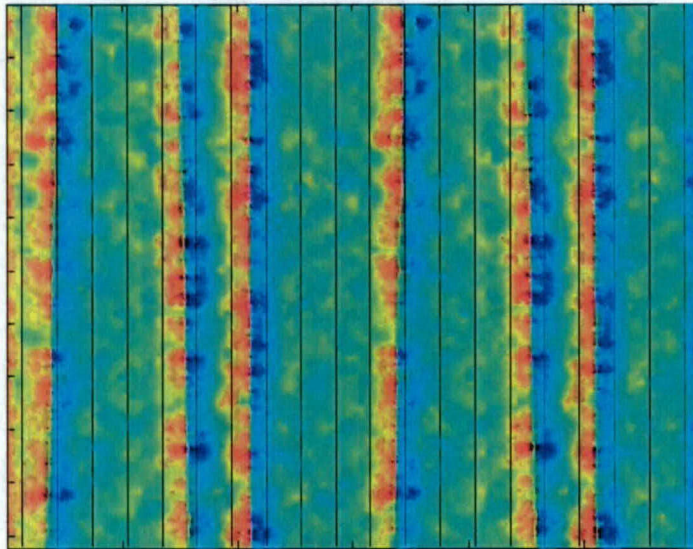


Figure 3-11: Linear templates translate and rotate in the gradient field. Points on the template in the blue tend to move to the left, in the red to the right. Translational and rotational moments on the templates can be calculated by integrating the gradient field along the templates.

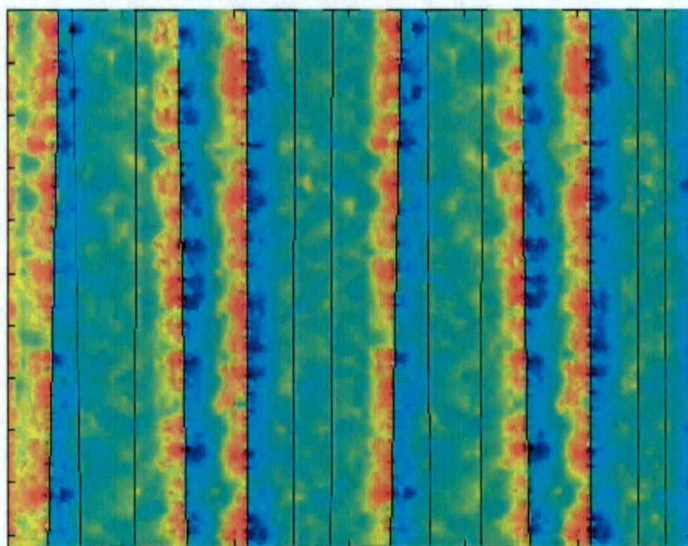


Figure 3-12: At least one of the templates converges to each of the lines in the lag display, located at the positive-negative (red-blue) boundaries in the gradient field.

Coherent Rake Receiver Implementation

Those templates which converge correspond to optimum trajectories of sampling times for the correlator outputs. Each template corresponds to one of the multipath arrivals. The curves in Figure 3-13 are the real and imaginary parts of the sampled correlator outputs along those trajectories. There are twice as many points in these curves as there are bits because there is a pair of correlators for each bit, one for a zero and one for a one. The correct correlations of each of the pairs should produce a string of large positive numbers, modulated by a slowly varying complex phase. The slowly varying phase is due to various clock errors. The incorrect one of each of the pairs should produce a zero mean complex random process. It is easy to see in Figure 13 that the larger values in the real and complex parts have sinusoidal behavior. This sinusoidal behavior may be analyzed by Fourier transform to estimate the time varying phase of the correlator output.

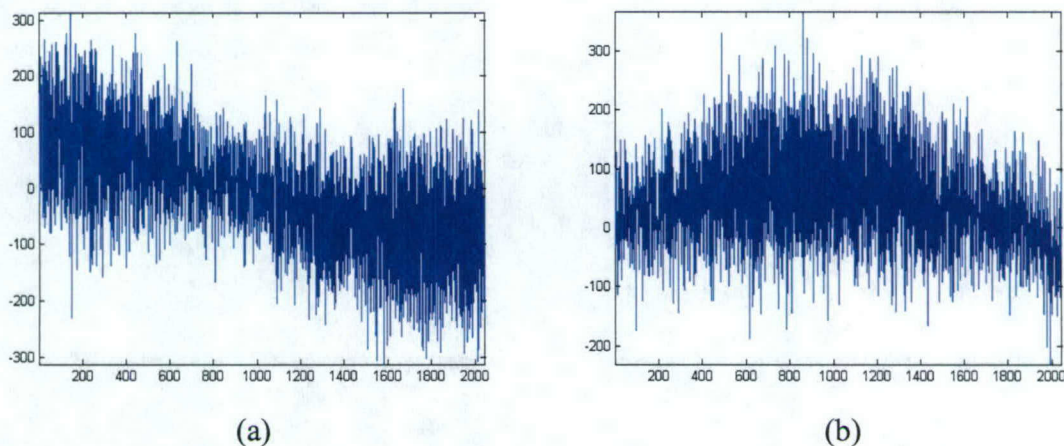


Figure 3-13: The real and imaginary parts of the sampled correlator outputs along the template trajectories. Note the sinusoidal envelope.

On correcting for this complex phase, the signals can be decoded by comparing the real parts of each pair of correlators. This gives a 3 dB improvement over comparing the complex magnitudes. If the real part of the first correlator is largest, it is decoded as a zero; if the real part of the second one is largest, it is decoded as a one.

At this point the complex correlator outputs from the different multipath arrivals can be coherently combined. For this example with three multipath arrivals at approximately equal levels, this combination will produce a 5 dB improvement in SNR. The only remaining problem is to select those templates whose outputs should be combined. This is accomplished by cross-correlating all the decoded outputs, then clustering those that are highly correlated. The sampled, phase corrected complex outputs of the largest cluster are optimally linearly combined, with weighting coefficients appropriate for the mean power in each sample.

Segments of a set of six decoded output sequences from six templates are shown in Figure 3-14. It is easy to see that the second, third, and fourth output sequences are very similar. They are not identical because the SNR on each individual path is low. The first output sequence is similar to the middle three, but shifted by one to the right. The fifth and sixth output sequences are similar to the middle three but shifted one to the left. As the largest cluster of similar sequences, the complex outputs of the second, third, and fourth sequences would be coherently combined.

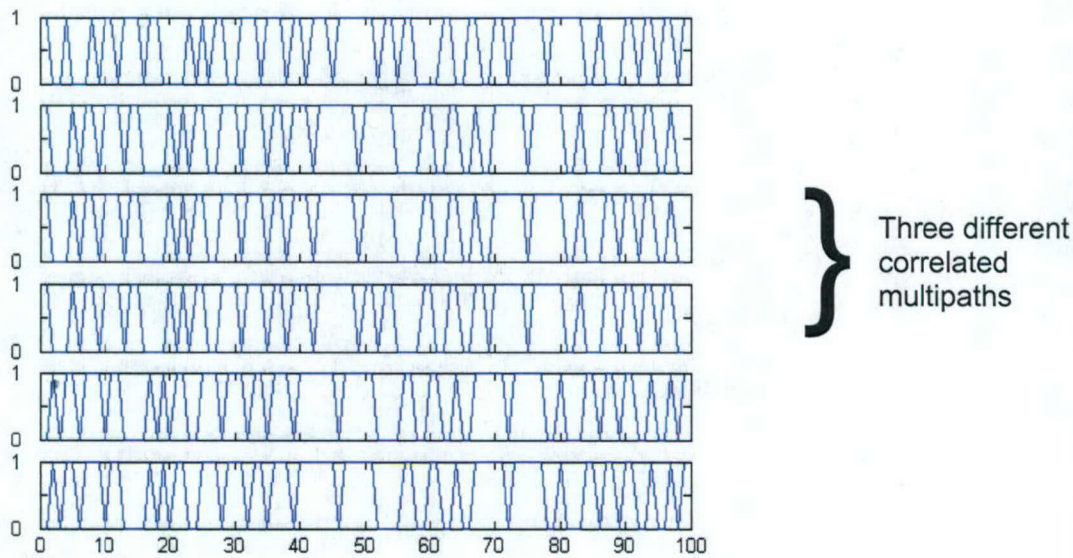


Figure 3-14: Segments of a set of six decoded output sequences from six templates. It is easy to see that the second, third, and fourth output sequences are very similar. They would be coherently combined.

The classifier that clusters the multipath output sequences then determines the modulation parameters used by the transmitter. Based on the analysis of the clusters of multipaths, the receiver can estimate the total spread of the multipath. When this information is fed back to the transmitter (Partial Channel Information in Figure 1), the transmitter can determine how many sets of spreading code pairs must be used to minimize Inter-Symbol Interference, namely the length of the total spread of the cluster in symbol periods.

FAM Concepts for Low Probability of Intercept Applications

The FAM approach outlined above has the capability to recover unknown timing information by curve fitting in a lag display formed from the spread spectrum correlator output. The process requires knowledge of the spreading codes in order to build the lag display and recover the timing information. An adversary, trying to exploit this signal would need to recover the spreading codes from the signal (or else resort to brute force). However, algorithms for recovering the spreading codes are generally based on exploiting the regular periodic retransmission of the spreading codes in the message. This begs the question: Can the intentional smooth variation of the timing (chip length) in the signal add enough irregularity to prevent recovery of the spreading codes?

The following example shows a periodic artificial Doppler-like warble applied to the signal. The warble depends on two unknown parameters, the amplitude, and the phase of a cosine function relative to the start of the frame.

The first example is received at -9 dB in white noise and appears at the correlator output at a level of 6 dB. The display of the absolute value squared of the correlator output data is shown in Figure 3-15. The presence of the signal is very hard to see. A very slight expression of several sinusoidal curves can be seen running vertically down the lag display. The goal is to accurately estimate the timing data for these curves and coherently combine the three resulting multipaths for an additional 5 dB of gain, decoding at +11 dB.

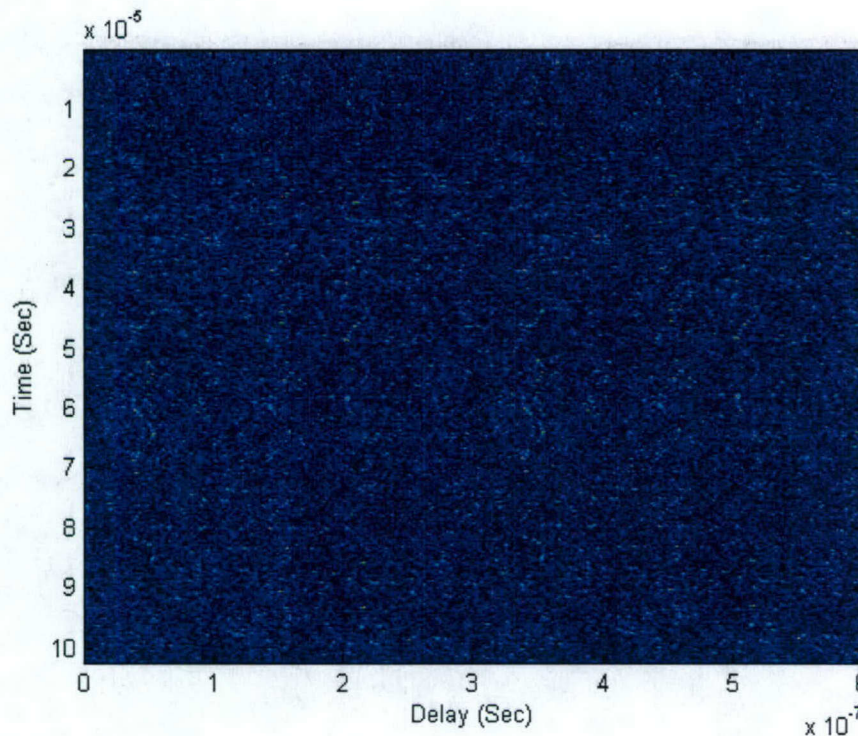


Figure 3-15: A lag display of the absolute value squared of the complex correlator output data. A very slight expression of several sinusoidal curves can be seen running vertically down the display.

An added benefit of the convolution with $\log(r)$ is that it can enhance linear features in noisy displays. In this example, the sinusoidal curves in the display can be enhanced by the convolution of the x-derivative of $\log(r)$ with the complex version of the lag display in Figure 15. The result shown in Figure 3-16 clearly shows the presence of sinusoidal structure at the red-blue boundaries in the real (left) and imaginary (right) parts of the convolution. This result is similar in effect to a Hough transform and results from the integrative properties of log convolution. An enhanced version of the lag display, shown in Figure 3-17, can be calculated by applying a horizontal difference of both displays in Figure 3-16, then taking the sum of the squares of each. Six multipath arrivals are now clearly seen in the new display, the second, third, and fourth forming the cluster of interest, as before.

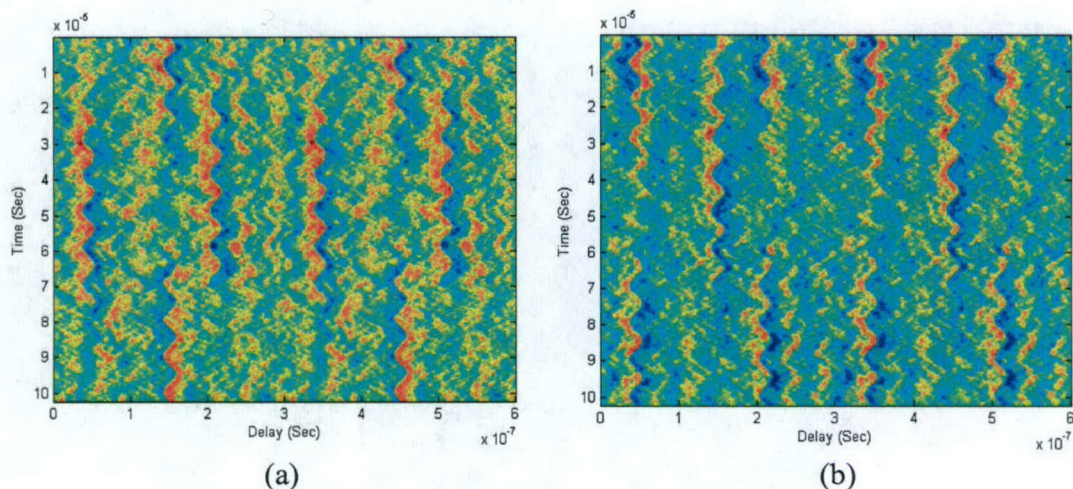


Figure 3-16: The sinusoidal curves in the lag display are enhanced by the convolution of the derivative of $\log(r)$ with the real (a) and imaginary (b) parts of the complex version of the lag display.

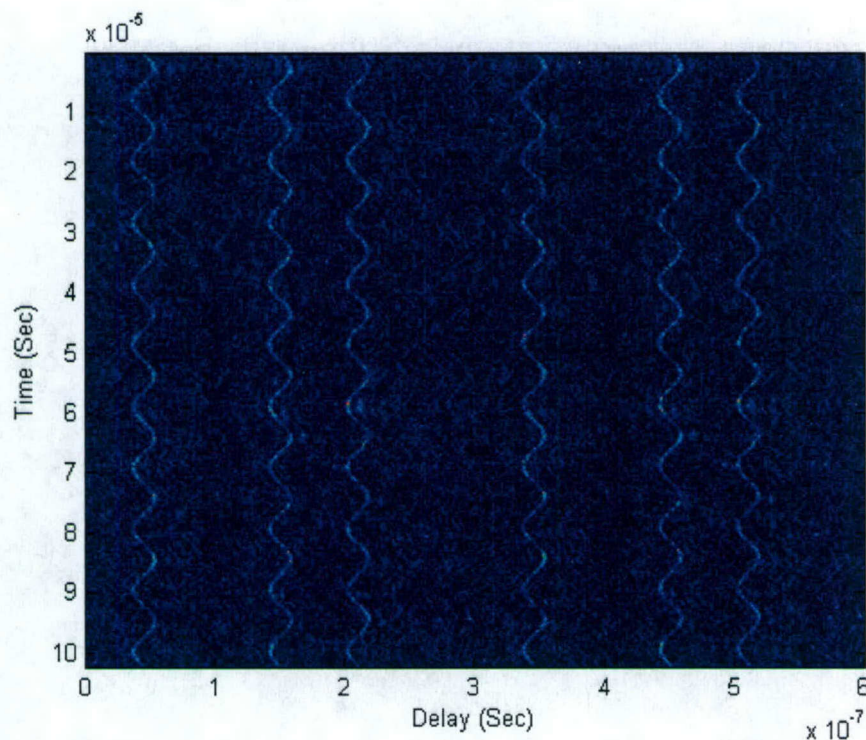


Figure 3-17: Six multipath arrivals are now clearly seen in the enhanced lag display. The second, third, and fourth arrivals form the cluster of interest, as before.

Templates with three parameters, delay with the amplitude and phase of the artificially induced cosine Doppler delay, easily converge to an accurate estimate of the timing, as shown in Figure 3-18. After decoding, the three multipaths in the cluster have error rates of 29, 48, and 47 out of 1023 bits. This is reduced to 0 errors after coherent combination of the three multipaths improves the signal-to-noise ratio to 11 dB.

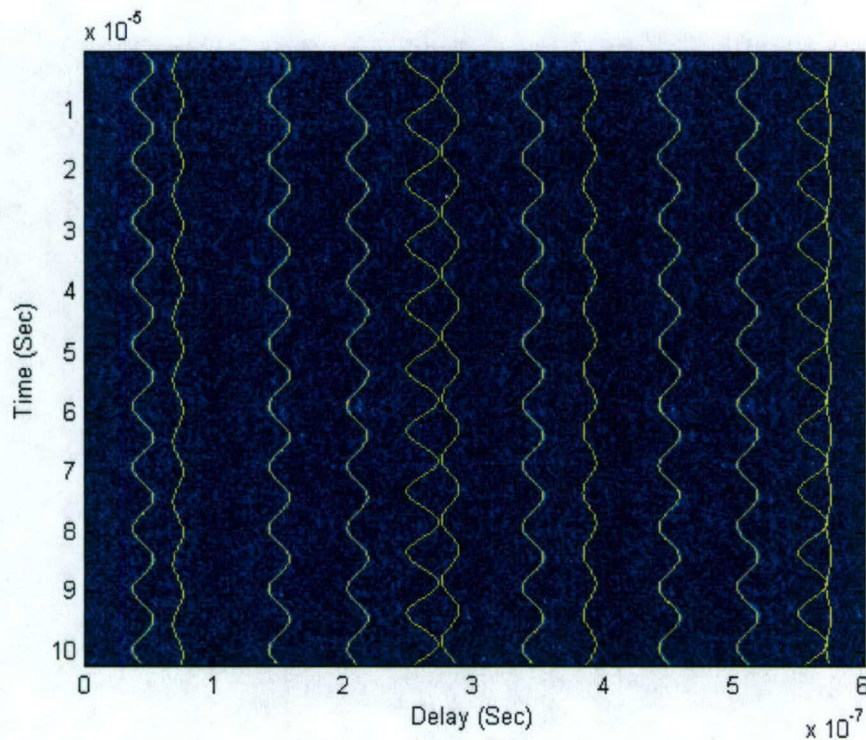


Figure 3-18: Templates with three parameters, delay with the amplitude and phase of the artificially induced cosine Doppler delay, easily converge to an accurate estimate of the timing.

An additional LPI application is the coherent combination of transmissions from several tight beam transmitters whose signals converge on a central location. This is shown in Figure 3-19, where five transmitters, labeled “T”, are sending tight beam signals at a receiver, labeled “R”. The concept here is that, outside the area of overlap, the signal from any one transmitter would be too weak to decode, or perhaps to detect.

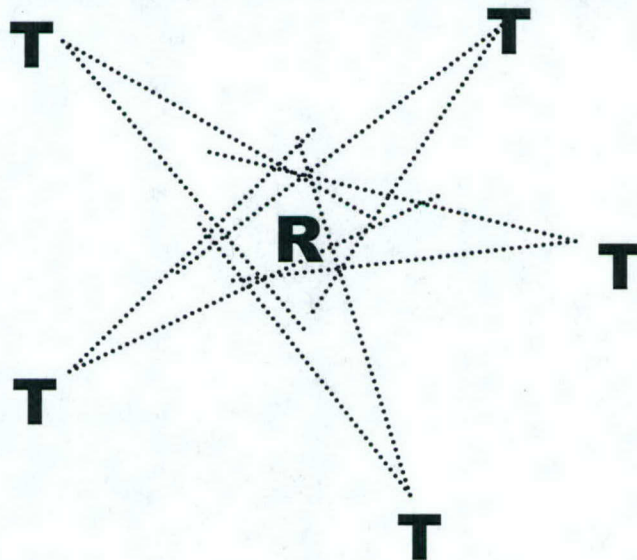


Figure 3-19: An additional LPI application is the coherent combination of transmissions from several tight beam transmitters whose signals converge on a central location.

The lag display for a scenario with seven multipath arrivals, received at -12 dB, and appearing at 3 dB at the correlator output is shown in Figure 3-20. The individual paths are difficult to see at this SNR level. The algorithm successfully enhances the seven multipaths in Figure 3-21. It then curve fits, as shown in Figure 3-22, clusters and coherently combines the seven paths to achieve an SNR of 11.5 dB. For this realization, the error rates for decoding on the seven multipath arrivals individually were: 116, 119, 122, 162, 147, 143, and 199 out of 1023 bits. After coherent combination, the error rate is reduced to 2 out of 1023.

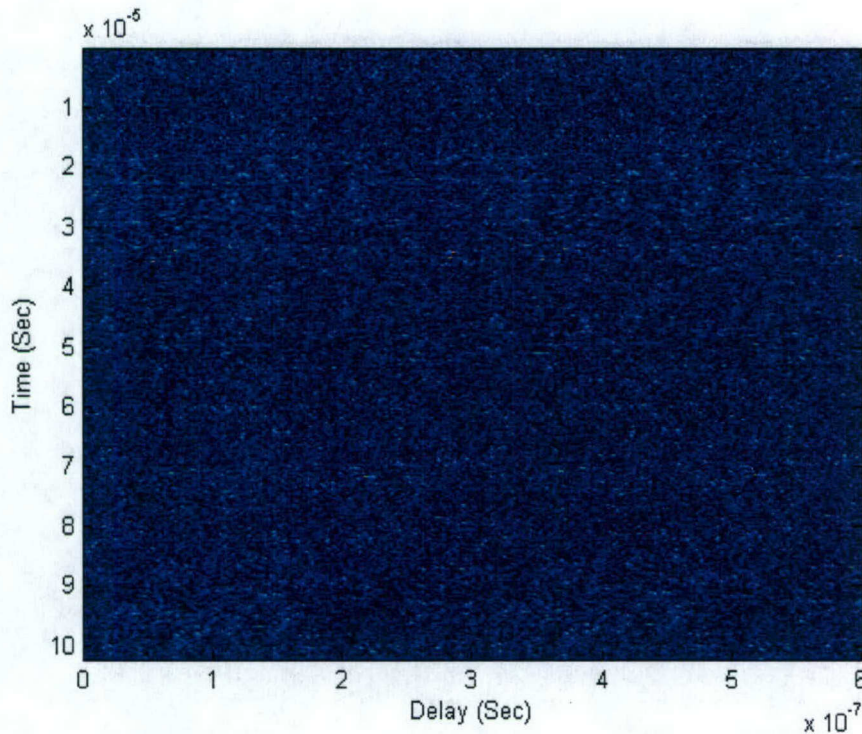


Figure 3-20: The lag display for a scenario with seven multipath arrivals, received at -12 dB, and appearing at 3 dB at the correlator output.

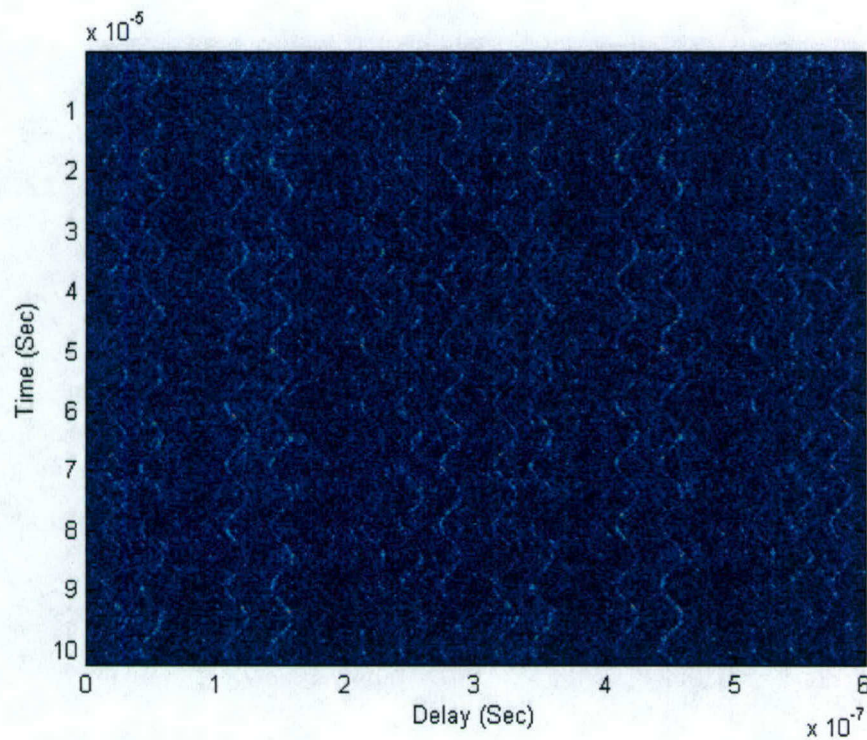


Figure 3-21: The algorithm successfully enhances the seven multipaths in the lag display.

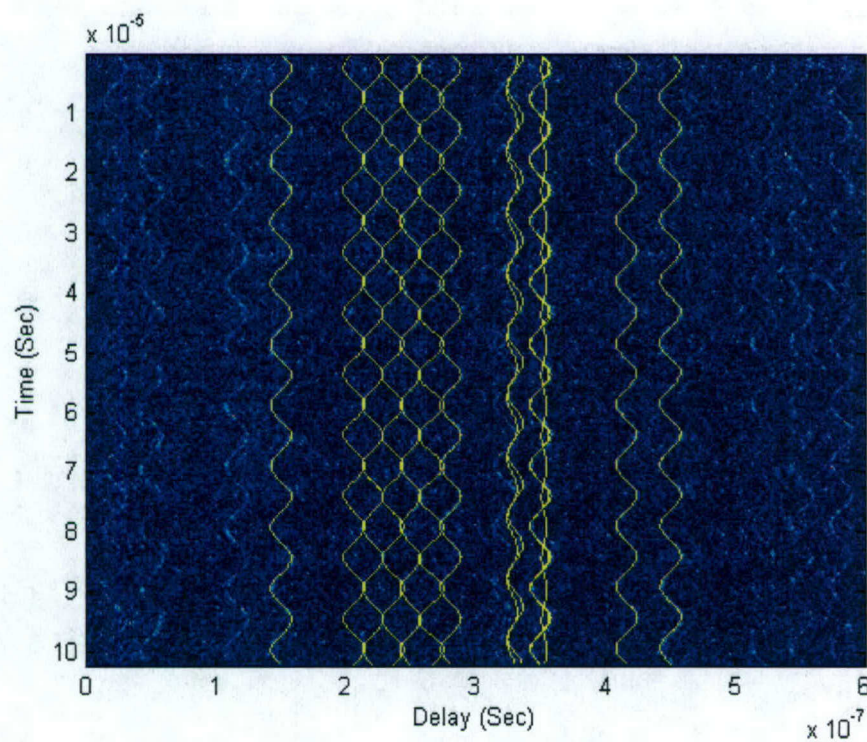


Figure 3-22: At least one template successfully fits each of the curves in the enhanced lag display.

Results and Estimate of Technical Feasibility

The Fast Adaptive Modulation task has demonstrated a novel method of nonlinear optimization that can be applied to end-to-end optimization of a radio communications system. The optimization methodology applied here is optimization in "least-logs" rather than the more familiar "least-squares" optimization. Least-logs optimization has the advantage that it is less sensitive to outliers than least-squares optimization. As such, it can be applied to curve-fitting and curve-finding applications in noisy images and displays. It follows that this methodology measures features, not pixels, with the resulting SNR gain due to integration over the feature in question.

In this application, the optimization procedure allows the estimation of a small number of parameters that characterize the combined effects of variations in clock rates that depend on components in both radios as well as the motion of the platforms and number and types of scatterers in the propagation environment. A single successful optimization process enables several different functionalities:

- 1) A classifier at the receiver that clusters the signal output sequences into related multipath arrivals then selects the modulation parameters used by the transmitter,
- 2) A pilot-signal-free (blind) Rake receiver,
- 3) Coherent combination of signals from multiple transmitters for low probability of interception applications,
- 4) Recovery of hidden timing information for low probability of exploitation applications.

Enhancement of linear features in displays based on a log transform is achieved as an intermediate product in the optimization process.

This project has developed a simulator that can estimate Bit Error Rates. These simulations suggest that these functionalities can be achieved at acceptable Bit Error Rates. The authors are pursuing transition opportunities within DOD and DARPA to test these applications versus appropriate baseline applications in realistic complex environments.