

1. Data Compression Development: A Comparison Of Two Floating-Aperture Data Compression Schemes, J. C. Wilson

1. Introduction

A number of methods have been suggested for compressing data, such as telemetry records from a spacecraft, into a form that is more economical to transmit. For example, instead of sending back sample readings of temperature every minute or every second, the mean and variance of temperature readings for that hour might be transmitted once each hour. Some recent work has been done at JPL in this area of data compression (Ref. 23).

In most cases, however, the equipment required to compress the data must be as simple as possible; and in some cases, the desired information must approximate the actual waveform of the signal. The methods discussed below have the advantage that they are quite simple. They operate directly on data that would normally be transmitted in an uncompressed mode so that an approximation to the exact waveform is sent.

2. Zeroth and First-Order Data Compression

The basic method is the following (Fig. 9). The data are given as samples, taken τ seconds apart, of a signal $f(t)$. The samples are taken often enough to define the signal, through some interpolation procedure, to the accuracy desired. At $t = 0$ (arbitrary), $f(0)$ is transmitted. If $f(\tau)$ and $f(0)$ differ by more than some preassigned constant K , called the (half) aperture, then $f(\tau)$ is transmitted at $t = \tau$. If $f(\tau)$ and $f(0)$ differ by less than K , then $f(\tau)$ is not transmitted, and the next sample, $f(2\tau)$, is compared with $f(0)$. It will then be sent or not sent (at $t = 2\tau$), depending on how much it differs from $f(0)$.

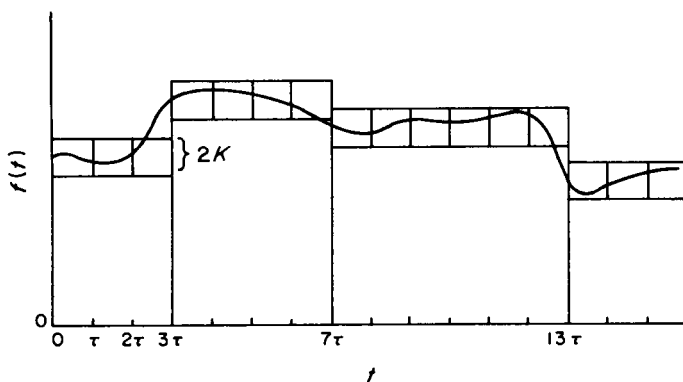


Fig. 9. Zeroth-order compression scheme

Each time a sample, say $f(j\tau)$, is transmitted, it becomes the new reference to which succeeding samples are compared. Since the aperture effectively "floats" around the last transmitted sample, this is often called the floating-aperture scheme (SPS 37-17, Vol. IV, pp. 81-84). We shall also refer to it as zeroth-order data compression.

We can readily extend the basic idea to what could be called a first-order compression scheme, where the slope of the function is taken into account. In one possible method, $f(0)$ is first sent. Then the slope of the curve between $f(0)$ and $f(\tau)$ can be found, and a line extended beyond $f(\tau)$ with that slope. If the next sample, $f(2\tau)$ differs by more than K from the line, it is transmitted; otherwise it is not, and the next sample is examined. As before, once a sample value has been transmitted, it becomes the new reference for succeeding samples, with a new slope determination (see Fig. 10).

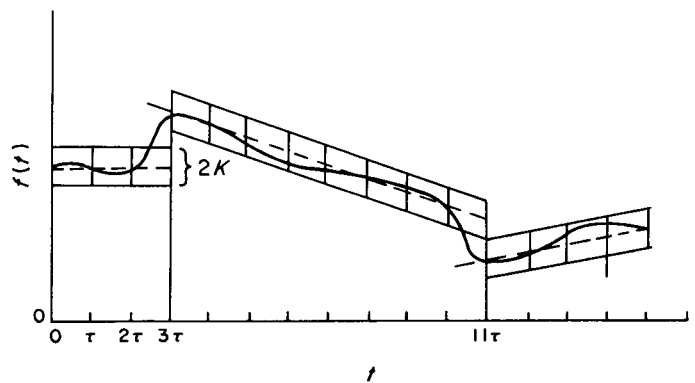


Fig. 10. First-order compression scheme

A variation on this first-order scheme has been suggested by T. O. Anderson (JPL Section 339). This involves storing a small number of slopes for comparison with the actual slope of the function. Whenever a sample $f(j\tau)$ is transmitted, the slope between it and the next sample is compared with the stored slopes, and the closest stored slope is used as the reference to which succeeding samples are compared.

It is fairly clear that higher compression ratios (the ratio of the total number of samples observed to the number of transmitted samples) would be expected for the first-order compression scheme than for the zeroth-order one. However, we might also ask how much error we will suffer at the receiver when the (compressed) signal is reconstructed.

For the zeroth-order scheme, the error due to compression is bounded by $\pm K$ about the last transmitted sample. The first-order case is not quite so simple, since we do not know the value of the slope of the curve at the last transmitted sample. If we were to send two values at each transmission, to cover both $f(j\tau)$ and the slope of

the curve at $f(j\tau)$, the (compression) error would be $\pm K$, as before. But if only one sample is sent, the error is indeterminate, although the actual error after reconstruction will not differ much from that of the zeroth-order scheme if a simple interpolation procedure is used. Fig. 11 shows the possible errors resulting from the use of these data compression schemes.

3. Experimental Setup

A good idea of the range of compression ratios available can be obtained by performing the various compression schemes on a random signal. For this purpose, colored Gaussian noise was generated on a computer by smoothing white Gaussian noise of zero mean and unit variance (also generated on the computer). The two compression schemes were programmed, and various combinations of τ (sample spacing) and K (aperture) tried. In addition, for some choices of these parameters, the computer plotted the uncompressed data and the zeroth-order and first-order approximations.

Here something should be said about subjectiveness in interpreting the data. We can think of many fidelity criteria or metrics that will give us a numerical measure of how well the compressed data approximates the original data. From each of these metrics we can make some kind of comparison between various compression schemes. However, two problems arise. First of all, one scheme may give best results under one metric while another scheme gives best results under another metric. Consequently, there is a subjectiveness in deciding which representation of fidelity is most appropriate for our particular purposes.

There is also the related problem of determining what types of signals are most likely to be received. For example, for certain types of signals, sampling uniformly at a slow rate and using a good interpolation scheme may be preferable to using, say, the first-order data-compression method. However, for signals consisting of long steady periods, with occasional transient bursts, some other method would be more economical. The waveforms used were supposed to be short samples of the latter type of process.

4. Results

The experimental results are graphically displayed in Figs. 12–15, which illustrate typical behavior with changes in aperture-width and sample spacing. As mentioned before, the *compression ratio* is the ratio of the number of samples observed to the number actually transmitted after

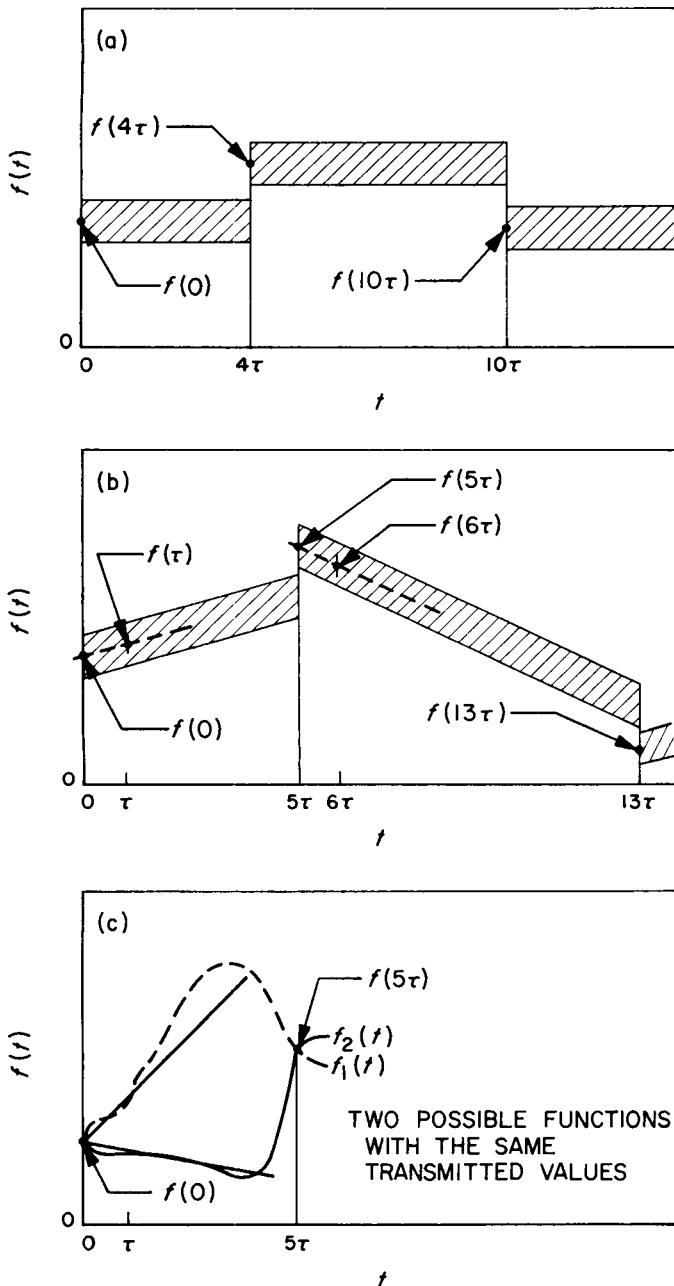


Fig. 11. Bands of possible values for $f(j\tau)$ at receiver: (a) zeroth-order compression; (b) first-order compression with value of slope given; (c) first-order compression with value of slope not given

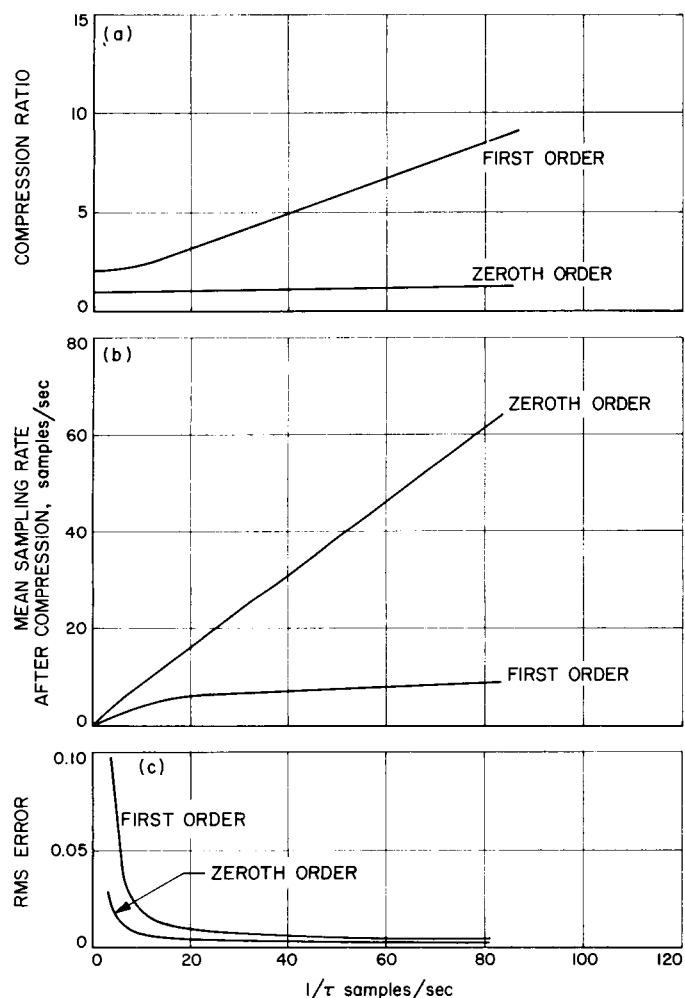


Fig. 12. Compression ratio, mean sampling rate, and rms error vs. sampling rate ($1/\text{sample spacing}$) for a half-aperture of 0.0125

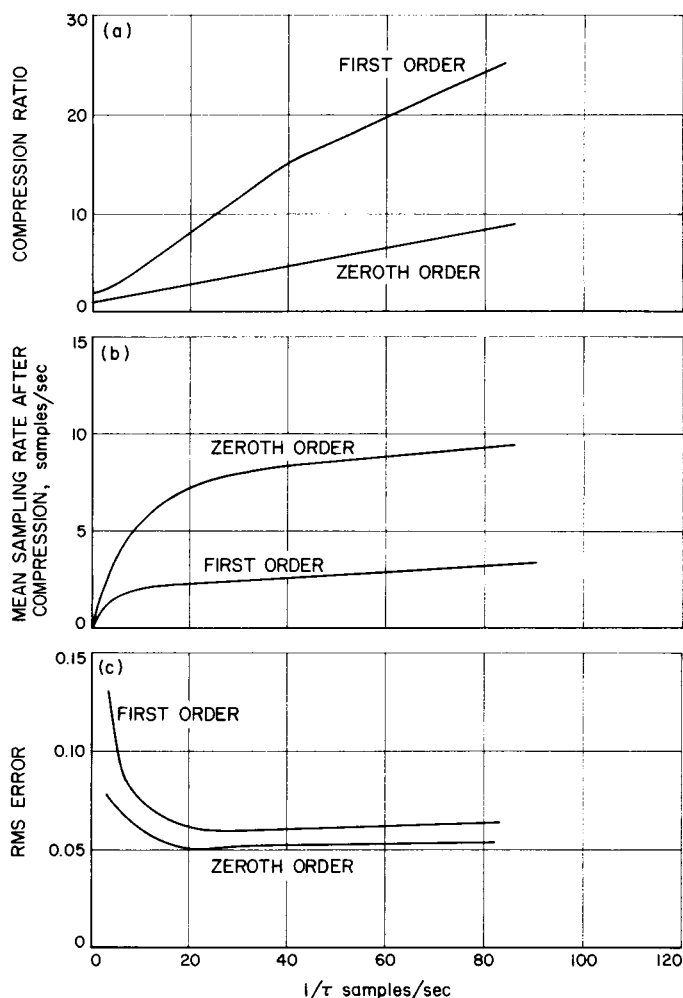


Fig. 13. Compression ratio, mean sampling rate, and rms error vs. sampling rate for a half-aperture of 0.125

compression. Considered alone, the compression ratio is somewhat meaningless, since we can achieve arbitrarily high ratios simply by taking samples closer and closer together. A more meaningful measure of compression is the *mean sampling rate*, which is the average rate at which samples are transmitted after data compression has taken place; it can be derived from the compression ratio as:

Sampling rate before compression $\div$ compression ratio.

The *efficiency* of the system relative to some theoretical optimum would be a better parameter, but the optimum is not usually known.

Finally, *rms error* refers to the rms value of the difference between the uncompressed data and the compressed

data, after a linear interpolation between compressed points. This, then, is our measure of how well the transmitted samples approximate the original data.

Of primary interest in this report is the comparison between the zeroth-order and first-order compression schemes (slope not transmitted). In all cases, substantially lower mean sampling rates were achieved with the first-order scheme than with the zeroth-order one (Figs. 12(b) and 13(b)). In most cases, this was not at the expense of a larger error (Figs. 12(c) and 13(c)). For small sampling rates in the neighborhood of less than 10 samples/sec, the difference in error became less pronounced as the aperture was increased, as shown in Fig. 14(c), which plots error versus aperture for a rate of 4 samples/sec. In general, then, the first-order scheme is probably to be preferred over the zeroth-order one from the

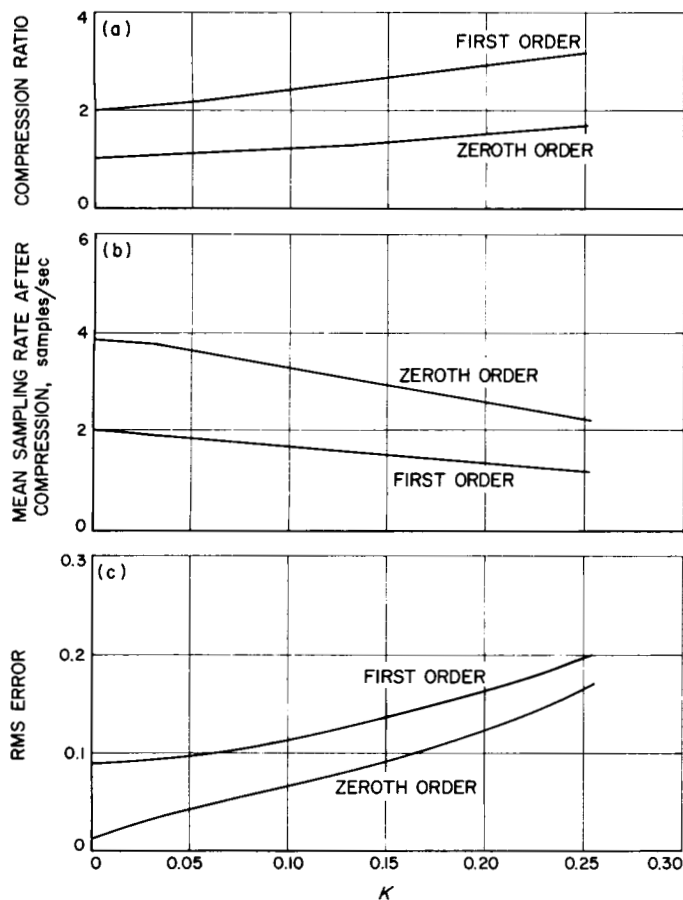


Fig. 14. Compression ratio, mean sampling rate, and rms error vs. half-aperture (K) for a sample spacing of 0.25 sec.

standpoint of economy in transmission, disregarding the various hardware requirements of the two methods.

There are some interesting results aside from a comparison of the two schemes. The first relates to the observation made earlier that a high compression ratio, considered by itself, can be misleading. Fig. 12(a) and 12(b) show that, even though the compression ratio increases with the sampling rate, the mean sampling rate is also increasing, which is not desirable.

A second observation is that although the rms error changes sharply with changes in the sampling rate for $1/\tau < 10$ samples/sec., it levels off above this point. Taking samples closer together does not lessen our error substantially, but it does bring about an increase in the mean sampling rate after compression. There is definitely room in this area for design to achieve a trade-off between small errors and low average sampling rates.

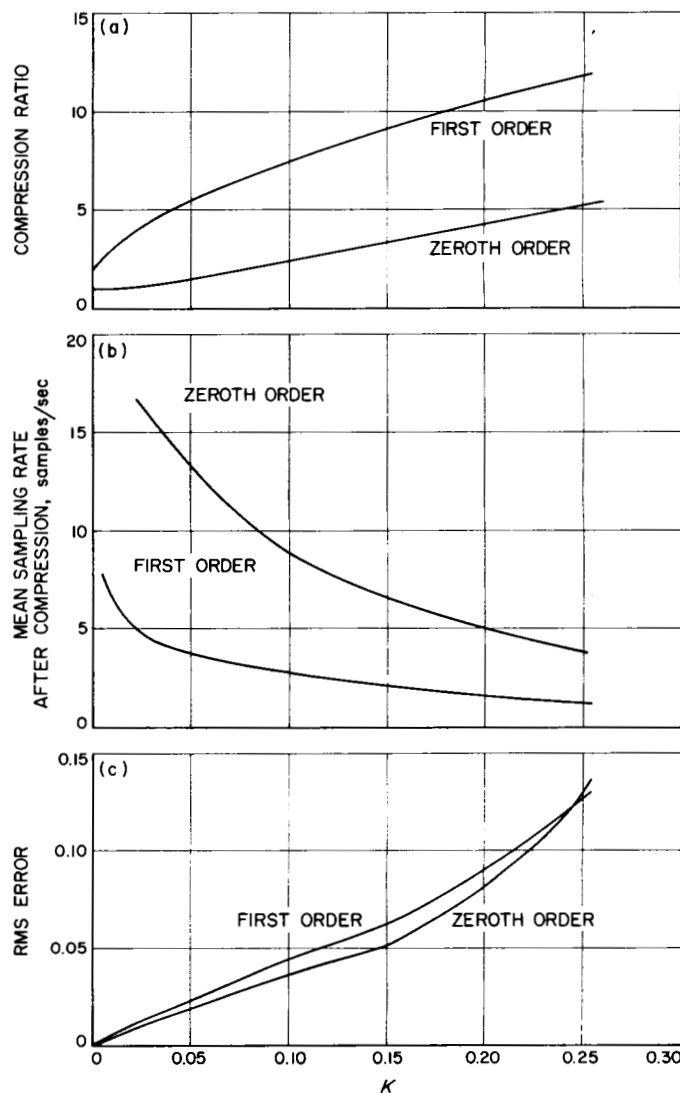


Fig. 15. Compression ratio, mean sampling rate, and rms error vs. half-aperture for a sample spacing of 0.05 sec.

A third observation is that the rms error for the first-order scheme was achieved with no knowledge of the slope of the signal at the receiver. Transmitting the slope will insure that we can reconstruct the signal with a maximum absolute error of K at each of the original sampling points.

Finally, we can look at some actual waveforms (Figs. 16–18) to compare the compression schemes. In all the plots, a linear interpolation was made between data points. In Fig. 16, it will be seen that even a very small aperture does not necessarily mean good fidelity. Fig. 17, however, illustrates the trend, mentioned above, of a smaller error with a smaller τ (higher sampling rate).

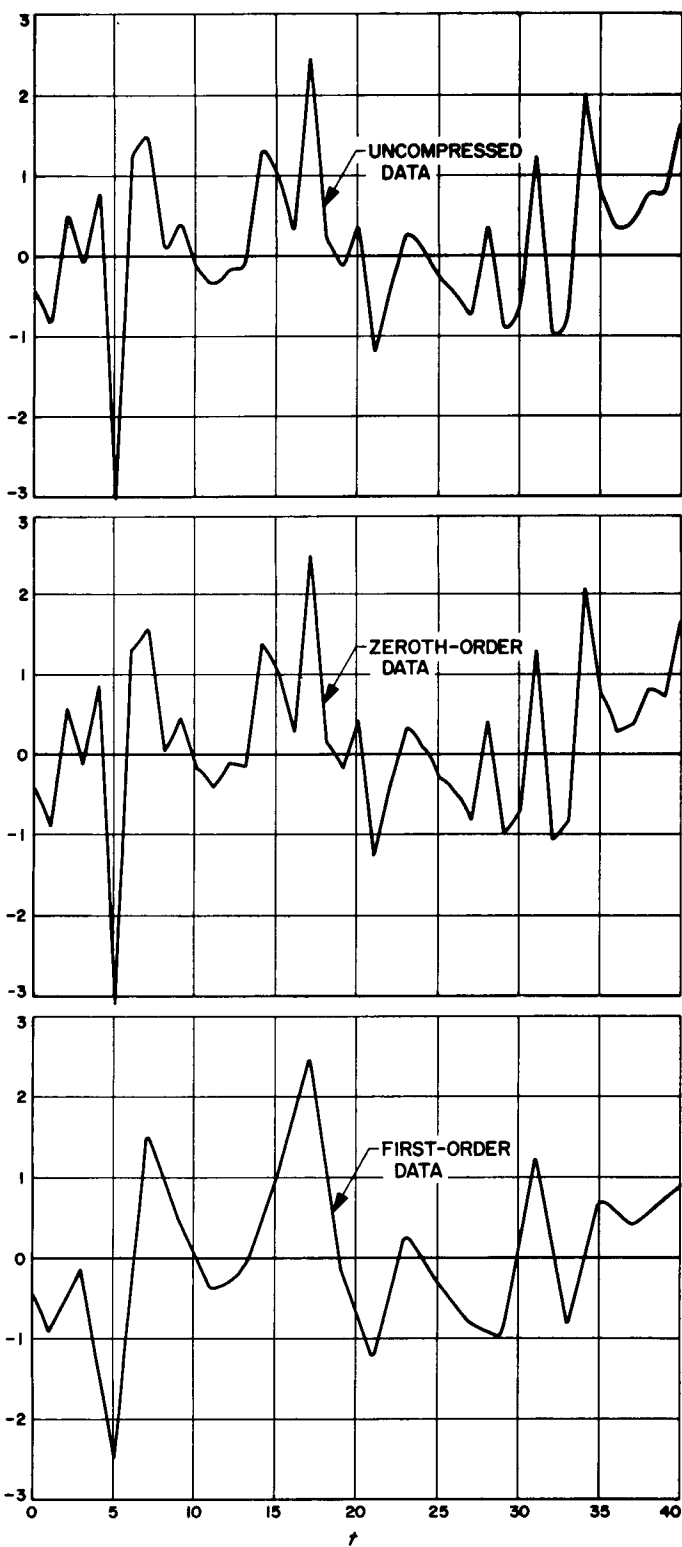


Fig. 16. Plots of uncompressed data, and the zeroth and first-order approximations to them:
 $\tau = 1.0$; $K = 0.00625$

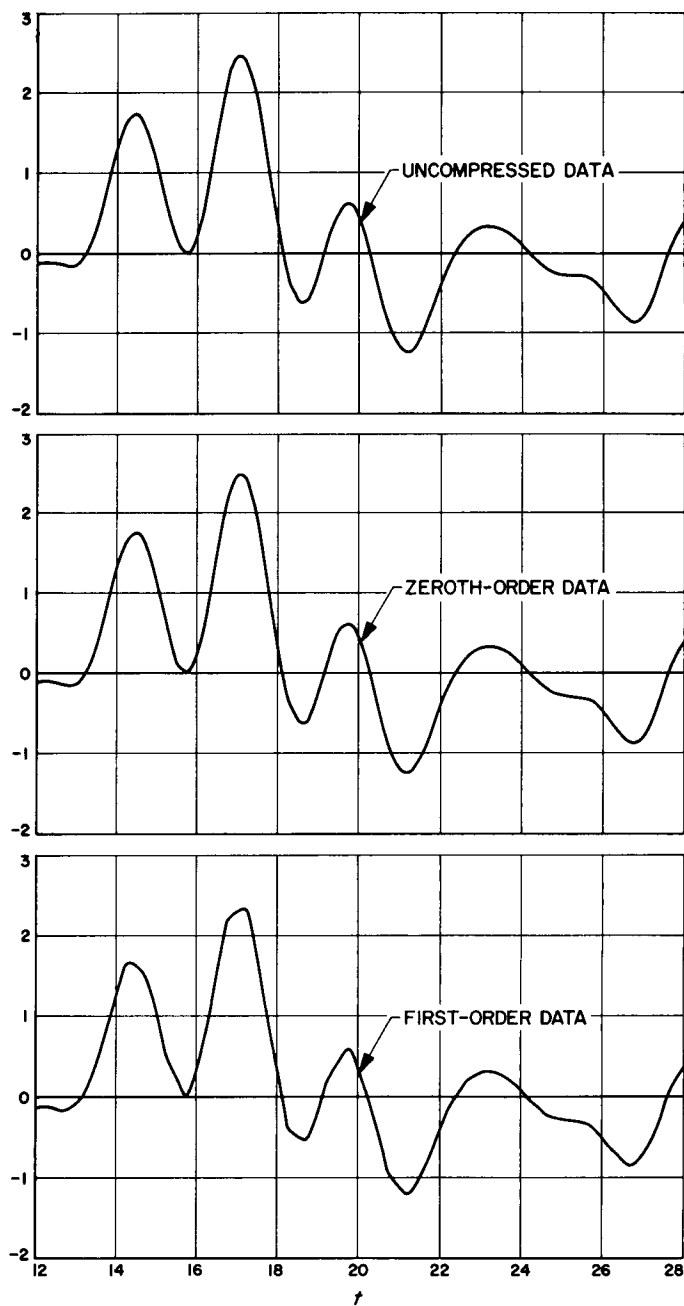


Fig. 17. Plots of uncompressed data, and the zeroth and first-order approximations to them:
 $\tau = 0.25$; $K = 0.00625$

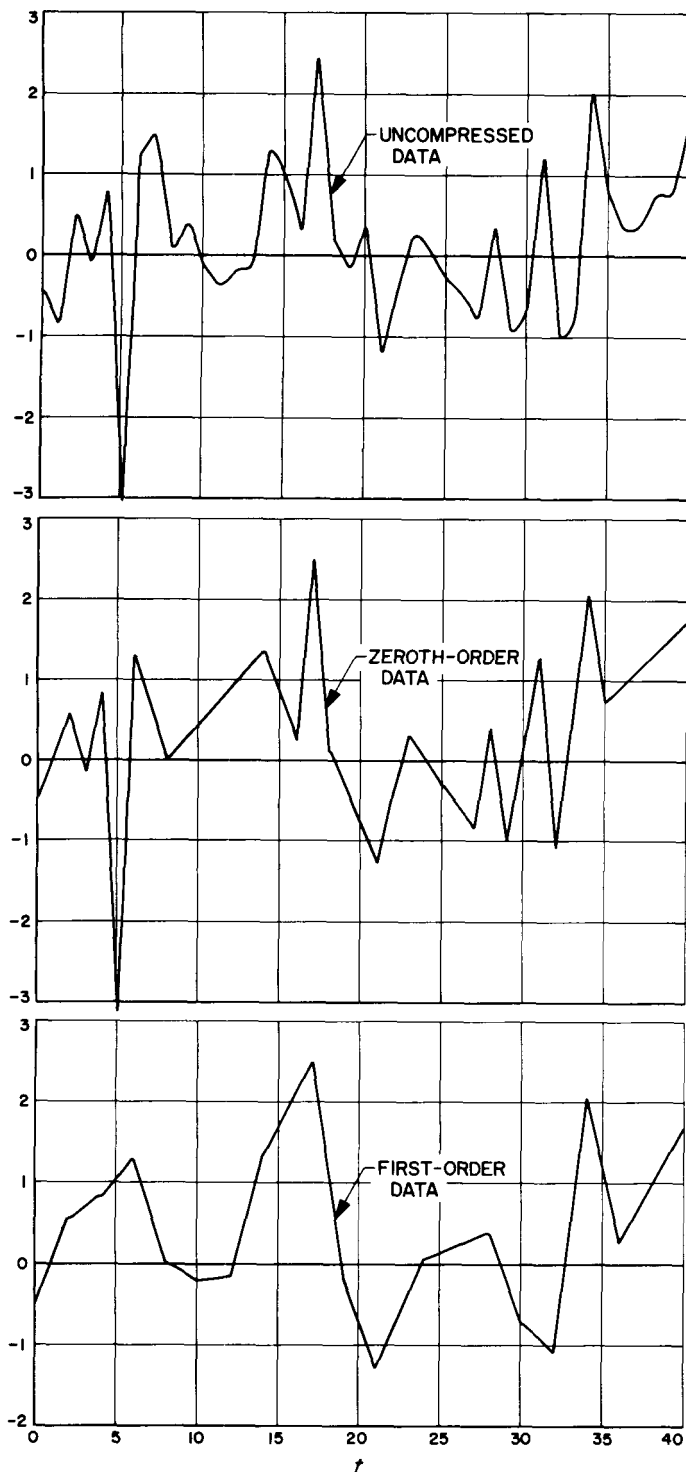


Fig. 18. Plots of uncompressed data and the zeroth and first-order approximations to them:
 $\tau = 1.0; K = 0.5$

Figures 16 and 18, when compared, illustrate the loss of fidelity when the aperture is increased. Note particularly the zeroth-order case.

We have seen that the first-order data compression scheme is generally superior to the zeroth-order one. A reduction of a factor of 3 to 6 or more in the number of samples sent occurs when the first-order system is chosen over the zeroth-order system, for a wide choice of parameters. Moreover, the first-order system requires fewer timing bits, since samples occur less frequently. It appears safe to say that, for a large class of signal sources, that resemble colored Gaussian sources, a factor of at least 3 can be gained in choosing a first-order system over a zeroth-order one.

One last point; for the purposes of comparing data-compression schemes, instead of sending out the compressed version of the data samples as they are observed, the usual procedure is to store them in a buffer to be sent out at a uniform rate. Consequently, timing information will have to be sent along, decreasing the net transmission rate. Also consider the quantization of the data for coding purposes. The results indicate that there is much room for trade-off design once a scheme has been chosen. For example, a large sample spacing (low rate) may result in a large error, while a very small spacing may yield too much compressed data to transmit economically.

N67-15763

J. Digital Devices Development: Integrated Circuits — Delay Line Interface Circuitry,

R. A. Winklestein and E. Lutz

1. Introduction

This article describes a unique and flexible system of interfacing ultrasonic delay lines within digital systems using integrated-circuit logic. To avoid misunderstanding, the term "NRZ" in this document refers only to input-output waveform relations which are similar to those achieved by conventional Non-Return-to-Zero detection schemes. However, all the advantages of the NRZ mode of operation are maintained while certain distinct disadvantages are completely eliminated.

The basic technique is to store digital information of an aperiodic nature in the delay medium and to retrieve an exact replica of that information delayed by a certain amount. The conventional method of detection is shown by the timing diagram in Fig. 19 where the amplified waveform from the delay line is "sliced" by two threshold levels: one level to set and the other to reset a flip-flop. A change of state at the output of the flip-flop takes place only if a "1" follows a "0," or a "0" follows a "1." This method produces an output waveform which is an exact