

ω_2 , the limit points of the pass band. The system being minimum phase, the corresponding delay is obtained by the Hilbert transform of the slope of the attenuation and becomes¹

$$\begin{aligned}\beta'(\omega) &= \frac{1}{\pi} \int_{-\infty}^{\infty} \left\{ \delta(\lambda - \omega_1) - \delta(\lambda + \omega_1) \right. \\ &\quad \left. - \delta(\lambda - \omega_2) + \delta(\lambda + \omega_2) \right\} \frac{d\lambda}{\lambda - \omega} \\ &= -\frac{2}{\pi} \left\{ \frac{\omega_2}{\omega^2 - \omega_2^2} - \frac{\omega_1}{\omega^2 - \omega_1^2} \right\}. \quad (1)\end{aligned}$$

The delay tends to infinity at the edges, whereas the point of minimum delay drifts towards the low-frequency edge ω_1 if $\omega_2 \gg \omega_1$.

A similar calculation can be performed for a minimum phase system which is assumed to have an ideal delay behavior in the pass band. It shows that the ideal delay (i.e., constant except for delta singularities at ω_1 and ω_2) appears with a distorted amplitude response. In actual systems, where delay distortion is more undesirable than amplitude distortion, one sometimes applies a moderate and acceptable amount of this minimum phase amplitude distortion to improve the delay.

The above discussion leads to the simple result:

Lemma 1: A long distance band pass communication system of the minimum phase type has an inherent delay distortion if it is constant amplitude in the pass band.

Lemma 2: A long distance system with a constant delay in the pass band has always an amplitude distortion if the system is minimum phase.

Note that these phase and amplitude distortions have singularities at the edges of the pass band; in practice, amplitude singularities at the edges as implied by lemma 2 are not permissible because actual systems will only be linear for finite amplitudes.

It follows from lemma 1 and 2 that

Theorem 1: A minimum phase long distance system with finite pass band always imparts a linear distortion to a band-limited signal within the pass band.

Hence, the transmission of a band-limited signal without linear distortion, if possible at all, requires at least a nonminimum phase system. For practical reasons we will consider the delay correction provided by the all-pass elements described by the partial fraction expansion (6a) mentioned in conjunction with the constant amplitude band-pass system.¹

Fig. 1 gives a graphical illustration of how the delay distortion can be decreased by adding extra delay.

The procedure demonstrated in Fig. 1 can, of course, be refined by an appropriate choice of the number of terms in the expansion (6a) and the positions of the corresponding zeros, i.e., x_n and y_n .² The delay correction can be pushed towards the edges of the

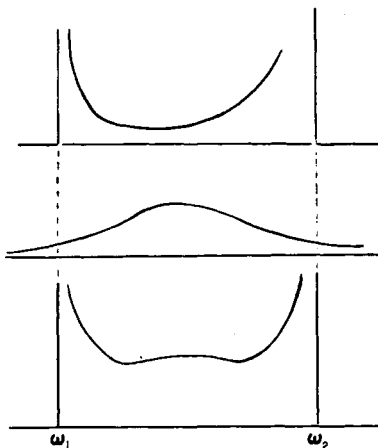


Fig. 1—Delay correction of a minimum phase long distance system by adding all-pass delay.

pass band by adding in more and more delay, however, one can never utilize the whole band. Thus one can trade effective distortion-free bandwidth for extra delay. Formulated in the form of a theorem:

Theorem II: The delay distortion of a minimum phase, long distance system with constant amplitude in the pass band ($\omega_2 - \omega_1$) can be arbitrarily reduced in an interval $\Delta\omega < (\omega_2 - \omega_1)$ by adding an appropriate all-pass delay. The interval $\Delta\omega$ can approach the actual pass band if, and only if, the all-pass delay tends to infinity.

The above treatment of delay distortion differs from customary discussions primarily in the emphasis on two points. They are:

1) The considerations are based on the delay and slope relations rather than on the somewhat more limited phase-attenuation relations.¹

2) For a real long distance system it is not possible to use the transmission characteristics outside the pass band for the elimination of linear distortion.

Corollary: It is essential in a long distance system to transmit "edge-bands" for equalization purposes. These edge-bands can be made very small with respect to the actual pass band.

E. J. Post
AF Cambridge Research Labs.
Bedford, Mass.

The Motivation and Technique of Writing Scientific Contributions*

INTRODUCTION

The problem of scientific publication has been treated from a number of points of view, but many aspects of this important field appear to have been neglected. It is, for example, somewhat surprising that the great

advances made by motivational research in the last decade have made so little impact on this subject. There are undoubtedly numerous books and pamphlets on the subject of how to write a paper but these are invariably concerned either with broad generalities (write intelligibly, say what you mean, do not digress, etc.) or with trivial matters of preparation (reasonable margin must be left at the edge of the manuscript, figure captions must be typewritten, illustrations must not exceed the 4"×6" size, etc.). While not denying the importance of these recommendations I feel that they represent a rather narrow view and leave the basic questions of motivation completely unanswered. The purpose of this communication is not to introduce new concepts but rather to present some of my personal experiences in writing technical papers and to pass on the valuable comments which I have received from a number of friends.

SOME THOUGHTS ON MOTIVATION

The motives of paper writing are involved, ranging from a simple love of writing to the most complicated cases of status seeking. I do not wish to go into detail here and shall confine the study to the following four categories: 1) dissemination of knowledge, 2) establishment of priority, 3) professional reputation, and 4) status.

Motive 1) applies in the main to young men who may be preparing their first scientific contribution. The numbers involved here are rather small and few of them will ever write a single paper. Therefore 1) cannot be treated on the same level as the other—more compelling—motives. Nevertheless we should not forget that 1) provides the pretext for the existence of any paper and although its role in motivation is negligible its effect on the presentation of a paper is still considerable.

Priority

This motive applies again only to a narrow section of the paper-writing community, but its importance far exceeds that of any other motive. The desire to attach his own name to a discovery has long been a distinguishing characteristic of a research worker. Since the proof of discovery lies in a publication, there is a tendency to publish quickly. The author has to keep in mind, however, the possibility of further exploitation. If he publishes his findings, someone else might follow up his ideas and might deprive him of the further fruits. The ideal solution is to ensure priority by announcing the discovery and then delay publication until a thorough assessment of its potentialities has been made. The first scientist known to have applied this method was Galileo Galilei who sent to Kepler the description of his astronomical discoveries in anagram form, providing the solution only a year later. Since modern scientific journals are not in the habit of publishing anagrams the modern discoverer (or inventor) has to turn to a different technique. To begin with I would recommend an impressive title because the more resounding the title the less information need be conveyed in the paper. For example, a title "The Stressed Negative Induc-

¹ Note that the numerical factors of the impulse functions have been taken "one" for simplicity. E. J. Post, "Note on phase-amplitude relations," *Proc. IEEE (Correspondence)*, vol. 51, p. 627; April, 1963.
² See for instance C. H. Dagnall and P. W. Rounds, *Bell Sys. Tech. J.*, vol. 28, p. 181; 1949.

* Received December 17, 1962, revised manuscript received January 8, 1963.

tance Amplifier" would immediately convince every reader that an important new principle has been discovered. The author will then be forgiven if he is rather vague on the essential points and gives only a broad outline of the discovery.

A further argument in favor of vagueness is our moral obligation to our descendants. In a few generations a nation may emerge with a profound desire to uphold the honor of its ancestors. It will wish to prove that every discovery and every invention, no matter how small, was initiated by the very citizens of that nation. If we are not sufficiently vague now, we make its future work more difficult.

Professional Reputation

Professional reputation can be obtained in a number of different ways. It is sufficient, for example, to make a major invention or, even better, to receive the Nobel Prize and one's professional competence is no longer challenged. For the majority of research workers, however, the only way open is to write a large number of papers each contributing a little to the advance of science. It is advisable to confine the first few papers to a narrow subject (for example to waveguide discontinuities) to obtain recognition, but later the author has also to give evidence of his versatility by writing numerous papers covering a wider subject (for example microwaves). Having reached the three dozen mark the author's reputation can no longer gain by writing additional papers. This is the time to cut off publication suddenly (a few survey papers are still permissible) and find a comfortable position in management.

Status

Scientific status can be obtained through professional reputation and professional reputation by publishing papers. Thus if this simple relationship were always to apply there would be no need to mention the motive of status seeking separately. There is, however, a growing body of opinion which maintains that the intermediate step of professional reputation is not essential. It is suggested that a person's status can be elevated by publishing papers whose scientific value is zero or even negative, and it is emphasized that only the total number of papers is significant. While I lack the evidence to disprove these claims statistically, I think that the long-term gains are very questionable. I would therefore tend to recommend this method only as an emergency measure when one's creative genius has temporarily dried out.

SOME THOUGHTS ON PRESENTATION

So far I have treated the underlying motives and have only lightly touched upon the problem of presentation. I would now like to investigate the problem of young authors (without a powerful co-author) whose main concern is to pass unscathed through the line of reviewers.

How to Have a Paper Accepted

Reviewers are selected from the leading scientists in order to filter the flood of manu-

scripts reaching the editor. Unfortunately, leading scientists are generally busy, have many obligations and administrative burdens as well. They are unable to spend the major part of an afternoon on reading a single paper; but nonetheless, they make comments.

A young author should be aware of this situation and—instead of wasting his time on complaints—should write his paper so as to satisfy the reviewer, whose sharp eyes detect the slightest irregularity. If the paper is too long the author will be accused of verbosity; if the paper is too short he will be recommended to collect further additional material. If he reports on purely experimental work the "foundations" will be criticized; if he propounds a simple theory he will be called "superficial." If he lists a too long bibliography he will be regarded as "unoriginal," if he refers to no one at all he will be branded "conceited." Thus I would suggest a compromise. The length of the paper should be between 8 and 12 type-written pages (double spacing, of course) and about one third of it should be covered with mathematical formulas richly decorated with integral signs and special functions. The number of references should vary between 6 and 10, half of them referring to famous (the reviewer has heard of them) and the other half to unknown (the reviewer has not heard of them) works.

Following the above recommendations the author has a fair chance that his paper will not be refused *a priori*. On a superficial look the reviewer is favorably inclined. Now everything depends on his reactions in the next half hour. If within this time he can quickly comment on three minor errors, the paper will be passed. If he finds nothing on which to comment his resistance will only harden. He will take the next assumption meeting his eye, will call it unfounded (which assumption is invulnerable), and will recommend the return of the paper.

Thus clearly the task of the author is to supply the material for the three minor comments. To facilitate this choice a few examples are given below:

- 1) Choose an inappropriate title (every reviewer is fond of suggesting titles)
- 2) Forget to define one of the symbols in equation (1)
- 3) Deviate from the usual notations (only in the case of one parameter)
- 4) Misspell a word (only one) which is often misspelt (preferable for British journals)
- 5) Write $\exp x$ and e^x alternately.

The rules for an advanced contributor (at least ten papers published) are much more relaxed. He may write colorful introductions, may hide a few cracks in the main body of the text, may admit that he does not quite understand the results, etc.

I hope the above note will help in reaching some understanding of the nature of paper writing and will at the same time provide some guidance for young authors.

L. SOLYMAR

Standard Telecommun. Labs. Ltd.
London Road
Harlow, Essex, England

Theoretical Considerations anent Pattern Recognition by Means of Random Masks*

INTRODUCTION

An optical pattern recognizer has been described by Palmieri, *et al* [1], which functions as follows for dichotomic separations: The luminous pattern to be recognized is projected on a multiplicity of masks having clear and opaque portions, and for each projection, a note is made whether the light transmitted by the mask exceeds a given threshold specific to each mask, thus causing a decision circuit to "fire." A weighted tally of the firings serves to determine whether the exposed pattern belongs to class *A* or *B*.

The weights of the tally are determined by first exposing several times each mask, M_1, M_2 , etc., to several patterns A_1, A_2 , etc., B_1, B_2 , etc., belonging to the two classes *A* and *B*, the relative orientation of each mask and pattern being varied either at random or by regular angular increment, and by noting, for the *i*th mask, the fractional number of firings with the two classes, designated by a_i and b_i . The weights of the tally are then taken as

$$\lg \frac{a_i(1 - b_i)}{b_i(1 - a_i)},$$

the weighted tally is added to

$$\sum \lg \frac{1 - a_i}{1 - b_i}.$$

and the unknown pattern is said to belong to class *A* or *B* depending upon whether the figure obtained is positive or negative.

The masks are random in the sense that the opaque portion of each mask contains a few straight or curved lines in no particular orientation to each other, or a bit of writing, etc.

Remarkable recognizing performance has been reported for this recognizer (designated by the acronym P.A.P.A.) in the publication cited.¹ The intriguing nature of the principle involved has been further underlined by the use of anthropomorphic expressions such as "intelligence" of the device, its "ability to infer," etc., and by a reference to the human brain. However, there has been little attempt to construct a theory of the real *modus operandi* involved. The purpose of these notes is to indicate what this *modus operandi* probably is, in simple experiments, and to draw tentative conclusions therefrom.

DISCUSSION

If the various light outputs of a mask exposed to several patterns in all possible orientations had a Gaussian distribution with a given mean for one class of patterns, and a different mean with the same variance for another class, an efficient procedure would be to expose that mask to the unknown pattern in all possible orientations, determine whether the average light output then obtained is closer to the mean of one or

* Received November 30, 1962.