

Communications of Discrete Information Models with Best 1:1 Mean Codeword Lengths and Suitable Codes

Retneer Sharma¹, Om Parkash², Rakesh Kumar³, Vikramjeet Singh⁴

¹Department of Mathematics, Guru Nanak Dev University Amritsar, India,

sharma.ret@gmail.com

²Department of Mathematics, Guru Nanak Dev University Amritsar, India,

omparkash777@yahoo.co.in

³Department of Mathematics, Guru Nanak Dev University Amritsar, India,

rakeshmaths@yahoo.in

⁴Department of Mathematics, I.K. Gujral Punjab Technical University Amritsar Campus, India,

vikram31782@gmail.com

Article Info

Page Number: 6572 - 6596

Publication Issue:

Vol 71 No. 4 (2022)

Abstract

The well recognized reality concerning the discrete information models indicates the practical magnitude of their relevance towards countless information processing systems. On the other hand possible codeword lengths and their possible lower bounds in addition happen to be a subject of matter for the furtherance of research. At the same time coding theory formulates transactions with a variety of codes accessible with the literature of information theory including uniquely decipherable codes, instantaneous codes, possible codes, suitable codes and the best 1:1 codes. The primary objective of our learning is to extend the literature on these codes along with the illustration of involvement sandwiched discrete entropic, divergence and inaccuracy models. The foremost intention of our communiqué is to make available the correspondence between information theoretic entropic models and the best 1:1 code for revealing the fruitful results. Additionally, our principal objective is to deliberate contributions of discrete divergence and inaccuracy models for the development of suitable codes.

Keywords: Kraft's inequality, Codeword length, Noiseless coding theorem, Suitable codes, Binary code, Best one-one codes.

Article History

Article Received: 25 March 2022

Revised: 30 April 2022

Accepted: 15 June 2022

Publication: 19 August 2022

INTRODUCTION

It is well celebrated actuality that to facilitate sensible applications of discrete information models in the control of probabilistic coding theory, a broad-spectrum approach has been endowed with in a structure pedestal on entropic model pioneered by Shannon [25]. This quantitative entropic model persuades several advantageous self-evident prerequisites and furthermore it can be capable of allocation with an outfitted consequence in imperative convenient optimization problems of numerous disciplines. This is the distinguished authenticity that information models are significant for convenient relevance of information dispensing a broad-spectrum advancement in statistical structure. This entropy persuades some advantageous self-evident requirements and additionally can be dispensed outfitted connotation in imperative realistic problems.

The well-acknowledged and authoritative authenticity about the Coding theory provides the exploration of mixture of codes through discrete probabilistic entropic models and makes deliberations towards applications in extraordinary disciplines. Shannon [25] structured the hypothetical environment upon introducing the crucial conception of entropy $H(P)$ attached with the discrete probability spaces. The fundamentally well-acknowledged perception of probabilistic entropy premeditated by Shannon [25] enriched the literature of coding theory with the facilitation of numerous entropic models. This entrenched advancement arranged the stone of discrete information entropic model with astonishingly agreeable properties and was well acknowledged by means of subsequent quantitative model:

$$H(P) = - \sum_{i=1}^n p_i \log p_i \quad (1.1)$$

To boost the literature of discrete entropic models, Parkash and Kakkar [15, 16] organized protracted efforts for the investigation of abundant entropic models for the discrete probability spaces from application point of observation and consequently enhanced the literature of discrete entropy models by means of twisting the subsequent quantitative entropic models:

$${}_{{}_\beta} S(P) = \frac{\sum_{i=1}^n p_i \beta^{\log_D p_i} - 1}{1 - \beta}, \quad \beta > 1 \quad (1.2)$$

$$S_{\beta}(P) = \frac{\beta^{\sum_{i=1}^n p_i \ln p_i} - 1}{1 - \beta}, \quad \beta > 1 \quad (1.3)$$

and

$$S_a^b(p) = \frac{a^{\frac{1}{b-1} \ln \left(\sum_{i=1}^n p_i^b \right)} - 1}{1-a}, \quad a > 1, \quad 0 < b < 1 \quad (1.4)$$

There subsist a colossal assortment of discrete entropic models but still expectedness arises to originate amplification in the manuscript of these models. Additionally, there become noticeable exceedingly strong involvement connecting entropy and Statistics from application position of thoughtfulness. To carry out this purpose, Parkash, Sharma and Singh [20] produced a new-fangled pioneering discrete entropic model and by this development, the authors improved the application area of well identified principle accountable for the clarification of abundant optimization problems known as “maximum entropy principle”.

To enhance the literature of discrete weighted entropic models, Parkash, Kumar, Mukesh and Kakkar [19] well thought-out prolonged efforts for the exploration of plentiful weighted parametric entropic models for the discrete probability spaces from application point of view in the field of coding theory. Through their cooperative efforts, the authors delivered numerous observations and consequently enhanced the literature of such models by means of twisting the two subsequent quantitative outward show:

$$H_\alpha^1(P; W) = \frac{1}{2^{1-\alpha} - 1} \left[\sum_{i=1}^n w_i p_i^\alpha - \left\{ \sum_{i=1}^n w_i p_i^{\frac{1}{\alpha}} \right\}^\alpha \right], \quad \alpha > 1 \quad (1.5)$$

and

$$H_{\alpha, \beta}(P; W) = \frac{1}{\beta - \alpha} \log \left[\frac{\sum_{i=1}^n w_i p_i^\alpha}{\sum_{i=1}^n w_i p_i^\beta} \right], \quad \alpha < 1, \beta > 1 \text{ or } \alpha > 1, \beta < 1 \quad (1.6)$$

This is to highlight that in addition to entropic models, divergence models also cooperate with momentous responsibility for the development of countless optimization problems. An additional viewpoint pertaining to the inevitability of such models is that the miscellaneous attempts have been made to broaden the perception of distance in assorted fields other than mathematical sciences where distance models can be effectively made functional. But, distance in such cases may not necessarily be geometrical and hence there is prerequisite for its modification. Keeping in mind the present conception, abundant distance models have been scrutinized by an assortment of researchers. Further, it is added that magnitude of events with which these outcomes happen cannot

be overlooked and consequently weights can be put together to these outcomes for the development of weighted divergence models. Motivated by the technique deliberated by Kapur [5], Parkash, Kumar and Kakkar [18] created certain innovative divergence models desired to be investigated for the discrete weighted distributions.

Additionally, Huang and Zhang [3] delivered a surprising interpretation with reference to Shannon's [25] mutual information and emphasized that it has been extensively second-handed its effective computation and it is over and over again complicated for innumerable realistic problems. As a consequence, the authors pointed out that asymptotic modulus operandi pedestal on Fisher information frequently makes available wonderful estimates to such information and deliberated the estimates pedestal on certain divergence models in the probability spaces. Furthermore, the authors carried out numerical duplication and established that their estimated modulus operandi were extraordinarily wonderful with burgeoning convenience to numerous realistic and hypothetical problems. This is supplementary additional that the discrete entropy models discover marvelous applications in abundant many disciplines. Certain pioneer towards application areas of entropy models include Renyi [21], Kapur [5], Parkash and Kakkar [15, 16] etc.

It is well recognized that the discipline of information theory provides transactions with earth-shattering and a crucial areas concerned with the examination of a collection of codes for their applicability in divergent circumstances of the theory of different codes. More uncomplicatedly, the two mandatory insights, that is, entropy and coding are forcefully associated to each other and an outstandingly significant relation between the two was first accomplished by Shannon [25]. Additionally, it is meaningfully experimental that Kraft's inequality participates with an imperative accountability in demonstrating coding theorems and is outstandingly established through the stipulation of exclusive decipherability. It cannot be personalized in a subjective approach, aggravated by the aspiration to demonstrate a new consequence. If we make amendment, we shall acquire codes with a dissimilar configuration other than UD codes dissatisfying our motive of study. This inequality is capable of contemplation in provisions of controlled budget to be exhausted on codewords with shorter codewords being supplementary exclusive. Along with D as alphabet size, let l_i be the codeword length pleasant through well approved Kraft's [11] inequality and specified by the subsequent mathematical manifestation:

$$D^{-l_1} + D^{-l_2} + \dots + D^{-l_n} \leq 1 \quad (1.7)$$

After the commencement of mean length established by Shannon [25] himself, it was Campbell [1] who initially provided the inspiration of parametric mean length for uniquely decipherable (UD) codes with the establishment of the subsequent expression:

$$L_\alpha = \frac{\alpha}{1-\alpha} \log_D \left[\sum_{i=1}^n p_i D^{(1-\alpha)n_i/\alpha} \right] \quad (1.8)$$

and additionally demonstrated that the lower bound for this standard length lies between $H^\alpha(P)$ and $H^\alpha(P) + 1$ where

$$H^\alpha(P) = (1-\alpha)^{-1} \log_D \left[\sum_{i=1}^n p_i^\alpha \right]; \alpha > 0, \alpha \neq 1 \quad (1.9)$$

represents parametric discrete entropic model in the probability spaces investigated and derived through the efforts payable to Renyi [21]. After the well twisted mean length established by Shannon [25], Campbell [1] was the former to reflect upon an exceptional mean specified by the well ascribed equality pointed out in equation (1.9). With the support of Campbell's [1] exponentiated mean, Parkash and Kakkar [14] made furtherance of research attached with other exponentiated means for the progression of application areas in abundant disciplines. To enhance the literature of such means, the authors twisted the subsequent two mean codeword lengths:

$$L(a, b) = \frac{a}{1-a} \log_D \left(\frac{\sum_{i=1}^n p_i^{\frac{b}{a}} D^{\frac{l_i(1-a)}{a}}}{\sum_{i=1}^n p_i^{\frac{b}{a}}} \right), \quad a > 0, b > 0, a \neq 1 \quad (1.10)$$

and

$$L(b) = \frac{\sum_{i=1}^n p_i^b l_i}{\sum_{i=1}^n p_i^b}, \quad b > 0 \quad (1.11)$$

From the point of observation that divergence models in addition cannot be overlooked for the expansion of mean codeword lengths, Parkash and Kumar [17] made escalation of new weighted divergence model to provide new meaningful lengths and extorted the surviving ones. This ground-breaking learning provided a profound participation flanked by weighted divergence model and the mean lengths given by the subsequent appearance:

$$L^\alpha(W) = \frac{1}{(\alpha-1)} \log_D \left[\frac{\sum_{i=1}^n w_i p_i^\alpha D^{-l_i(1-\alpha)}}{\sum_{i=1}^n w_i p_i^\alpha} \right], \alpha > 1 \quad (1.12)$$

In the manuscript of coding theory, the source codes are UD codes in view of the fact that these can traditionally be embedded as a sequence of source symbols. The distinctive decodability of UD codes endows the assurance that a faultless resurgence of the transmitted symbols from the acknowledged concatenated codewords is achievable. Although the foremost spotlight in the source coding text is on the category of UD codes, there has been certain research concentration in a less restraining but larger category of codes, the group of one-to-one codes. As such, unlike UD codes, one-to-one codes need not acquire distinctive decodability.

Leung-Yan-Cheong and Cover [13] well thought-out on the literature of the one-to-one codes and demonstrated that the least amount of best one-to-one code symbolized by $L_{1:1}$ convince the subsequent inequality fascinated in the company of logarithmic function:

$$L_{1:1} \geq H(P) - \log \sum_{i=1}^n \left(\frac{2}{i+2} \right) \quad (1.13)$$

where $L_{1:1} = \sum_{i=1}^n p_i \left\lceil \log \left(\frac{i}{2} + 1 \right) \right\rceil$, $\lceil x \rceil$ designates the least integral value larger than or equal to x .

Additionally, in the companionship of logarithmic function, these authors established a ground-breaking and first and foremost fundamental inequality with the facilitation of subsequent mathematical appearance:

$$L_{1:1} \geq H(P) - \log \log n - 3 \quad (1.14)$$

The investigations delivered by the authors made furtherance of communication flanked by the two selections of codes. Again in the companionship of logarithmic function, the authors rightfully furnished the subsequent transformations from best one-one codes to UD codes:

$$T_1 : l_i = b_i + b \lceil \log b_i \rceil + \log \left(\frac{2^b - 1}{2^b - 2} \right), b > 1 \quad (1.15)$$

$$T_2 : l_i = b_i + 2 \lceil \log(b_i + 1) \rceil \quad (1.16)$$

$$T_3 : l_i = b_i + \lceil \log b_i + \log(\log b_i) + \log(\log(\log b_i)) + \dots \rceil + 3 \quad (1.17)$$

where $b_i = \left\lceil \log \left(\frac{(D-1)i}{D} + 1 \right) \right\rceil$, $i = 1, 2, \dots, n$. is the best one-one codeword length, l_i , is the UD codeword length. The authors after making deliberations on these unique transformations verified the same equipped to the satisfaction of well accredited inequality payable to Kraft [11].

The explorations conveyed by Kapur and Sharma [7] completed furtherance of communiqué flanked by the two collection of codes. To enhance the further literature on the above pointed out dissimilar codes, the authors by employing logarithmic function, conveniently and meticulously delivered certain convinced and significant annotations related with such transformations, that is, the connections of best one-one codes to UD codes by means of facilitation of subsequent mathematical outward show:

$$T_1' : l_i = b_i + b \lfloor \log b_i \rfloor + \log \left(\frac{D^b - 1}{D^b - 2} \right) \quad (1.18)$$

$$T_1'' : l_i = b_i + b \lfloor \log b_i \rfloor + \log \left(\frac{(D^b - 1)D^{b-1}}{D^{b-1} - 1} \right) \quad (1.19)$$

$$T_2' : l_i = b_i + b \lfloor \log(b_i + 1) \rfloor + \log(D^b - 2D^{b-1} + 1) - \log(D^{b-1} - 1) \quad (1.20)$$

$$T_2'' : l_i = b_i + b \lfloor \log(b_i + 1) \rfloor + \log(D^{b+1} - D^b - D^{b-1}) - \log(D^{b-1} - 1) - b \quad (1.21)$$

The discrete entropic models and the discrete divergence models convince a number of functionally obvious requirements and furthermore can be dispense operational association in many crucial pragmatic problems connected with numerous well surviving disciplines pertaining to dissimilar branches of mathematics. This is supplementary added at this juncture that some convinced coding techniques with the facilitation of the distance model in the discrete probability spaces and graph theory have been deliberated by well celebrated researchers Reviewed and Ferreira [22].

Recently, Schulte et al. [24] remarked numerous applications in communication system necessitate resembling target distributions to be surrounded by undersized informational divergence. The supplementary prerequisite of predictability frequently show the way to using

encoders that are one-to-one mappings. Nevertheless, even the best one-to-one encoders have divergences which produce logarithmically with the block length. To conquer this drawback, an encoder is projected that has an invertible one-to-many mapping and a low-rate random number generator. The authors have wrought two algorithms which confer information rates approaching the entropic value of the target distribution with exponentially declining divergence. Some additional pioneers who have completed momentous and fascinating deliberations for the facilitation of applications of their entropic models in an assortment fields of coding theory include Kawan and Yüksel [8], Yamamoto et al. [27], Yang et al. [28], Kochman et al. [10], Jose and Kulkarni [4], Chen et al. [2] etc.

In section 2, we have demonstrated some fascinating deliberations for the communication between discrete entropic models and the best 1:1 code. In section 3, we have endowed with the contribution of divergence models for the development of suitable codes whereas section 4 makes transactions with the outline of suitable codes through inaccuracy model. In the sequel, we provide discussions for the association between discrete entropic models and the best 1:1 code.

2. RELATIONS AMONG DISCRETE ENTROPIC MODELS AND BEST 1:1 MEAN LENGTHS

It is valuable to point out here that one-to-one codes are constantly believed to be nonsingular codes that allocate a distinctive codeword to every one source representation. These are as well acknowledged as “one-shot” encodings as these could be engaged when one merely desires to transmit a single source representation rather than a sequence of source cryptogram. Such a state of affairs can take place when the last message must be accredited previous to the next message. In view of the fact that our apprehension is to reduce the average length of the codeword, therefore for the preparation of these one-to-one codes, we should designate codewords to the letters which take place most repeatedly, that is, the codewords which have largest probability of happening. In developing the best 1:1 code, we first create the subsequent parametric mean:

$$L^{\alpha,a} = \frac{\alpha}{2(1-\alpha)} \log_D \left[\frac{\sum_{i=1}^n p_i D^{(1-\alpha)l_i/\alpha}}{\left\{ \sum_{i=1}^n p_i D^{(1-\alpha)l_i a/\alpha} \right\}^{-1/a}} \right]; \alpha \neq 1, \alpha > 0 \quad (2.1)$$

The exponentiated mean articulated in equation (2.1) convince the subsequent attractive properties of being a legitimate mean:

(i) If $l_1 = l_2 = l_3 = \dots = l_n = l$, then the generalized mean get hold of the subsequent appearance:

$$\begin{aligned} L^{\alpha,a} &= \frac{\alpha}{2(1-\alpha)} \log_D \left[\frac{D^{(1-\alpha)l/\alpha}}{\left\{ D^{(1-\alpha)la/\alpha} \right\}^{-1/a}} \right] \\ &= \frac{\alpha}{2(1-\alpha)} \left[(1-\alpha)l/\alpha + \frac{1}{a}(1-\alpha)la/\alpha \right] = l \end{aligned}$$

(ii) $L^{\alpha,a}$ must lie between extreme values of $l_1, l_2, l_3, \dots, l_n$

(iii) To find the limiting value of $L^{\alpha,a}$, we proceed as follows:

$$L^{\alpha,a} = \frac{\alpha \log_D \sum_{i=1}^n p_i D^{(1-\alpha)l_i/\alpha}}{2(1-\alpha)} + \frac{\alpha \log_D \sum_{i=1}^n p_i D^{(1-\alpha)l_i a/\alpha}}{2(1-\alpha)a}$$

$$\text{Thus } \lim_{\alpha \rightarrow 1} L^{\alpha,a} = \frac{\sum_{i=1}^n p_i l_i}{2} + \frac{\sum_{i=1}^n p_i l_i}{2} = \sum_{i=1}^n p_i l_i = L$$

$$\text{where } L = \sum_{i=1}^n p_i l_i$$

Consequently, we claim that exponentiated mean (2.1) persuade the mandatory properties of a proper mean and hence (2.1) provides an acceptable formula of mean.

Next, we exploit the length $L^{\alpha,a}$ deliberated in (2.1) and define the best 1:1 code by means of subsequent mathematical manifestation:

$$L_{1:1}^{\alpha,a} = \frac{\alpha}{2(1-\alpha)} \log_D \left(\frac{\sum_{i=1}^n p_i D^{\left\lceil \log_D \left(\frac{i}{2} + 1 \right) \right\rceil \frac{(1-\alpha)}{\alpha}}}{\left\{ \sum_{i=1}^n p_i D^{\left\lceil \log_D \left(\frac{i}{2} + 1 \right) \right\rceil \frac{(1-\alpha)}{\alpha} a} \right\}^{-1/a}} \right); \alpha \neq 1, \alpha > 0 \quad (2.2)$$

The subsequent theorem provides a lower bound for $L_{1:1}^{\alpha,a}$. For proving the theorem, we make use of the modified version of Renyi's [21] discrete entropic model denoted by $R_\alpha(P)$ and provided by means of subsequent precise materialization:

$$R_{\alpha,a}(P) = \frac{1}{(1-\alpha)a} \log_D \left[\sum_{i=1}^n p_i^\alpha \right]; \alpha > 0, \alpha \neq 1 \quad (2.3)$$

Theorem-2.1: The lower bound for the best 1:1 code $L_{1:1}^{\alpha,a}$ persuades the subsequent inequality connecting the discrete entropic model and $L_{1:1}^{\alpha,a}$:

$$L_{1:1}^{\alpha,a} \geq \frac{1}{2} \left[R_\alpha(P) + R_{\alpha,a}(P) - \log_D \left\{ \left\{ \sum_{i=1}^n \left(\frac{2}{i+2} \right) \right\} \left\{ \sum_{i=1}^n \left(\frac{2}{i+2} \right)^{-a} \right\} \right\} \right]; \alpha \neq 1, \alpha > 0 \quad (2.4)$$

Proof: By employing the definition provided in equation (2.2), we have the subsequent structure of the best 1:1 code $L_{1:1}^{\alpha,a}$:

$$L_{1:1}^{\alpha,a} \geq \frac{\alpha}{2(1-\alpha)} \log_D \left(\frac{\sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha}}}{\left\{ \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha} a} \right\}^{-1/a}} \right)$$

Thus, the relation surrounded by the three entities $R_\alpha(P)$, $R_{\alpha,a}(P)$ and $L_{1:1}^{\alpha,a}$ can be prepared capable of the precise mathematical appearance as deliberated below:

$$\begin{aligned} \frac{1}{2} [R_\alpha(P) + R_{\alpha,a}(P)] - L_{1:1}^{\alpha,a} &\leq \frac{1}{2} \left[R_\alpha(P) - \frac{\alpha}{2(1-\alpha)} \log_D \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha}} \right] \\ &+ \frac{1}{2} \left[R_{\alpha,a}(P) - \frac{\alpha}{2(1-\alpha)a} \log_D \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha} a} \right] \end{aligned} \quad (2.5)$$

To provide the solution, we put $\frac{1-\alpha}{\alpha} = t$ in the above brought up equation so as to make available the term of equation (2.5) in the subsequent appearance:

$$\frac{1}{2} \left[R_\alpha(P) - \frac{\alpha}{2(1-\alpha)} \log_D \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha}} \right]$$

$$= \frac{1}{2} \left[\frac{1+t}{t} \log_D \sum_{i=1}^n p_i^{1/1+t} - \frac{1}{2t} \log_D \sum_{i=1}^n p_i D^{t \log_D \left(\frac{i}{2} + 1 \right)} \right]$$

Thus

$$\frac{1}{2} \left[R_\alpha(P) - \frac{\alpha}{2(1-\alpha)} \log_D \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha}} \right] = \frac{1}{2} \log_D \left[\left\{ \sum_{i=1}^n p_i^{1/1+t} \right\}^{\frac{1+t}{t}} \left\{ \sum_{i=1}^n p_i \left(\frac{i}{2} + 1 \right)^t \right\}^{-1/2t} \right] \quad (2.6)$$

Now, we make use of Holder's inequality by substituting

$$x_i = p_i^{\frac{1}{t}}, \quad y_i = p_i^{-\frac{1}{t}} \left\{ \frac{i}{2} + 1 \right\}^{-1}, \quad p = \frac{t}{1+t}, \quad q = -t, \quad \text{we get hold of the subsequent relation:}$$

$$\left\{ \sum_{i=1}^n p_i^{1/1+t} \right\}^{\frac{1+t}{t}} \left\{ \sum_{i=1}^n p_i \left(\frac{i}{2} + 1 \right)^t \right\}^{-1/2t} \leq \sum_{i=1}^n \left\{ \frac{2}{i+2} \right\}$$

Furthermore, upon attracting logarithmic function in the company of the above inequality and by way of uncomplicated computations, we get hold of the subsequent quantitative outward show:

$$\frac{1}{2} \log_D \left[\left\{ \sum_{i=1}^n p_i^{1/1+t} \right\}^{\frac{1+t}{t}} \left\{ \sum_{i=1}^n p_i \left(\frac{i}{2} + 1 \right)^t \right\}^{-1/2t} \right] \leq \frac{1}{2} \log_D \sum_{i=1}^n \left\{ \frac{2}{i+2} \right\} \quad (2.7)$$

Using (2.7) in (2.6), we get the consequent manifestation:

$$\frac{1}{2} \left[R_\alpha(P) - \frac{\alpha}{2(1-\alpha)} \log_D \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha}} \right] \leq \frac{1}{2} \log_D \sum_{i=1}^n \left\{ \frac{2}{i+2} \right\} \quad (2.8)$$

Proceeding on similar lines and taking the subsequent substitution

$$x_i = p_i^{\frac{1}{t}}, \quad y_i = p_i^{-\frac{1}{t}} \left\{ \frac{i}{2} + 1 \right\}^{-a}, \quad p = \frac{t}{1+t}, \quad q = -t \quad \text{in Holder's inequality, we acquire the second term on}$$

the R.H.S. of equation (2.5) in the succeeding appearance:

$$\frac{1}{2} \left[R_{\alpha,a}(P) - \frac{\alpha}{2(1-\alpha)a} \log_D \sum_{i=1}^n p_i D^{\log_D \left(\frac{i}{2} + 1 \right) \frac{(1-\alpha)}{\alpha} a} \right] \leq \frac{1}{2} \log_D \sum_{i=1}^n \left\{ \frac{2}{i+2} \right\}^{-a} \quad (2.9)$$

By employing equations (2.9) and (2.8) in (2.5), we get hold of the consequent quantitative outward show:

$$\frac{1}{2} [R_{\alpha}(P) + R_{\alpha,a}(P)] - L_{1:1}^{\alpha,a} \leq \frac{1}{2} \left[\log_D \left\{ \left\{ \sum_{i=1}^n \left(\frac{2}{i+2} \right) \right\} \left\{ \sum_{i=1}^n \left(\frac{2}{i+2} \right)^{-a} \right\} \right\} \right] \text{ which proves the theorem.}$$

We, next define another mean involving two parameters and connecting the two categories of codes, that is, the best 1:1 code $L_{1:1}^{\alpha,\beta}$. By our hypothesis this new-fangled and an inventive definition of mean capture the successive configuration:

$$L_{1:1}^{\alpha,\beta} = \frac{1}{(\alpha-1)} \log_D \left(\frac{\sum_{i=1}^n p_i^{\beta} D^{(\alpha-1) \left\lceil \log_D \left(\frac{i}{2} + 1 \right) \right\rceil}}{\sum_{i=1}^n p_i^{\beta}} \right); \alpha \neq 1, \alpha > 0 \quad (2.10)$$

The above mean length (2.10) corresponds to well acknowledged mean length already deliberated by Kapur's [6] in the succeeding form

$$L_{UD}^{\alpha,\beta} = \frac{1}{\alpha-1} \log_D \left(\frac{\sum_{i=1}^n p_i^{\beta} D^{(\alpha-1)l_i}}{\sum_{i=1}^n p_i^{\beta}} \right) \quad (2.11)$$

and well accredited Kapur's [5] discrete entropic model previously accessible in the literature and shaped in consequent manifestation:

$$E_{\alpha,\beta}(P) = \frac{1}{\alpha-1} \log_D \left(\frac{\left(\sum_{i=1}^n p_i^{\frac{\beta}{\alpha}} \right)^{\alpha}}{\sum_{i=1}^n p_i^{\beta}} \right) \quad (2.12)$$

The following theorem provides a correspondence between $L_{1:1}^{\alpha,\beta}$ and $L_{UD}^{\alpha,\beta}$.

Theorem-2.2: The generated codes $L_{1:1}^{\alpha,\beta}$ and $L_{UD}^{\alpha,\beta}$ convince the subsequent inequality:

$$L_{UD}^{\alpha,\beta} - L_{1:1}^{\alpha,\beta} \leq 2 + \log_D \sum_{i=1}^n \left(\frac{1}{i+2} \right) \quad (2.13)$$

Proof: From equation (2.10), we have the subsequent mathematical nature of $L_{1:1}^{\alpha,\beta}$:

$$L_{1:1}^{\alpha,\beta} \geq \frac{1}{(\alpha-1)} \log_D \left(\frac{\sum_{i=1}^n p_i^\beta \left(\frac{i}{2} + 1 \right)^{(\alpha-1)}}{\sum_{i=1}^n p_i^\beta} \right)$$

Putting $\alpha-1=t$ in the above equation and employing (2.12), we may acquire the primarily well-built communication by employing Holder's inequality with the subsequent substitution:

$$x_i = p_i^{\frac{\beta}{t}}, \quad y_i = p_i^{-\frac{\beta}{t}} \left\{ \frac{i}{2} + 1 \right\}^{-1}, \quad p = \frac{t}{1+t}, \quad q = -t$$

Consequently, upon attracting logarithmic employment and moreover by means of straightforward computations, we acquire the subsequent quantitative appearance providing association flanked by discrete entropic model and the best 1:1 code:

$$E_{\alpha,\beta}(P) - L_{1:1}^{\alpha,\beta} \leq \frac{1}{t} \log_D \left(\frac{\left(\sum_{i=1}^n p_i^{\frac{\beta}{t+1}} \right)^{t+1}}{\sum_{i=1}^n p_i^\beta} \right) - \frac{1}{t} \log_D \left(\frac{\sum_{i=1}^n p_i^\beta \left(\frac{i}{2} + 1 \right)^t}{\sum_{i=1}^n p_i^\beta} \right)$$

Moreover by way of straightforward mathematical computations, we get hold of the subsequent inequality wrought in the well built manifestation:

$$E_{\alpha,\beta}(P) - L_{1:1}^{\alpha,\beta} \leq \log_D \left[\sum_{i=1}^n \frac{2}{i+2} \right]$$

$$\Rightarrow L_{1:1}^{\alpha,\beta} \geq E_{\alpha,\beta}(P) - \log_D \left(\sum_{i=1}^n \frac{2}{i+2} \right)$$

Employing equation (2.12), we search out the subsequent consequence:

$$L_{UD}^{\alpha,\beta} - L_{1:1}^{\alpha,\beta} < 1 + E_{\alpha,\beta}(P) - L_{1:1}^{\alpha,\beta} \leq 1 + \log \left(\sum_{i=1}^n \frac{2}{i+2} \right) = 2 + \log \left(\sum_{i=1}^n \frac{1}{i+2} \right)$$

Consequently, the theorem gets proved.

Important Observations: The study of correspondence between discrete information entropic models and the mean codeword lengths reveals the following most advantageous consequences:

- (i) Own mean codeword length and the best 1:1 code length.
- (ii) Development of communication flanked by best 1:1 code and UD code.

Note-I: It is worth mentioning at this juncture that a code is believed to be suitable if the codeword lengths l_i persuade a suitable connection in such a manner that the minimum of the specified length for each and every one code satisfying the given association has a specific value or the lower bound of the mean for every single one code lies stuck between two specific values. In addition, it has to be investigational that Kraft's inequality which contributes with a crucial answerability in demonstrating a noiseless coding theorem cannot be tailored in a subjective approach, provoked by the ambition to demonstrate a new literature for the development of UD codes. If one formulates adjustment in this inequality, the formulated codes will appear with different configuration other than rewarding the stipulation of unique decipherability. Kapur [6] emphasized that each UD code ought to persuade Kraft's inequality nevertheless for the development of suitable codes customized this inequality. The modified version of this inequality is prescribed subsequently:

$$\frac{1}{D} \leq \sum_{i=1}^n D^{-l_i} \leq 1$$

With this development, Kapur [6] completed scrupulous study for the environment of dissimilar suitable codes depending upon diverse circumstances. These suitable codes can be prepared by the employment of an assortment of information models including entropic models, divergence models and inaccuracy models. It is supplementary emphasized that if we desire to prepare suitable codes through entropic models, then the subsequent inequality should hold good:

$$\sum_{i=1}^n D^{-l_i} \leq 1$$

In case our aspiration is to develop suitable codes via discrete distance models, then the subsequent inequality should be satisfied:

$$\sum_{i=1}^n p_i q_i D^{-l_i} \leq 1$$

In case we wish for the development of suitable codes all the way through inaccuracy models, then the subsequent inequality should be employed:

$$\sum_{i=1}^n p_i q_i^{-1} D^{-l_i} \leq 1$$

Note-II: It is frequently enviable to compute the discrepancy flanked by two probability distributions for a given random variable. This happens recurrently in machine learning problems when we may be concerned in manipulating the dissimilarity between real and observed probability distribution. This can be accomplished by means of procedures usually available in the field of information theory, such as the Kullback-Leibler [12] divergence (KL divergence) and Sibson [23] divergence also called Jensen-Shannon divergence that make available a normalized and symmetrical description of the KL divergence. These achievable procedures can be employed as shortcuts in the computations of other extensively used procedures such as mutual information for feature selection earlier to modeling and cross-entropy second-handed as a loss function for numerous dissimilar classifier models. From the above conversation, one can meaningfully formulate the subsequent observations:

- (i) Statistical distance is the universal encouragement of influencing the differentiation sandwiched statistical substances resembling dissimilar probability distributions for a random variable.
- (ii) Kullback-Leibler divergence works out a score that measures the divergence of one probability distribution commencing from an additional one.
- (iii) Jensen-Shannon divergence broadens KL divergence to compute a unprejudiced score and distance measure of one probability distribution commencing from an additional one.

Additionally, it is added that in real life situation, there are countless circumstances where we may desire to compare two probability distributions. Particularly, we might have a single random variable and two dissimilar probability distributions for the variable, such as an accurate distribution and an estimate of that distribution. In such situations, it can be constructive idea to quantify the distinction flanked by the distributions. Commonly, this is referred to as the problem of manipulating the statistical distance flanked by the two statistical objects. One advancement is to compute a distance measure connecting the two distributions. This can be challenging as it can be complicated to understand the measure. As an alternative, it is more widespread to compute a divergence flanked by the two probability distributions. A divergence is similar to a measure but is not symmetrical. This means that a divergence is a scoring of how one distribution fluctuates

from another, where computing the divergence for distributions P and Q would furnish a diverse score from Q and P .

Below, we illustrate the association surrounded by the suitable codes and the weighted and non-weighted discrete divergence models.

3. CONTRIBUTION OF DISCRETE DIVERGENCE MODELS FOR THE DEVELOPMENT OF SUITABLE CODES

Theorem-3.1: If $l_1, l_2, l_3, \dots, l_n$ are the lengths of a code, then the subsequent inequality is true:

$$D_{\alpha}^{\beta}(P:Q) \leq \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}}; \alpha \neq \beta, \alpha \leq 1, \beta \geq 1 \text{ or } \alpha \geq 1, \beta \leq 1 \quad (3.1)$$

where

$$D_{\alpha}^{\beta}(P:Q) = \frac{1}{\alpha - \beta} \log_D \left[\frac{\sum_{i=1}^n p_i^{\alpha} q_i^{1-\alpha}}{\sum_{i=1}^n p_i^{\beta} q_i^{1-\beta}} \right]; \alpha \neq \beta, \alpha \leq 1, \beta \geq 1 \text{ or } \alpha \geq 1, \beta \leq 1 \quad (3.2)$$

is a discrete parametric divergence model developed by Kapur [5].

Proof: To provide evidence for the establishment of the theorem, we make use of Holder's inequality and Kapur's [6] inequality given by the subsequent manifestation:

$$\sum_{i=1}^n p_i q_i D^{-l_i} \leq 1 \quad (3.3)$$

Substituting $x_i = p_i^{\frac{1-\alpha^2}{\alpha-1}} q_i^{1+\alpha} D^{-l_i}$, $y_i = p_i^{\frac{\alpha^2}{\alpha-1}} q_i^{-\alpha}$, $p = 1-\alpha$, $q = \frac{\alpha-1}{\alpha}$ in well known inequality

payable to Holder, we acquire the subsequent materialization:

$$\sum_{i=1}^n p_i q_i D^{-l_i} \geq \left[\sum_{i=1}^n \left(p_i^{\frac{1-\alpha^2}{\alpha-1}} q_i^{1+\alpha} D^{-l_i} \right)^{1-\alpha} \right]^{\frac{1}{1-\alpha}} \left[\sum_{i=1}^n \left(p_i^{\frac{\alpha^2}{\alpha-1}} q_i^{-\alpha} \right)^{\frac{\alpha-1}{\alpha}} \right]^{\frac{\alpha}{\alpha-1}}$$

$$\text{or} \left[\sum_{i=1}^n \left(p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right]^{\frac{1}{1-\alpha}} \left[\sum_{i=1}^n \left(p_i^{\alpha} q_i^{1-\alpha} \right) \right]^{\frac{\alpha}{\alpha-1}} \leq \sum_{i=1}^n p_i q_i D^{-l_i}$$

Taking logarithms both sides and applying the inequality (3.3), the above equation furnishes the subsequent nature:

$$\frac{1}{1-\alpha} \log_D \left[\sum_{i=1}^n \left(p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right] + \frac{\alpha}{\alpha-1} \log_D \left[\sum_{i=1}^n \left(p_i^\alpha q_i^{1-\alpha} \right) \right] \leq 0$$

or $\frac{1}{1-\alpha} \log_D \left[\sum_{i=1}^n \left(p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right] \leq \frac{\alpha}{1-\alpha} \log_D \left[\sum_{i=1}^n \left(p_i^\alpha q_i^{1-\alpha} \right) \right]$ (3.4)

If $\alpha \leq 1$, then equation (3.4) provides

$$\log_D \left[\sum_{i=1}^n \left(p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right]^{1/\alpha} \leq \log_D \left[\sum_{i=1}^n \left(p_i^\alpha q_i^{1-\alpha} \right) \right] \quad (3.5)$$

Similarly, for $\beta \geq 1$, we can proceed to prove the following consequence:

$$-\log_D \left[\sum_{i=1}^n \left(p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right) \right]^{1/\beta} \leq -\log_D \left[\sum_{i=1}^n \left(p_i^\beta q_i^{1-\beta} \right) \right] \quad (3.6)$$

From equations (3.5) and (3.6), we get hold of the subsequent outcome:

$$\log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}} \leq \log_D \frac{\sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}}{\sum_{i=1}^n p_i^\beta q_i^{1-\beta}}$$

Now, since $\alpha \leq 1$ and $\beta \geq 1$, we acquire $\alpha - \beta < 0$. By the employment of these results, the above equation entails that

$$\frac{1}{\alpha - \beta} \log_D \frac{\sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}}{\sum_{i=1}^n p_i^\beta q_i^{1-\beta}} \leq \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}} \quad (3.7)$$

Proceeding on similar lines, if $\alpha \geq 1$, then after multiplying equation (3.4) by -1, we acquire the subsequent form:

$$\log_D \left[\sum_{i=1}^n \left(p_i^\alpha q_i^{1-\alpha} \right) \right] \leq \log_D \left[\sum_{i=1}^n \left(p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right]^{1/\alpha} \quad (3.8)$$

Similarly, for $\beta \leq 1$, we can proceed to prove the subsequent outcome:

$$\log_D \left[\sum_{i=1}^n \left(p_i^\beta q_i^{1-\beta} \right) \right] \geq \log_D \left[\sum_{i=1}^n \left(p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right) \right]^{1/\beta}$$

Again, multiplying the above equation by -1, we acquire

$$-\log_D \left[\sum_{i=1}^n \left(p_i^\beta q_i^{1-\beta} \right) \right] \leq -\log_D \left[\sum_{i=1}^n \left(p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right) \right]^{1/\beta} \quad (3.9)$$

Adding equations (3.8) and (3.9), we acquire the subsequent appearance:

$$\log_D \frac{\sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}}{\sum_{i=1}^n p_i^\beta q_i^{1-\beta}} \leq \log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}} \quad (3.10)$$

Now, since $\alpha \geq 1$ and $\beta \leq 1$, we have $\alpha - \beta > 0$. By the employment of these results, the above equation (3.10) brings about the subsequent inequality:

$$\frac{1}{\alpha - \beta} \log_D \frac{\sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}}{\sum_{i=1}^n p_i^\beta q_i^{1-\beta}} \leq \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}} \text{ which is (3.7).}$$

Consequently, in every case we observe that the subsequent correspondence holds good:

$$D_\alpha^\beta(P:Q) \leq \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}} \text{ which proves the theorem.}$$

Theorem-3.2: If $l_1, l_2, l_3, \dots, l_n$ are the lengths of a code, then the subsequent weighted inequality always holds good:

$$D_{\alpha}^{\beta}(P:Q;W) \leq \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n w_i^{\alpha} p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n w_i^{\beta} p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}}; \alpha \neq \beta, \alpha \leq 1, \beta \geq 1 \text{ or } \alpha \geq 1, \beta \leq 1 \quad (3.11)$$

where

$$D_{\alpha}^{\beta}(P:Q;W) = \frac{1}{\alpha - \beta} \log_D \left[\frac{\sum_{i=1}^n w_i p_i^{\alpha} q_i^{1-\alpha}}{\sum_{i=1}^n w_i p_i^{\beta} q_i^{1-\beta}} \right]; \alpha, \beta > 0; \alpha \leq 1, \beta \geq 1 \text{ or } \alpha \geq 1, \beta \leq 1 \quad (3.12)$$

is a discrete generalized parametric divergence model developed by Parkash and Kumar [17] for the discrete weighted distribution.

Proof: To provide evidence for the survival of the theorem, we once more bring into play the well recognized Holder's inequality and Kapur's [6] inequality provided in equation (3.2) to be made applicable for the construction of suitable code. Moreover, as accustomed we formulate the subsequent substitution for confirmation of the coding theorem:

$$x_i = w_i^{\frac{\alpha}{\alpha-1}} p_i^{\frac{1-\alpha^2}{\alpha-1}} q_i^{1+\alpha} D^{-l_i}, \quad y_i = w_i^{\frac{\alpha}{\alpha-1}} p_i^{\frac{\alpha^2}{\alpha-1}} q_i^{-\alpha}, \quad p = 1 - \alpha, \quad q = \frac{\alpha - 1}{\alpha} \quad (3.13)$$

With this substitution (3.13) in well known Holder's inequality, we get hold of the subsequent materialization:

$$\left[\sum_{i=1}^n \left(w_i^{\alpha} p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right]^{\frac{1}{1-\alpha}} \left[\sum_{i=1}^n \left(w_i p_i^{\alpha} q_i^{1-\alpha} \right) \right]^{\frac{\alpha}{\alpha-1}} \leq \sum_{i=1}^n p_i q_i D^{-l_i}$$

The employment of logarithmic function and the inequality (3.3), the above formulated equation delivers the subsequent character:

$$\frac{1}{1-\alpha} \log_D \left[\sum_{i=1}^n \left(w_i^{\alpha} p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right] \leq \frac{\alpha}{1-\alpha} \log_D \left[\sum_{i=1}^n \left(w_i p_i^{\alpha} q_i^{1-\alpha} \right) \right] \quad (3.14)$$

If $\alpha \leq 1$, then equation (3.14) provides the subsequent form:

$$\log_D \left[\sum_{i=1}^n \left(w_i^\alpha p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right) \right]^{1/\alpha} \leq \log_D \left[\sum_{i=1}^n \left(w_i p_i^\alpha q_i^{1-\alpha} \right) \right] \quad (3.15)$$

Similarly, for $\beta \geq 1$, we can proceed to prove the following consequence:

$$-\log_D \left[\sum_{i=1}^n \left(w_i^\alpha p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right) \right]^{1/\beta} \leq -\log_D \left[\sum_{i=1}^n \left(w_i p_i^\beta q_i^{1-\beta} \right) \right] \quad (3.16)$$

Now, since $\alpha \leq 1$ and $\beta \geq 1$, we must have $\alpha - \beta < 0$. By the employment of this actuality, the above equation (3.15) and (3.16) bring about the subsequent mathematical appearance:

$$\frac{1}{\alpha - \beta} \log_D \left[\frac{\sum_{i=1}^n w_i^\alpha p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)}}{\sum_{i=1}^n w_i p_i^\beta q_i^{1-\beta}} \right] \leq \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n w_i^\alpha p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n w_i p_i^\beta q_i^{1-\beta} \right\}^{1/\beta}} \quad (3.17)$$

Proceeding on comparable lines, that is, if $\alpha \geq 1$ and $\beta \leq 1$, then $\alpha - \beta > 0$, the expression (3.17) immediately gets verified. Consequently, in every case we scrutinize that the manifestation (3.17) always holds good and the theorem gets established.

Note-III: From the above theorem, we have developed subsequent suitable codes, one through Kapur's [5] discrete divergence model and another one through Parkash and Kumar's [17] discrete weighted parametric divergence model for the discrete probability space:

$$1. L_1^s = \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}}$$

$$2. {}_w L_1^s = \frac{1}{\alpha - \beta} \log_D \frac{\left\{ \sum_{i=1}^n w_i^\alpha p_i^{1-\alpha+\alpha^2} q_i^{(1-\alpha^2)} D^{-l_i(1-\alpha)} \right\}^{1/\alpha}}{\left\{ \sum_{i=1}^n w_i^\beta p_i^{1-\beta+\beta^2} q_i^{(1-\beta^2)} D^{-l_i(1-\beta)} \right\}^{1/\beta}}$$

Proceeding on comparable defenses, one can create countless new innovative weighted and non-weighted parametric suitable codes through weighted and non-weighted discrete parametric divergence models.

Note-IV: It is further observed that in an experimentation dealing with the proclamation about probabilities of dissimilar events, two varieties of errors are plausible, explicitly one because of the nonappearance of adequate data or indistinctness in test results and other from erroneous data. Shannon's information model can be second-handed to enlighten the error because of ambiguity only whereas the both types of errors can be explained by using a measure identified as measure of inaccuracy which ascertains applications in statistical inference and a concept anticipated by Kerridge [9]. Different instigators anticipated new inaccuracy models for the reasons that of their applicability in statistics, coding theory and supplementary associated fields.

Sathar et. al [26] made investigations about the past inaccuracy model and consequently recommended nonparametric estimators for these models. The authors made rigorous study of the asymptotic properties of these estimators under convinced appropriate and reliability conditions. Additionally, the authors made comparisons for the performance of the projected estimators by employing Monte-Carlo simulation technique. Numerous investigators projected their own inaccuracy models for delivering their applications in the discipline of coding theory. Proceeding on comparable lines and following Kapur's [6] advancement, one can prepare suitable codes for the supplementary information models including inaccuracy models. In the sequel, we exemplify the significant association surrounded by the suitable codes and the inaccuracy models.

4. DEVELOPMENT OF SUITABLE CODE VIA INACCURACY MEASURE

Theorem-4.1: If $l_1, l_2, l_3, \dots, l_n$ are the lengths of a code, then the subsequent inequality is factual:

$$K^\alpha(P:Q) \leq \frac{1}{\alpha(\alpha-1)} \log_D \left[\frac{\sum_{i=1}^n p_i^{\alpha^2-\alpha+1} q_i^{-(1-\alpha^2)} D^{-(1-\alpha)l_i}}{\left\{ \sum_{i=1}^n p_i^\alpha \right\}^\alpha} \right]; \alpha > 1 \quad (4.1)$$

where

$$K^\alpha(P:Q) \leq \frac{1}{(\alpha-1)} \log_D \frac{\sum_{i=1}^n p_i^\alpha q_i^{(1-\alpha)}}{\sum_{i=1}^n p_i^\alpha}; \alpha > 1 \quad (4.2)$$

is Kapur's [5] inaccuracy model of order α .

Proof: To establish the above theorem, we bring into play Holder's inequality and Kapur's [6] inequality prearranged by the subsequent demonstration:

$$\sum_{i=1}^n p_i q_i^{-1} D^{-l_i} \leq 1 \quad (4.3)$$

Substituting $x_i = p_i^{1-\frac{\alpha^2}{\alpha-1}} q_i^{\alpha-1} D^{-l_i}$, $y_i = p_i^{\frac{\alpha^2}{\alpha-1}} q_i^{-\alpha}$, $p = 1-\alpha$, $q = \frac{\alpha-1}{\alpha}$ in Holder's inequality, we get hold of the subsequent inequality:

$$\sum_{i=1}^n p_i q_i^{-1} D^{-l_i} \geq \left[\sum_{i=1}^n \left(p_i^{1-\frac{\alpha^2}{\alpha-1}} q_i^{\alpha-1} D^{-l_i} \right)^{1-\alpha} \right]^{\frac{1}{1-\alpha}} \left[\sum_{i=1}^n \left(p_i^{\frac{\alpha^2}{\alpha-1}} q_i^{-\alpha} \right)^{\frac{\alpha-1}{\alpha}} \right]^{\frac{\alpha}{\alpha-1}}$$

With the facilitation of logarithmic function and the application of inequality (4.3), we acquire the subsequent relation:

$$\frac{1}{1-\alpha} \log_D \left[\sum_{i=1}^n p_i^{\alpha^2-\alpha+1} q_i^{-(1-\alpha)^2} D^{-(1-\alpha)l_i} \right] + \frac{\alpha}{\alpha-1} \log_D \left[\sum_{i=1}^n p_i^{\alpha} q_i^{1-\alpha} \right] \leq 0$$

$$\text{or } \frac{1}{(\alpha-1)} \log_D \left[\sum_{i=1}^n p_i^{\alpha} q_i^{1-\alpha} \right] \leq \frac{1}{\alpha(1-\alpha)} \log_D \left[\sum_{i=1}^n p_i^{\alpha^2-\alpha+1} q_i^{-(1-\alpha)^2} D^{-(1-\alpha)l_i} \right]$$

Thus, we have the subsequent expression:

$$K^{\alpha}(P:Q) \leq \frac{1}{\alpha(\alpha-1)} \log_D \left[\frac{\sum_{i=1}^n p_i^{\alpha^2-\alpha+1} q_i^{-(1-\alpha)^2} D^{-(1-\alpha)l_i}}{\left\{ \sum_{i=1}^n p_i^{\alpha} \right\}^{\alpha}} \right]; \alpha > 1$$

Consequently the theorem gets established.

Note-V: Commencing from the above theorem, we have shaped a subsequent suitable code through Kapur's [5] inaccuracy model:

$$L_2^s = \frac{1}{\alpha(\alpha-1)} \log_D \left[\frac{\sum_{i=1}^n p_i^{\alpha^2-\alpha+1} q_i^{-(1-\alpha^2)} D^{-(1-\alpha)l_i}}{\left\{ \sum_{i=1}^n p_i^\alpha \right\}^\alpha} \right]; \alpha > 1$$

With the analogous arguments discussed above, one can produce innumerable new-fangled parametric suitable codes through inaccuracy models.

Concluding Remarks: In the literature of theory of coding, a multiplicity of ground-breaking mean lengths can be produced subsequent to the inspection of a collection of discrete entropic, divergence and inaccuracy models. The foremost intention of our learning in this communication is to widen the literature on best 1:1 codes, the UD codes and suitable codes. The present comprehensive study is an accurate footstep in this direction in which we have completed the illustration for providing the association between discrete entropic models and the best 1:1 codes. Our wide-ranging study reveals the information that there survive an involvement sandwiched between own mean codeword length and best 1:1 code. Additionally, we have fabricated a well-built relation flanked by best 1:1 and UD code. Furthermore, we have provided the contribution of divergence and inaccuracy models for the development of suitable codes. This innovative inspiration can be made capable of extension with the generation of an assortment of new-fangled information entropic models along with the conversation of their association with coding theory. With analogous judgment and by engendering a multiplicity of continuous entropic models, this comprehensive learning can be completed for an assortment of codes together with best 1:1, suitable codes and UD codes.

REFERENCES

- [1] Campbell, L.L. (1965). A coding theorem and Renyi's entropy. *Information and Control*, 8: 423-429.
- [2] Chen, J., Xie, L. and Chang, Y., Wang, J. and Wang, Y. (2020). Generalized Gaussian Multiterminal Source Coding: The Symmetric Case. *IEEE Transactions on Information Theory* 66(4): 2115-2128.
- [3] Huang, W. and Zhang, K. (2019). Approximations of Shannon mutual information for discrete variables with applications to neural population coding. *Entropy* 21(3): Paper No. 243, 21 pp.

- [4] Jose, S. T. and Kulkarni, A. A. (2020). Shannon Meets von Neumann: A Minimax Theorem for Channel Coding in the Presence of a Jammer. *IEEE Transactions on Information Theory* 66(5): 2842-2859.
- [5] Kapur, J.N. (1994). *Measures of Information and their Applications*. Wiley Eastern, New York.
- [6] Kapur, J.N. (1998). *Entropy and Coding*. Mathematical Sciences Trust Society, New Delhi.
- [7] Kapur, J.N. and Sharma, R. (1994). Construction of Efficient Uniquely decipherable codes via Best One-One codes. MSTS Res Report No. 684.
- [8] Kawan, C. and Yüksel, S. (2018). On optimal coding of non-linear dynamical systems. *IEEE Trans. Inform. Theory* 64 (10): 6816–6829.
- [9] Kerridge, D. F. (1961). Inaccuracy and inference. *Journal of Royal Statistical Society Series B* 23:184-194.
- [10] Kochman, Y., Ordentlich, O. and Polyanskiy, Y. (2020). A Lower Bound on the Expected Distortion of Joint Source-Channel Coding. *IEEE Transactions on Information Theory*, DOI: 10.1109/TIT.2020.2983148
- [11] Kraft, L.G. (1949). A Device for Quantizing Grouping and Coding Amplitude Modulated Pulses. M.S. Thesis, Electrical Engineering Department, MIT.
- [12] Kullback, S. and Leibler, R. A. (1951). On information and sufficiency. *Annals of Mathematical Statistics* 22: 79-86.
- [13] Leung-Yan-Cheong, S.K. and Cover, T.M. (1978). Some equivalences between Shannon entropy and Kolmogorov complexity. *IEEE Trans. Inf. Theory* 24:331–338.
- [14] Parkash, O. and Kakkar, P. (2012). Development of two new mean codeword lengths. *Information Sciences* 207: 90-97.
- [15] Parkash, O. and Kakkar, P. (2014). New information theoretic models, their detailed properties and new inequalities. *Canadian Journal of Pure and Applied Sciences* 8(3): 3115-3123.
- [16] Parkash, O. and Kakkar, P. (2014). Generating entropy measures through quasilinear entropy and its generalized forms, In: Parkash. O. (ed). *Mathematical Modeling and applications*, Lambert Academic Publishers, pp. 42-53.
- [17] Parkash, O. and Kumar, R. (2020). Advancement of mean codeword lengths through weighted divergence model. *International Journal of Advanced Science and Technology* 29(5s): 1903-1910.

- [18] Parkash, O., Kumar, R. and Kakkar, P. (2017). Two new weighted divergence measures in probability spaces. *International Journal of Pure and Applied Mathematics* 116(19), 643-653.
- [19] Parkash, O., Kumar, R., Mukesh and Kakkar, P. (2019). Weighted measures and their correspondence with coding theory. *A Journal of Composition Theory* 12(9): 1670-1680.
- [20] Parkash, O., Sharma, R. and Singh, V. (2022). A new Discrete Information Model and its Applications for the study of Contingency Tables. *Journal of Discrete Mathematical Sciences and Cryptography*. 25(3): 785-792.
- [21] Renyi, A. (1961). On measures of entropy and information. *Proceedings 4th Berkeley Symposium on Mathematical Statistics and Probability* 1: 547-561.
- [22] Reviewed, O, K. and Ferreira, H. C. (2019). New distance concept and graph theory approach for certain coding techniques design and analysis. *Commun. Appl. Ind. Math.* 10: 53–70.
- [23] Sibson, R. (1969). Information radius. *Zeitschrift fur Wahrscheinlichkeitstheorie und Verwandte Gebiete* 14:149-160.
- [24] Schulte, P., Amjad, R.A., Wiegart, T. and Kramer, G. (2022). Invertible low-divergence coding. *IEEE Transactions on Information Theory* 68(1): 178–192.
- [25] Shannon, C. E. (1948). A mathematical theory of communication. *Bell. Syst. Tech. J.* 27: 379–423, 623–659.
- [26] Sathar, E. I. A., Viswakala, K. V. and Rajesh, G. (2021). Estimation of past inaccuracy measure for the right censored dependent data. *Communications in Statistics: Theory and Methods* 50(6): 1446–1455.
- [27] Yamamoto, N., Nakano, K. and Ito, Y., Takafuji, D., Kasagi, A. and Tabaru, T. (2020). A New Huffman Code for Accelerating GPU Decompression (preliminary version). *Bulletin of Networking, Computing, Systems, and Software* 9(1): 91–93.
- [28] Yang, Y., Guo, S. and Liu, G. and Wang, Q. (2020). Joint Source Coding Rate Allocation and Flow Scheduling for Data Aggregation in Collaborative Sensing Networks. *Computer Networks* 175:107269.