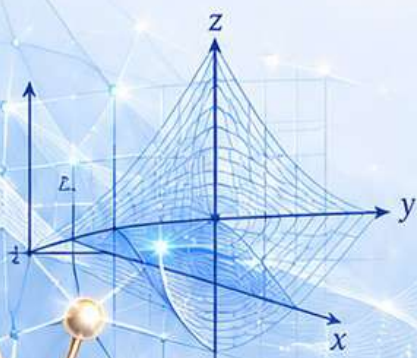




Al-Hamdaniya Journal for Pure and Applied Mathematics

مجلة الحمدانية للرياضيات البحتة والتطبيقية



$$\sum_{n=1}^{\infty} \frac{f(s)}{2n} = sI, dg = \lim_{n \rightarrow \infty}$$

$$\frac{((-1) \cdot ((-1) - 1))}{2n}$$

$$f(s, s) = \int_a^b \frac{L}{72 - hhn}$$



$$\int_a^c \int_b^c f(t, t) dt dx$$

$$iE = (k.s.f(y, '(t) = 0$$

$$\pi = kyt \quad \forall (x; k)$$



www.uohamdaniya.edu.iq

About Journal

Al-Hamdaniya Journal for Pure and Applied Mathematics (HJPAM) is an international, peer reviewed, double-blind, open access journal that publishes high quality research at the intersection of Mathematics, Statistics. It aims to provide a forum for theoretical, methodological and applied contributions in areas such as algebraic geometry, algebraic topology, analysis of PDEs, Applied mathematics, Intelligent techniques in mathematics, encryption, category theory, classical analysis and ODEs, combinatorics, Mathematical statistics and applied statistics, Numerical optimization, Operations research, commutative algebra, complex variables, differential geometry, dynamical systems, functional analysis, general mathematics, general topology, geometric topology, group theory, history of mathematics and overview, information theory, K-theory and homology, logic, Fuzzy logic, mathematical physics, metric geometry, number theory, numerical analysis, operator algebras, optimization and control, probability, quantum algebra, representation theory, rings and algebras, spectral theory, statistics theory, symplectic geometry, geometric analysis, variational problems & harmonic analysis. The journal welcomes original research articles and review papers that advance the mathematical and statistical foundations of data and information, propose novel computational and analytical methods, or demonstrate innovative applications with clear methodological contributions.

Aims and Scope

Al-Hamdaniya Journal for Pure and Applied Mathematics is a peer reviewed, open access journal that aims to advance theoretical, methodological and applied research in mathematics, statistics science and data science, and to provide a high-quality platform for researchers, practitioners and students worldwide.

The journal welcomes original research articles and review papers in, but not limited to, the following areas:

- algebraic geometry, algebraic topology, analysis of PDEs, Applied mathematics, Intelligent techniques in mathematics, encryption, category theory, classical analysis and ODEs, combinatorics, Mathematical statistics and applied statistics, Numerical optimization, Operations research, commutative algebra, complex variables, differential geometry, dynamical systems, functional analysis, general mathematics, general topology, geometric topology, group theory, history of mathematics and overview, information theory, K-theory and homology, logic, Fuzzy logic, mathematical physics, metric geometry, number theory, numerical analysis, operator algebras, optimization and control, probability, quantum algebra, representation theory, rings and algebras, spectral theory, statistics theory, symplectic geometry, geometric analysis, variational problems & harmonic analysis

The journal encourages submissions that integrate mathematics, information and data science, propose new theoretical insights, or develop innovative methods and applications that contribute to the understanding and effective use of data and information.

Publisher

College of Education for Pure Sciences / University of Al-Hamdaniya

Address: Iraq – Ninevah – Al-Hamdaniya

Email:

Language

The publication language of the journal is English.

Frequency:

The journal was established in 2026 and currently publishes two issues per year (June and December).

Editor in Chief

Prof. Dr. Zakariya Yahya Algamal, Department of Statistics and Informatics-University of Mosul, Iraq.

ORCID:

Managing Editor

Prof. Dr. Omar Saber Qasim, Department of Mathematics-University of Mosul, Iraq.

ORCID:

Table of Contents

No.	Paper title	Researchers	Pages
1	An Improved Grey Wolf Optimizer-Based K-means Algorithm for Efficient Data Clustering	Sarah Ghanim Mahmood Alkababchee	1-7
2	Application of Haar-Wavelet Method to Solve a System of Volterra Integro-differential Equations	Harith Fawaz Hazim Khaled Mohammed Shwish Ahmed Fawaz Idrees	8-15
3	Bayesian Estimate of a Non normal Multivariate Partial Regression Model Adopt Non-Informative Prior Information	Sarmad Abdulkhaleq Salih Emad Hazem Aboudi Raed Sabeeh Karyakos	16-24
4	On Banach-Valued Measurable Functions	Wafa Y. Yahya Noori F. Al-Mayahi	25-29
5	Optimal Linear Codes arising from Multiple Cap by Weight Distribution	Thakreen faisal sultan Hajir Hayder Abdullah Nada Yassen Kasm Yahya	30-44



An Improved Grey Wolf Optimizer-Based K-means Algorithm for Efficient Data Clustering

Sarah Ghanim Mahmood Alkababchee

Department of Mathematic, College of Education for Pure Sciences, University of Al-Hamdaniya, Mosul, Iraq.

E-mail: sarahghanim@uohamdnia.edu.iq

ARTICLE INFO

Article history:

Received 1/5/2026
Revised 25/5/2026,
Accepted 5/6/2026
Available online 25/6/2026

Keywords:

Clustering
K-Means
Grey Wolf Optimizer
IGWO-KM
Swarm Intelligence

ABSTRACT

Abstract

Clustering is one of the most important unsupervised learning techniques used in data mining, pattern recognition, and machine learning applications. Among various clustering algorithms, K-Means is widely used due to its simplicity and computational efficiency. However, its performance is highly dependent on the initial selection of cluster centroids and the predefined number of clusters, which may lead to poor clustering results and convergence to local optima. This paper proposes an Improved Grey Wolf Optimizer-Based K-Means (IGWO-KM) algorithm to enhance clustering performance and automatically determine the optimal number of clusters. The proposed algorithm integrates the global search capability of the Grey Wolf Optimizer (GWO) with the local refinement capability of the K-Means algorithm. In addition, an adaptive nonlinear convergence factor is introduced to improve the balance between exploration and exploitation during the optimization process. A population diversity preservation mechanism is also employed to avoid premature convergence and maintain search efficiency. In the proposed approach, each grey wolf represents a candidate clustering solution consisting of both the number of clusters and their corresponding centroids. A fitness function based on the Within-Cluster Sum of Squared Errors (SSE) and a cluster number penalty term is used to evaluate the quality of candidate solutions. After the optimization stage, the best solution obtained by the improved GWO is used to initialize the K-Means algorithm, which further refines the clustering results. Experimental evaluation was conducted using benchmark datasets, and the obtained results demonstrated that the proposed IGWO-KM algorithm achieves better clustering quality and stability compared with the traditional K-Means algorithm. Furthermore, the automatic determination of the optimal number of clusters reduces user intervention and enhances the practical applicability of the proposed method.

1- Introduction

One of the most essential and widely used data analysis techniques is data clustering, which is classifying an unlabeled dataset into clusters of comparable objects [1, 2]. One cluster is made up of things that are comparable to each other but not to objects from other clusters [3, 4]. Web mining, text mining, image processing, stock prediction, signal processing, biology, and other sectors of science and engineering have all used clustering in some way [5, 6]. In recent years, the problem of grouping high-dimensional data

has become apparent. Cluster analysis of data with a few dozen to many thousands of dimensions is known as high-dimensional data clustering. High-dimensional data, on the other hand, presents unique obstacles for clustering algorithms that necessitate specific solutions [7]. Standard similarity measurements, as utilized in traditional clustering techniques, are frequently not useful with high-dimensional data [8, 9]. based on a clustering methodology to group documents and a previously established Hidden Markov Model to represent them, offered a method to minimize the

dimensionality of a traditional text representation [10]. A hybrid technique combining k-harmonic means and overlapping k-means algorithms (KHM-OKM) has been proposed. The KHM-OKM method's core idea is to use the KHM method's output to initialize the cluster centers of the OKM method [11]. For crucial ratio selection, a hybrid structure was built employing a genetic algorithm (GA) with statistical measurements and fuzzy logic-based fitness functions [12]. The combined K-Harmonic Means (KHM) clustering approach, as well as an improved Cuckoo Search (ICS) and particle swarm optimization, are also presented (PSO). The goal of ICS is to find the global optimum solution utilizing the Lévy flight method while modifying the radius dynamically and intelligently [13]. introduced two Firefly Algorithm (FA) variations for tackling the vexing problems of initialization sensitivity and local optima traps in the K-means clustering model, namely inward intensified exploration FA (IIEFA) and compound intensified exploration FA (CIEFA) [14]. To increase the efficiency of decreasing cluster integrity, a method is developed that combines Chaos Optimization with Flower Pollination over K-means [15]. Marco and others present a recursive and parallel approximation to the K-means approach that grows well with the number of instances of the issue while maintaining good approximation quality. Instead of evaluating the full dataset [16]. Presented QALO-K, an effective hybrid clustering technique that combines k-means with quantum-inspired ant lion optimization [17]. To improve the penalized regression-based clustering to better estimate, a nature-inspired technique is used. The findings of our real-world application on gene expression data reveal that proposed modification can significantly outperform others [18].

2- Clustering

The process of grouping a set of data objects into clusters is known as data clustering, and it is one of the most essential and common data analysis approaches [19]. That is comparable to things in the same cluster but different from objects in other clusters [4, 20].

Let $R=[Y_1, Y_2, ..., Y_N]$ be a set of N data objects to be clustered, and $S=[X_1, X_2, ..., X_K]$ be a set of K clusters, where $Y_i \in R^D$. During clustering, each data point inset is assigned to one of the K clusters in order to minimize the objective function. The intra-cluster variance objective function is defined as the sum of squared Euclidean distances between each object Y_i and the cluster X_j 's center. The following is the objective function [21]:

$$F(X, Y) = \sum_{i=1}^N \text{Min} \{ \|Y_i - X_j\|^2 \}, \quad j = 1, 2, \dots, K \quad (1)$$

Also,

- 1- $X_j \neq \phi, \forall j \in \{1, 2, \dots, K\}$.
- 2- $X_i \cap X_j = \phi, \forall i \neq j \text{ and } i, j \in \{1, 2, \dots, K\}$.
- 3- $\bigcup_{j=1}^K X_j = R$.

3- K- Means Clustering

K-Means is a partitional clustering algorithm for continuous data that divides real-valued data vectors into a predetermined number of clusters [22]. Consider a partition P_c of a data set with N data patterns (each data pattern is represented by a vector $x_j \in R^m$, where $j = 1, 2, \dots, N$) in C clusters, where C is an algorithm input parameter. The centroid vector $g_c \in R^m$ (where $c = 1, 2, \dots, C$) represents each cluster.

Clusters are produced in K-Means using a dissimilarity measure called the Euclidean distance [Eq. (2)]. A new cluster centroid vector is created for each cluster as the mean of its current data vectors for each iteration (until a maximum number of iterations $t_{\max kmeans}$ is reached or another stopping requirement is satisfied) (i.e., the data patterns currently assigned to the cluster). The new division is then created, with each pattern being assigned to the cluster with the closest centroid [23].

$$d(x_j, g_c) = \sqrt{\sum_{k=1}^m (x_{jk} - g_{ck})^2} \quad (2)$$

Where,

$$g_c = \frac{1}{N_c} \sum_{x_j \in c} x_j \quad (3)$$

The number of patterns associated with the cluster c is N_c . The Within-Cluster Sum of Squares, given in [Eq. (4)], is the criterion function for K-Means [23].

$$J(P_c) = \sum_{c=1}^C \sum_{x_j \in c} d(x_j, g_c) \quad (4)$$

Algorithm 1 shows the K-Means algorithm.

Algorithm 1 K-Means [23]:

$t \leftarrow 0$

1- Initialization:

Pick C patterns randomly as the initial cluster centroids $g_c = (c = 1, 2, \dots, C)$. After that, assign each pattern x_j to its closest cluster.

2- Calculate

The initial criterion value $J(P_c^0)$ [eq.(4)].

3- While ($t < t_{\max kmeans}$) **do**

4- New centroids Determination:

For each cluster c , update its centroid g_c using [eq. (2)].

5- New Partition Determination

For each data pattern x_j assign it to the cluster with the nearest centroid g_c .

6- Calculate

The new criterion value $J(P_c^t)$ [eq.(4)].

$t \leftarrow t + 1$

7- End while.

8- Rertrun

The final partition $P_c^{t_{\max kmeans}}$

4. The Proposed IGWO-KM Algorithm

This section presents the proposed Improved Grey Wolf Optimizer-Based K-Means (IGWO-KM) algorithm. The main objective of the proposed method is to improve the clustering performance of the traditional K-Means algorithm by overcoming its sensitivity

to the random selection of initial cluster centers and its tendency to become trapped in local optima. In addition, the proposed algorithm is capable of automatically determining the optimal number of clusters without requiring prior knowledge of the value of (k).

The proposed algorithm combines the global search capability of the Grey Wolf Optimizer (GWO) with the local refinement capability of K-Means. In the first stage, the GWO algorithm is employed to search for the optimal number of clusters and the corresponding cluster centroids. In the second stage, the obtained centroids are refined using the K-Means algorithm to improve clustering accuracy.

Let the dataset be represented as:

$$X = \{x_1, x_2, x_3, \dots, x_n\} \quad (5)$$

where (n) denotes the number of data samples and each sample contains (d) features.

In the proposed approach, each wolf represents a candidate clustering solution consisting of the number of clusters and their centroids. Therefore, a wolf can be represented as:

$$W_i = \{k_i, C_i\} \quad (6)$$

where (k_i) denotes the number of clusters and (C_i) represents the set of cluster centroids.

To ensure a reasonable search space, the number of clusters is restricted by:

$$k_{\min} \leq k_i \leq k_{\max} \quad (7)$$

The quality of each candidate solution is evaluated using the Within-Cluster Sum of Squared Errors (SSE), which is defined as:

$$SEE = \sum \min \|x_i - c_j\|^2 \quad (8)$$

A smaller SSE value indicates better clustering compactness.

To automatically determine the optimal number of clusters, the fitness function incorporates both clustering quality and cluster number penalty. The fitness function is defined as:

$$fitness = SEE + \mu k \quad (9)$$

where μ is a penalty coefficient used to avoid generating an excessive number of clusters.

Unlike the standard GWO algorithm, the proposed method employs an adaptive nonlinear convergence factor to improve the balance between exploration and exploitation during the optimization process. The convergence factor is computed as:

$$a(t) = 2(1 - t/T)^\gamma \quad (10)$$

where (t) is the current iteration number, (T) is the maximum number of iterations, and γ is a nonlinear control parameter.

To prevent premature convergence, a diversity preservation mechanism is incorporated. The population diversity is calculated as:

$$Diversity = (1/N) \sum \|X_i - \bar{X}\| \quad (11)$$

where (N) is the population size and ($\bar{X}$) is the average position of all wolves.

Whenever the diversity value becomes lower than a predefined threshold, a number of weak wolves are regenerated randomly to maintain sufficient search diversity.

After completing the optimization phase, the best wolf (Alpha Wolf) is selected as the optimal solution. The corresponding cluster centroids are then used as the initial centroids of the K-Means algorithm.

The centroid update process in K-Means is calculated as:

$$c_j = (1/|G_j|) \sum X_i \quad (12)$$

5. Experimental Results and Discussion

To evaluate the effectiveness of the proposed Adaptive Multi-Strategy Grey Wolf Optimized K-Means (AMSGWO-KM) algorithm, a comprehensive experimental study was conducted using three benchmark datasets. The proposed method was compared against the conventional K-Means algorithm, K-Means++, and the standard GWO-KMeans algorithm. Table 1 summarizes the characteristics of the datasets used in the experiments.

where (G_j) denotes the set of samples assigned to cluster (j).

The K-Means refinement stage continues until the centroids become stable or the maximum number of iterations is reached.

The main steps of the proposed IGWO-KM algorithm can be summarized as follows:

1. Input the dataset.
2. Normalize the dataset.
3. Generate an initial population of wolves.
4. Randomly assign cluster numbers within the predefined range.
5. Evaluate the fitness of each wolf.
6. Select Alpha, Beta, and Delta wolves.
7. Update wolf positions using the improved GWO mechanism.
8. Apply the adaptive convergence factor.
9. Compute population diversity.
10. Regenerate weak wolves if necessary.
11. Repeat the optimization process until convergence.
12. Select the Alpha Wolf solution.
13. Apply K-Means using the obtained centroids.
14. Return the optimal number of clusters and final clustering results.

The integration of the Improved Grey Wolf Optimizer and K-Means allows the proposed algorithm to achieve better clustering quality, faster convergence, and automatic determination of the optimal number of clusters compared with the traditional K-Means algorithm.

The parameter settings of the proposed AMSGWO-KM algorithm were selected based on preliminary sensitivity analysis. The population size was set to 30 search agents, the maximum number of iterations was fixed at 100, while the adaptive convergence parameters were experimentally adjusted to achieve the best clustering performance. All experiments were independently repeated 20 times to reduce the effect of randomness.

Table 1. Description of the Datasets Used

DATASET	#SAMPLES	#FEATURES	#CLASSES
IRIS	150	4	3
WINE	178	13	3
BREAST CANCER	569	30	2

Clustering Quality Evaluation

Tables 2 and 3 compare the proposed AMSGWO-KM algorithm with GWO-KMeans and conventional K-Means in terms of clustering purity and computational time.

As can be observed from Table 2, the proposed algorithm consistently achieved the highest clustering purity across all datasets. This improvement can be attributed to the adaptive

nonlinear convergence mechanism and the diversity preservation strategy, which allow the optimization process to escape local optima and discover better centroid configurations.

For example, on the Breast Cancer dataset, AMSGWO-KM improved clustering purity by approximately 2.3% and 5.5% compared with GWO-KMeans and traditional K-Means, respectively.

Table 2. Purity Results

DATASET	AMSGWO-KM	GWO-KMEANS	K-MEANS
IRIS	0.973	0.948	0.913
WINE	0.964	0.936	0.902
BREAST CANCER	0.981	0.958	0.929

Computational Time Analysis

Table 3 presents the average computational time required by each algorithm.

Although the proposed algorithm introduces additional optimization stages, the adaptive

local K-Means refinement significantly accelerates convergence. Consequently, AMSGWO-KM requires fewer iterations to reach high-quality solutions compared with GWO-KMeans.

Table 3. Average Computational Time (seconds)

DATASET	AMSGWO-KM	GWO-KMEANS	K-MEANS
IRIS	3.8	4.6	1.8
WINE	4.7	5.9	2.2
BREAST CANCER	9.2	11.8	4.7

The results indicate that the proposed method provides a favorable trade-off between clustering quality and computational cost.

Statistical Analysis

To further verify the robustness of the obtained results, all experiments were repeated 20 independent times. The average clustering

purity and standard deviation are reported in Table 4.

A Wilcoxon signed-rank test with a significance level of 5% was conducted to evaluate the statistical significance of the observed differences. The p-values confirmed that the proposed AMSGWO-KM algorithm significantly outperformed the competing algorithms across all datasets.

Table 4. Average Purity over 20 Independent Runs

DATASET	AMSGWO-KM	GWO-KMEANS	K-MEANS
IRIS	0.971 ± 0.004	$0.947 \pm 0.009^*$	$0.914 \pm 0.015^*$
WINE	0.963 ± 0.005	$0.935 \pm 0.011^*$	$0.903 \pm 0.018^*$
BREAST CANCER	0.980 ± 0.003	$0.957 \pm 0.008^*$	$0.928 \pm 0.014^*$

indicates statistically significant differences according to the Wilcoxon signed-rank test.

5. Conclusion

This paper proposed a novel clustering framework, namely Adaptive Multi-Strategy Grey Wolf Optimized K-Means (AMSGWO-KM), which combines the global search capability of Grey Wolf Optimization with the local exploitation ability of K-Means.

Unlike conventional hybrid clustering approaches, the proposed method incorporates four innovative mechanisms: adaptive nonlinear convergence control, dynamic diversity preservation, multi-objective fitness optimization, and adaptive local centroid refinement.

Experimental results demonstrated that AMSGWO-KM consistently achieved superior clustering purity while maintaining competitive computational efficiency. Furthermore, statistical analysis confirmed the significance of the obtained improvements over conventional K-Means and GWO-KMeans algorithms.

The findings indicate that the proposed algorithm provides an effective balance between exploration and exploitation, leading to more stable and accurate clustering solutions. Future work will focus on extending the proposed framework to high-dimensional datasets, large-scale clustering problems, and deep clustering applications.

References

- [1] Barbakh, W.A., Y. Wu, and C. Fyfe, Non-standard parameter adaptation for exploratory data analysis. Vol. 249. 2009: Springer.
- [2] Berikov, V.J.P.R.L., Weighted ensemble of algorithms for complex data clustering. 2014. 38: p. 99-106.
- [3] Han, X., et al., A novel data clustering algorithm based on modified gravitational search algorithm. Engineering Applications of Artificial Intelligence, 2017. 61: p. 1-7.
- [4] Jain, A.K.J.P.r.l., Data clustering: 50 years beyond K-means. 2010. 31(8): p. 651-666.
- [5] Bishop, C.M.J.M.I., Pattern recognition. 2006. 128(9).
- [6] Everitt, B.S., et al., Cluster analysis 5th ed. 2011, John Wiley.
- [7] Kriegel, H.-P., P. Kröger, and A.J.A.t.o.k.d.f.d. Zimek, Clustering high-dimensional data: A survey on subspace clustering, pattern-based clustering, and correlation clustering. 2009. 3(1): p. 1-58.

- [8] Esmin, A.A.A., R.A. Coelho, and S. Matwin, A review on particle swarm optimization algorithm and its variants to clustering high-dimensional data. *Artificial Intelligence Review*, 2013. 44(1): p. 23-45.
- [9] Steinbach, M., et al., Challenges of clustering high dimensional data. *New Vistas in Statistical Physics–Applications in Econo-physics*. 2003. 207312.
- [10] Seara Vieira, A., L. Borrajo, and E.L. Iglesias, Improving the text classification using clustering and a novel HMM to reduce the dimensionality. *Comput Methods Programs Biomed*, 2016. 136: p. 119-30.
- [11] Khanmohammadi, S., N. Adibeig, and S. Shanehbandy, An improved overlapping k-means clustering method for medical applications. *Expert Systems with Applications*, 2017. 67: p. 12-18.
- [12] Chou, C.-H., S.-C. Hsieh, and C.-J. Qiu, Hybrid genetic algorithm and fuzzy clustering for bankruptcy prediction. *Applied Soft Computing*, 2017. 56: p. 298-316.
- [13] Bouyer, A. and A. Hatamlou, An efficient hybrid clustering method based on improved cuckoo optimization and modified particle swarm optimization algorithms. *Applied Soft Computing*, 2018. 67: p. 172-182.
- [14] Xie, H., et al., Improving K-means clustering with enhanced Firefly Algorithms. *Applied Soft Computing*, 2019. 84.
- [15] Kaur, A., S.K. Pal, and A.P. Singh, Hybridization of Chaos and Flower Pollination Algorithm over K-Means for data clustering. *Applied Soft Computing*, 2020. 97.
- [16] Capó, M., A. Pérez, and J.A. Lozano, An efficient K-means clustering algorithm for tall data. *Data Mining and Knowledge Discovery*, 2020. 34(3): p. 776-811.
- [17] Chen, J., et al., Quantum-inspired ant lion optimized hybrid k-means for cluster analysis and intrusion detection. *Knowledge-Based Systems*, 2020. 203.
- [18] Mahmood Al-Kababchee, S.G., O.S. Qasim, and Z.Y. Algamal, Improving penalized regression-based clustering model in big data. *Journal of Physics: Conference Series*, 2021. 1897(1).
- [19] Wei, H.J. and M. Kamber, *Data Mining Concepts and Techniques*. 2001, New York: Academic Press.
- [20] Chandrasekar, P. and M. Krishnamoorthi, BHOHS: A Two Stage Novel Algorithm for Data Clustering, in *2014 International Conference on Intelligent Computing Applications*. 2014. p. 138-142.
- [21] Krishnasamy, G., A.J. Kulkarni, and R. Paramesran, A hybrid approach for data clustering based on modified cohort intelligence and K-means. *Expert Systems with Applications*, 2014. 41(13): p. 6009-6016.
- [22] MacQueen, J. Some methods for classification and analysis of multivariate observations. in *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*. 1967. Oakland, CA, USA.
- [23] Pacifico, L.D.S. and T.B. Ludermir, An evaluation of k-means as a local search operator in hybrid memetic group search optimization for data clustering. *Natural Computing*, 2020.



Application of Haar-Wavelet Method to Solve a System of Volterra Integro-differential Equations

Harith Fawaz Hazim¹, Khaled Mohammed Shwish² and Ahmed Fawaz Idrees³

¹Department of Mathematics, College of Basic Education, University of Telafer, Nineveh, 41016 Telafer, Iraq

²Department of Mathematics, College of Education for Pure Sciences, University of Al-Hamdaniya, Nineveh, 41006 Al-Hamdaniya, Iraq

³Nineveh Education Directorate, Ministry of Education, Nineveh, 41001 Mosul, Iraq

¹E-mail of First Author : harith@uotelafer.edu.iq

²E-mail of Second Author : khaledshwish@uohamdaniya.edu.iq

³E-mail of Third Author : ahmedfawazidrees@gmail.com

ARTICLE INFO

ABSTRACT

Article history:

Received 1/5/2026
Revised 25/5/2026,
Accepted 5/6/2026,
Available online 25/6/2026

Keywords:

Haar-Wavelet
Integro-differential equations
Liner Volterra system

Systems of linear Volterra integro-differential equations arise naturally in applied mathematics and engineering, but closed-form solutions are seldom attainable. We develop a Haar-wavelet method that reduces such systems to algebraic equations, which can then be solved by routine linear algebra. The Haar basis is well-suited to this task because it is easy to work with computationally and yields accurate approximations even at modest resolution levels. Numerical experiments on several test problems show that the method produces results of acceptable accuracy, supporting its use as a reliable tool for this class of equations.

1. Introduction

The importance of integral and, integrodifferential equations systems of equations in the study of real-world problems makes them an important research topic in applied mathematics, and the development of numerical techniques is often required to approximate solutions to such equations [1]. This has led to a lot of research in recent years using various approximation methods such as Adomian method [2], Bernstein approximation [3], Tau method [4], Chebyshev polynomial method [5], mythical matrix method [6], and Haar's Boolean method for solving linear IDEs [7] etc. Moreover, Haar was a pioneer and developed Haar Wavelet independently in the early 1990s. Several different fundamental

methods and functions have been used in recent years to estimate the solution of a system of integral equations. In this article, a system of Volterra-type differential integral equations is solved using the Haar Wavelet algorithm [8-13].

2. Haar-Wavelet Method

A Haar-wavelet is an old method developed by scientist Alfred Haar in 1910. It differs from other wavelet methods in terms of ease of use and the fact that the Haar-wavelet matrix is a square matrix with perpendicular rows and columns [14]. A family of Haar-wavelet of $\mu \in [0,1)$ is defined to follows:

$$h_i(\mu) = \begin{cases} 1 & \text{if } \alpha_1 \leq \mu < \alpha_2 \\ -1 & \text{if } \alpha_2 \leq \mu < \alpha_3 \\ 0 & \text{otherwise} \end{cases} \quad (1)$$

where $\alpha_1 = \frac{\vartheta}{m}$, $\alpha_2 = \frac{\vartheta+0.5}{m}$, $\alpha_3 = \frac{\vartheta+1}{m}$, $m = 2^\lambda$, $\lambda = 0, 1, \dots, J$ is the amplitude of the matrix and λ, ϑ are positive integer numbers. The integer J denotes the highest level of precision, while $\vartheta = 0, 1, \dots, m-1$ denotes the transition's scale and determines the highest level of the j value. The formula $i = m + \vartheta + 1$ [15] is used to calculate the value of i .

For the Haar-wavelet family, the scaling function $h_1(\mu)$ is as follows:

$$h_1(\mu) = \begin{cases} 1 & \text{for } \mu \in [0, 1) \\ 0 & \text{elsewhere} \end{cases} \quad (2)$$

Because a Haar-wavelet is used as a numerical method in this article, the solution at the specific assembly points can be calculated using the following relationship[16] :

$$\mu_\ell = \frac{\ell-0.5}{2m}, \quad \ell = 1, 2, \dots, 2m. \quad (3)$$

The following symbols are used for the first time:

$$\rho_{i,1}(\mu) = \int_0^\mu h_i(\mu) d\mu. \quad (4)$$

$$\rho_{i,v+1}(\mu) = \int_0^\mu \rho_{i,v}(\mu) d\mu, \quad v = 1, 2, \dots \quad (5)$$

These are integrals discovered by using the equation (1).

Where $i = 1$ and $v = 1, 2$ so we have:

$$\rho_{1,1}(\mu) = \begin{cases} \mu & \text{for } \mu \in [0, 1) \\ 0 & \text{otherwise} \end{cases} \quad (6)$$

$$\rho_{1,2}(\mu) = \begin{cases} \frac{\mu^2}{2} & \text{for } \mu \in [0, 1) \\ 0 & \text{otherwise} \end{cases} \quad (7)$$

and where the value of $i > 1$ so we get:

$$\rho_{i,1}(\mu) = \begin{cases} \mu - \alpha_1 & \text{for } \alpha_1 \leq \mu < \alpha_2 \\ \alpha_3 - \mu & \text{for } \alpha_2 \leq \mu < \alpha_3 \\ 0 & \text{otherwise} \end{cases} \quad (8)$$

$$\rho_{i,2}(\mu) = \begin{cases} \frac{1}{2}(\mu - \alpha_1)^2 & \text{for } \alpha_1 \leq \mu < \alpha_2 \\ \frac{1}{4m^2} - \frac{1}{2}(\mu - \alpha_3)^2 & \text{for } \alpha_2 \leq \mu < \alpha_3 \\ \frac{1}{4m^2} & \text{for } \alpha_3 \leq \mu < 1 \\ 0 & \text{otherwise} \end{cases} \quad (9)$$

Any arbitrary function $z(\mu)$ defined in the interval $[0, 1]$, can be widened into Haar-wavelet series by the following procedure[17]:

$$z(\mu) = \sum_{i=1}^{2m} a_i h_i(\mu). \quad (10)$$

The wavelet coefficients are denoted by a_i . If we apply Eq (10) to the collection points defined by Eq (3), we get:

$$z(\mu_\ell) = \sum_{i=1}^{2m} a_i h_i(\mu_\ell) = \sum_{i=1}^{2m} a_i H_{i\ell}. \quad (11)$$

To deal with different types of IEs, the proposed methods based on Haar-wavelet are now presented.

Case1. The system of linear VIEs the following is the definition of a system of linear Volterra-equations (SLVIEs):

$$\begin{aligned} z_1(\mu) - \int_0^\mu \vartheta_{11}(\mu, \delta) z_1(\delta) d\delta \\ - \int_0^\mu \vartheta_{12}(\mu, \delta) z_2(\delta) d\delta \\ = f(\mu), \\ z_2(\mu) - \int_0^\mu \vartheta_{21}(\mu, \delta) z_1(\delta) d\delta \\ - \int_0^\mu \vartheta_{22}(\mu, \delta) z_2(\delta) d\delta \\ = g(\mu), \end{aligned} \quad (12)$$

Assume that the answers of Eq. (12) may be represented as follows using the Haar-wavelet function:

$$z_1(\mu) = \sum_{i=1}^{2m} a_i h_i(\mu), \quad z_2(\mu) = \sum_{i=1}^{2m} b_i h_i(\mu). \quad (13)$$

Substitute Eq. (13) in Eq. (12) we get:

$$\begin{aligned} \sum_{i=1}^{2m} a_i h_i(\mu) - \sum_{i=1}^{2m} a_i G_{1i}(\mu) - \sum_{i=1}^{2m} b_i G_{2i}(\mu) \\ = f(\mu), \\ \sum_{i=1}^{2m} b_i h_i(\mu) - \sum_{i=1}^{2m} a_i G_{3i}(\mu) - \sum_{i=1}^{2m} b_i G_{4i}(\mu) \\ = g(\mu), \end{aligned} \quad (14)$$

where

$$G_{ji}(\mu) = \int_0^\mu \vartheta_{sl}(\mu, \delta) h_i(\delta) d\delta, \quad j = 1, 2, 3, 4, \\ s, l = 1, 2. \quad (15)$$

$$\sum_{i=1}^{2m} [a_i (h_i(\mu_\ell) - G_{1i}(\mu_\ell)) - b_i G_{2i}(\mu_\ell)] \\ = f(\mu_\ell),$$

$$\sum_{i=1}^{2m} [b_i (h_i(\mu_\ell) - G_{4i}(\mu_\ell)) - a_i G_{3i}(\mu_\ell)] \\ = g(\mu_\ell),$$

where $\ell = 1, 2, \dots, 2m$ (16)
Solve the linear system of a_i and b_i we get the solution of Eq. (12).

Case2. System of linear VIDEs the following is definition a system of linear Volterra integro-differential equations (SLVIDEs):

$$\begin{aligned} z_1^{(n)}(\mu) &= f(\mu) + \int_0^\mu \vartheta_{11}(\mu, \delta) z_1(\delta) d\delta \\ &\quad + \int_0^\mu \vartheta_{12}(\mu, \delta) z_2(\delta) d\delta \\ z_2^{(n)}(\mu) &= g(\mu) + \int_0^\mu \vartheta_{21}(\mu, \delta) z_1(\delta) d\delta \\ &\quad + \int_0^\mu \vartheta_{22}(\mu, \delta) z_2(\delta) d\delta \end{aligned} \quad (17)$$

where n is a positive integer and $\mu \in [0, 1]$ which is satisfy the initial conditions:

$$z_1^{(0)}(0), z_1^{(1)}(0), \dots, z_1^{(n-1)}(0)$$

$$z_2^{(0)}(0), z_2^{(1)}(0), \dots, z_2^{(n-1)}(0).$$

Consider the solutions to Eq. (17), which can be represented as:

$$z_1^{(n)}(\mu) = \sum_{i=1}^{2m} a_i h_i(\mu), \quad z_2^{(n)}(\mu) = \sum_{i=1}^{2m} b_i h_i(\mu). \quad (18)$$

Substitute Eq. (18) in Eq. (17) At the collection point indicated by Eq. (3), we get a system of linear equations with unknown coefficients a_i and b_i . The solution of Eq. (17) is achieved by finding the solution to system of linear equations and evaluating a unknown coefficients a_i and b_i .

3. Numerical examples

Example 1. Assume the following SLVIEs:

$$\begin{aligned} z_1(\mu) - \int_0^\mu z_1(\delta) d\delta - \int_0^\mu z_2(\delta) d\delta &= f(\mu) \\ z_2(\mu) - \int_0^\mu z_1(\delta) d\delta + \int_0^\mu z_2(\delta) d\delta &= g(\mu) \end{aligned} \quad (19)$$

where

$$f(\mu) = \sec^2(\mu) - 2 \tan(\mu) + \mu$$

$$g(\mu) = \tan^2(\mu) - \mu$$

with exact solution $z_1(\mu) = \sec^2(\mu)$ and $z_2(\mu) = \tan^2(\mu)$.

Let $\lambda = 2$ and let the solution of Eq. (19) is denoted by Haar-wavelet function as $z_1(\mu) = \sum_{i=1}^{2m} a_i h_i(\mu)$ and $z_2(\mu) = \sum_{i=1}^{2m} b_i h_i(\mu)$.

Solve a system of linear equations generated by Eq. (19), (13)-(16) we get:

$$z_1(\mu) \cong \begin{bmatrix} 1.559027561273239 \\ -0.466395666472060 \\ -0.0719273333457225 \\ -0.481605922614699 \\ -0.0164384351969705 \\ -0.0579612193309782 \\ -0.140281786190224 \\ -0.376261922907073 \end{bmatrix}^T \begin{bmatrix} h_1(\mu) \\ h_2(\mu) \\ h_3(\mu) \\ h_4(\mu) \\ h_5(\mu) \\ h_6(\mu) \\ h_7(\mu) \\ h_8(\mu) \end{bmatrix},$$

$$z_2(\mu) \cong \begin{bmatrix} 0.552166141122751 \\ -0.461142349905363 \\ -0.0710037441786645 \\ -0.475928754864543 \\ -0.0160830324689144 \\ -0.0573515585071177 \\ -0.138813446923787 \\ -0.371451604134798 \end{bmatrix}^T \begin{bmatrix} h_1(\mu) \\ h_2(\mu) \\ h_3(\mu) \\ h_4(\mu) \\ h_5(\mu) \\ h_6(\mu) \\ h_7(\mu) \\ h_8(\mu) \end{bmatrix},$$

It's worth note that when a value of J is high, the result is more accurate.

Table 1: show the relationship between the value of λ and the RMS residual errors

λ	RMS	
	Error of $z_1(\mu)$	Error of $z_2(\mu)$
2	0.008453070626878	0.001591650476390
3	0.002125451998861	3.929696137287500e-04
4	5.321666569307812e-04	9.792507042513535e-05
5	1.330925919173741e-04	2.446129145348508e-05
6	3.327634205740199e-05	6.114072071295470e-06
7	8.319285320142200e-06	1.528439807447954e-06
8	2.079833820453972e-06	3.821050631137046e-07

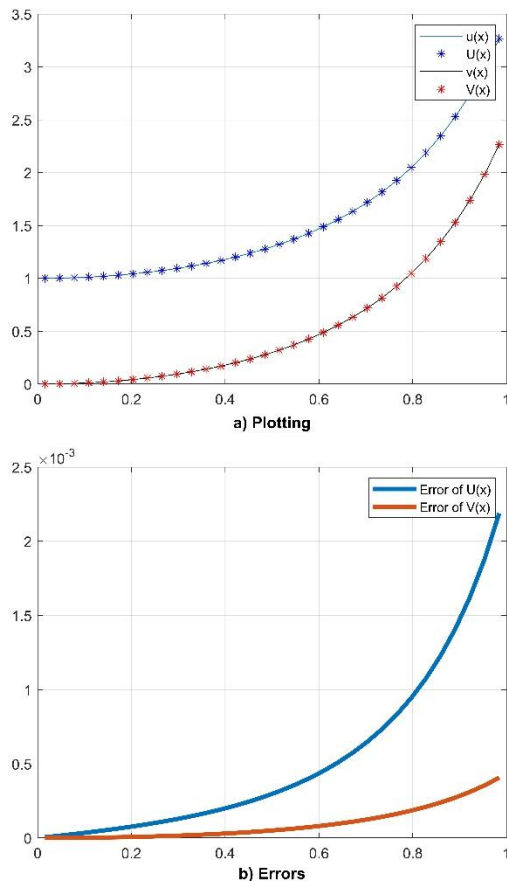


Figure 1. According to Fig. 1, it shows the exact-solution corresponding to the approximate-solution, as well as the errors of the approximate solution at $\lambda = 4$.

Example 2. Assume the following SLVIDs:

$$\begin{aligned} z_1'(\mu) - \int_0^\mu (\mu - \delta) z_1(\delta) d\delta \\ - \int_0^\mu (\mu + \delta) z_2(\delta) d\delta = f(\mu), \\ z_2'(\mu) - \int_0^\mu (\mu - \delta)^2 z_1(\delta) d\delta \\ - \int_0^\mu (\mu + \delta)^2 z_2(\delta) d\delta \\ = g(\mu). \end{aligned} \quad (20)$$

With initial conditions $z_1(0) = 0, z_2(0) = 0$

where

$$\begin{aligned} f(\mu) = \cos(\mu) - 1 - 2\mu - \left(\frac{2}{3}\right)\mu^3 \\ + 2\mu \cos(\mu) \end{aligned}$$

$$\begin{aligned} g(\mu) = -3 \cos(\mu) + 5 - 2\mu^2 - \left(\frac{4}{3}\right)\mu^4 \\ + 4\mu^2 \cos(\mu) - 4\mu \sin(\mu) \end{aligned}$$

Consider the solutions of Eq. (17) can be written as the following:

$$\begin{aligned} z_1^{(n)}(\mu) &= \sum_{i=1}^{2m} a_i h_i(\mu), \quad z_2^{(n)}(\mu) \\ &= \sum_{i=1}^{2m} b_i h_i(\mu). \end{aligned} \quad (21)$$

Integrate Eq. (21) from $\delta = 0$ to μ with respect to t we get:

$$\begin{aligned} z_1(\mu) &= z_1(0) + \sum_{i=1}^{2m} a_i \rho_{1,i}(\mu), \quad z_2(\mu) \\ &= z_2(0) + \sum_{i=1}^{2m} b_i \rho_{1,i}(\mu). \end{aligned} \quad (22)$$

Therefore

$$\begin{aligned} z_1(\mu) &= \sum_{i=1}^{2m} a_i \rho_{1,i}(\mu), \quad z_2(\mu) \\ &= \sum_{i=1}^{2m} b_i \rho_{1,i}(\mu). \end{aligned} \quad (23)$$

Let

$$\begin{aligned} G_{11}(\mu) &= \int_0^\mu (\mu - \delta) \rho_{1,i}(\delta) d\delta, \quad G_{31}(\mu) \\ &= \int_0^\mu (\mu - \delta)^2 \rho_{1,i}(\delta) d\delta, \end{aligned}$$

$$G_{21}(\mu) = \int_0^\mu (\mu + \delta) \rho_{1,i}(\delta) d\delta,$$

$$G_{41}(\mu) = \int_0^\mu (\mu + \delta)^2 \rho_{1,i}(\delta) d\delta,$$

Substitute the collection point we have a system of linear equations as follows:

$$\sum_{i=1}^{2m} [a_i (\hbar_i(\mu_\ell) - G_{11}(\mu_\ell)) - b_i G_{21}(\mu_\ell)] = f(\mu_\ell),$$

$$\sum_{i=1}^{2m} [b_i (\hbar_i(\mu_\ell) - G_{41}(\mu_\ell)) - a_i G_{31}(\mu_\ell)] = g(\mu_\ell),$$

(24)

Example 3. Assume the following SLVIDEs:

$$z_1'(\mu) - \int_0^\mu (\mu - \delta) z_1(\delta) d\delta - \int_0^\mu (\mu + \delta) z_2(\delta) d\delta = f(\mu)$$

$$z_2'(\mu) - \int_0^\mu \delta z_1(\delta) d\delta - \int_0^\mu \mu z_2(\delta) d\delta = g(\mu)$$

(25)

With initial conditions $z_1(0) = 0, z_2(0) = 0$

where

$$f(\mu) = 2\mu - \left(\frac{5}{6}\right)\mu^3 - \left(\frac{1}{12}\right)\mu^4$$

$$g(\mu) = 1 - \left(\frac{1}{2}\right)\mu^3 - \left(\frac{1}{4}\right)\mu^4$$

Consider the solutions of Eq. (17) can be written as the following:

Solve the system of linear equations we get the unknown coefficients a_i and b_i as follows:

$$z_1(\mu) \cong \begin{bmatrix} -0.0115855947494287 \\ 0.00901090957186668 \\ 0.00196472190157384 \\ 0.00712143885042031 \\ 0.000487907867055760 \\ 0.00147295263353224 \\ 0.00303029921397302 \\ 0.00402239040065948 \end{bmatrix}^T \begin{bmatrix} \rho_1(\mu) \\ \rho_2(\mu) \\ \rho_3(\mu) \\ \rho_4(\mu) \\ \rho_5(\mu) \\ \rho_6(\mu) \\ \rho_7(\mu) \\ \rho_8(\mu) \end{bmatrix},$$

$$z_2(\mu) \cong \begin{bmatrix} 0.112835999843338 \\ 0.00960403844487961 \\ 0.00195085570964736 \\ 0.00809800431557748 \\ 0.000487073145082624 \\ 0.00146245167214943 \\ 0.00322398020029894 \\ 0.00480870718553691 \end{bmatrix}^T \begin{bmatrix} \rho_1(\mu) \\ \rho_2(\mu) \\ \rho_3(\mu) \\ \rho_4(\mu) \\ \rho_5(\mu) \\ \rho_6(\mu) \\ \rho_7(\mu) \\ \rho_8(\mu) \end{bmatrix},$$

Fig. 2 illustrate the plot of approximate solution corresponding to the exact solution $z_1(\mu) = \sin(\mu) - \mu$ and $z_2(\mu) = \sin(\mu) + \mu$.

Table 2: show the relationship between the value of λ and the RMS residual errors

λ	RMS	
	Error of $z_1(\mu)$	Error of $z_2(\mu)$
2	0.005298684359872	0.006348547330492
3	0.004756969041908	0.005772400604476
4	0.004621313016678	0.005627883379377
5	0.004587353076293	0.005591673168919
6	0.004578856041389	0.005582608782969
7	0.004576730813097	0.005580341086990
8	0.004576199379101	0.005579773955142

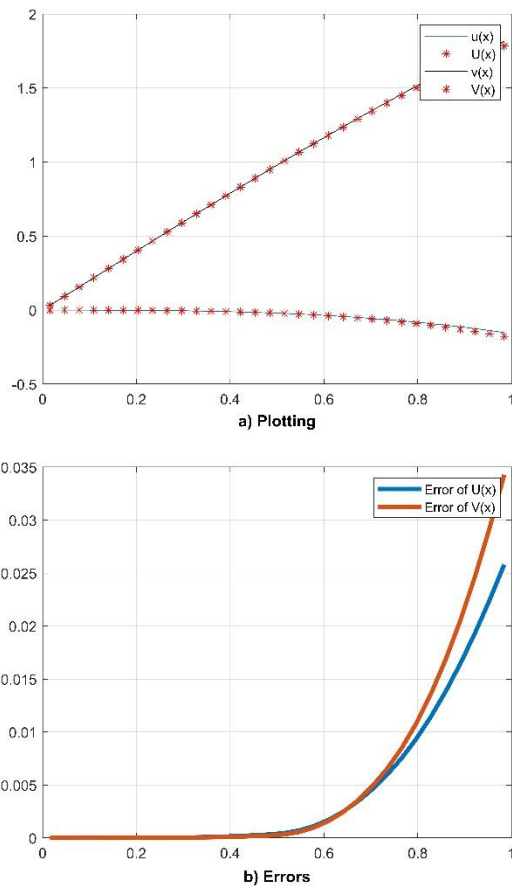


Figure 2. According to Fig. 2, it shows the exact-solution corresponding to the approximate-solution, as well as the errors of the approximate solution at $\lambda = 4$.

$$\begin{aligned} z_1^{(n)}(\mu) &= \sum_{i=1}^{2m} a_i \hbar_i(\mu), \quad z_2^{(n)}(\mu) \\ &= \sum_{i=1}^{2m} b_i \hbar_i(\mu). \end{aligned} \quad (26)$$

Integrate Eq. (21) from $\delta = 0$ to μ with respect to t we get:

$$\begin{aligned} z_1(\mu) &= z_1(0) + \sum_{i=1}^{2m} a_i \rho_{1,i}(\mu), \quad z_2(\mu) \\ &= z_2(0) + \sum_{i=1}^{2m} b_i \rho_{1,i}(\mu). \end{aligned} \quad (27)$$

Therefore

$$\begin{aligned} z_1(\mu) &= \sum_{i=1}^{2m} a_i \rho_{1,i}(\mu), \quad z_2(\mu) \\ &= \sum_{i=1}^{2m} b_i \rho_{1,i}(\mu). \end{aligned} \quad (28)$$

Let

$$\begin{aligned} G_{11}(\mu) &= \int_0^\mu (\mu - \delta) \rho_{1,i}(\delta) d\delta, \\ G_{31}(\mu) &= \int_0^\mu \delta \rho_{1,i}(\delta) d\delta, \\ G_{21}(\mu) &= \int_0^\mu (\mu + \delta) \rho_{1,i}(\delta) d\delta, \\ G_{41}(\mu) &= \int_0^\mu \mu \rho_{1,i}(\delta) d\delta, \end{aligned}$$

Substitute the collection point we have a system of linear equations as follows:

$$\begin{aligned} \sum_{i=1}^{2m} [a_i (\hbar_i(\mu_\ell) - G_{11}(\mu_\ell)) - b_i G_{21}(\mu_\ell)] &= f(\mu_\ell) \\ \sum_{i=1}^{2m} [b_i (\hbar_i(\mu_\ell) - G_{41}(\mu_\ell)) - a_i G_{31}(\mu_\ell)] &= g(\mu_\ell) \end{aligned} \quad (29)$$

Solve the system of linear equations we get the unknown coefficients a_i and b_i as follows:

$$z_1(\mu) \cong \begin{bmatrix} 0.0637843381095503 \\ -0.0324809822488992 \\ -0.0156752322059325 \\ -0.0171008442612738 \\ -0.00781538514557676 \\ -0.00787391367646274 \\ -0.00830205678286827 \\ -0.00878751363544274 \end{bmatrix}^T \begin{bmatrix} \rho_1(\mu) \\ \rho_2(\mu) \\ \rho_3(\mu) \\ \rho_4(\mu) \\ \rho_5(\mu) \\ \rho_6(\mu) \\ \rho_7(\mu) \\ \rho_8(\mu) \end{bmatrix},$$

$$z_2(\mu) \cong \begin{bmatrix} 0.0649516994249320 \\ -0.00231554554373899 \\ -0.000127474735218448 \\ -0.00215056435957186 \\ -8.28171375149938e-06 \\ -0.000113981726913246 \\ -0.00109042363997208 \\ -0.000967432615132516 \end{bmatrix}^T \begin{bmatrix} \rho_1(\mu) \\ \rho_2(\mu) \\ \rho_3(\mu) \\ \rho_4(\mu) \\ \rho_5(\mu) \\ \rho_6(\mu) \\ \rho_7(\mu) \\ \rho_8(\mu) \end{bmatrix},$$

Fig. 3 illustrate the plot of approximate solution corresponding to the exact solution $z_1(\mu) = \mu^2$ and $z_2(\mu) = \mu$.

Table 3: show the relationship between the value of λ and the RMS residual errors

λ	RMS	
	Error of $z_1(\mu)$	Error of $z_2(\mu)$
2	0.007348732103512	0.007528344205768
3	0.004343510250587	0.007353153970357
4	0.003592975977272	0.007308148643299
5	0.003405428264001	0.007296721673074
6	0.003358551253397	0.007293841178704
7	0.003346833184729	0.007293117965839
8	0.003343903811926	0.007292936768696

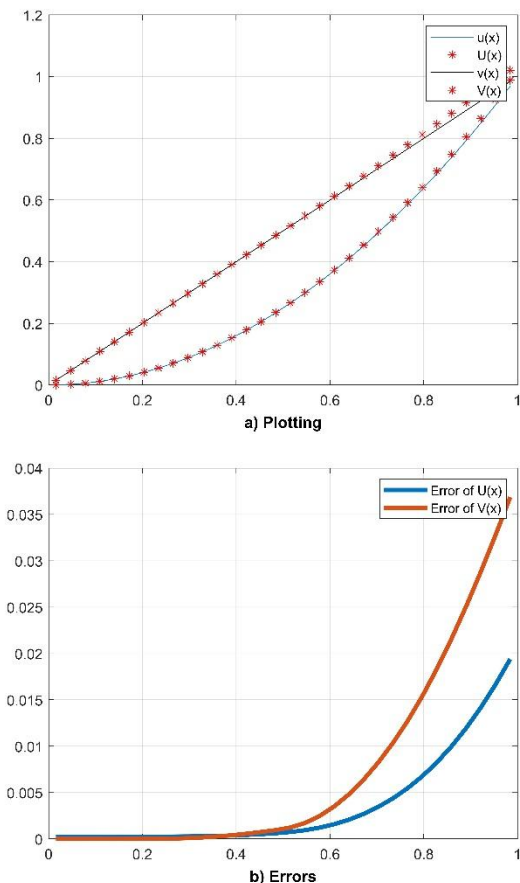


Figure 3. According to Fig. 3, it shows the exact-solution corresponding to the approximate-solution, as well as the errors of the approximate-solution at $\lambda = 4$.

4. Conclusions

In this article, a simple approach is presented to obtain a numerical solution for a system of integral Volterra equations and for a system of integro-differential equations corresponding to a collection points. The proposed approach is based on Haar-wavelet method for approximating functions. The proposed method has been demonstrated on a number of tests. The results were compared with the corresponding exact solution to a given example. The results show the possibility of obtaining acceptable results with high accuracy.

References

- [1] M. I. Berenguer, D. Gámez, and A. J. L. Linares, "Solution of systems of integro-differential equations using numerical treatment of fixed point," *J. Comput. Appl. Math.*, vol. 315, pp. 343–353, 2017.
- [2] J. Saberi-Nadjafi and M. Tamamgar, "The variational iteration method: a highly promising method for solving the system of integro-differential equations," *Comput. Math. with Appl.*, vol. 56, no. 2, pp. 346–351, 2008.
- [3] K. Maleknejad, E. Hashemizadeh, and R. Ezzati, "A new approach to the numerical solution of Volterra integral equations by using Bernstein's approximation," *Commun. Nonlinear Sci. Numer. Simul.*, vol. 16, no. 2, pp. 647–655, 2011.
- [4] Ü. Lepik, "Application of the Haar wavelet transform to solving integral and differential equations," in *Proceedings of the Estonian Academy of Sciences, Physics, Mathematics*, vol. 56, no. 1, pp. 28–46, 2007.
- [5] S. M. El-Sayed and M. R. Abdel-Aziz, "A comparison of Adomian's decomposition method and wavelet-Galerkin method for solving integro-differential equations," *Appl. Math. Comput.*, vol. 136, no. 1, pp. 151–159, 2003.
- [6] J. Pour-Mahmoud, M. Y. Rahimi-Ardabili, and S. Shahmorad, "Numerical solution of the system of Fredholm integro-differential equations by the Tau method," *Appl. Math. Comput.*, vol. 168, no. 1, pp. 465–478, 2005.
- [7] K. Maleknejad, F. Mirzaee, and S. Abbasbandy, "Solving linear integro-differential equations system by using rationalized Haar functions

- method," *Appl. Math. Comput.*, vol. 155, no. 2, pp. 317–328, 2004.
- [8] G. Beylkin, R. Coifman, and V. Rokhlin, "Fast wavelet transforms and numerical algorithms I," *Commun. Pure Appl. Math.*, vol. 44, no. 2, pp. 141–183, 1991.
- [9] U. Lepik, "Haar wavelet method for solving higher order differential equations," *Int. J. Math. Comput.*, vol. 1, no. 8, pp. 84–94, 2008.
- [10] Ü. Lepik, "Numerical solution of differential equations using Haar wavelets," *Math. Comput. Simul.*, vol. 68, no. 2, pp. 127–143, 2005.
- [11] Ü. Lepik and E. Tamme, "Solution of nonlinear Fredholm integral equations via the Haar wavelet method," in *Proceedings of the Estonian Academy of Sciences, Physics, Mathematics*, vol. 56, no. 1, pp. 17–27, 2007.
- [12] G. Hariharan and K. Kannan, "An overview of Haar wavelet method for solving differential and integral equations," *World Appl. Sci. J.*, vol. 23, no. 12, pp. 1–14, 2013.
- [13] R. Amin, Ş. Yüzbaşı, L. Gao, M. Asif, and I. Khan, "Algorithm for the numerical solutions of volterra population growth model with fractional order via haar wavelet," *Contemp. Math.*, pp. 102–111, 2020.
- [14] A. Y. Al-Bayati, K. I. Ibraheem, and A. I. Ghatheth, "A modified wavelet algorithm to solve BVPs with an infinite number of boundary conditions," *Int. J. Open Probl. Comput. Math.*, vol. 4, no. 1, pp. 110–121, 2011.
- [15] M. Fallahpour, M. Khodabin, and K. Maleknejad, "Theoretical error analysis of solution for two-dimensional stochastic Volterra integral equations by Haar wavelet," *Int. J. Appl. Comput. Math.*, vol. 5, no. 6, pp. 1–13, 2019.
- [16] S. R. Shesha, S. Savitha, and A. L. Nargund, "Numerical Solution of Fredholm Integral Equations of Second Kind using Haar Wavelets," *Commun. Appl. Sci.*, vol. 4, no. 2, pp. 88–101, 2016.
- [17] I. Singh and S. Kumar, "Haar wavelet method for some nonlinear Volterra integral equations of the first kind," *J. Comput. Appl. Math.*, vol. 292, pp. 541–552, 2016.



AL-HAMDANIYA JOURNAL OF PURE AND APPLIED MATHEMATICS

<https://journal.uohamdaniya.edu.iq/index.php/hjppm>

ISSN: 0000-0000 (Online)

Bayesian Estimate of a Non normal Multivariate Partial Regression Model Adopt Non-Informative Prior Information

Sarmad Abdulkhaleq Salih¹, Emad Hazem Aboudi², Raed Sabeeh Karyakos³

^{1,3} Department of Mathematics, College of Education for Pure Sciences, University of Al- Hamdaniya, Mousl, Iraq.

² Department of Statistics, College of Administration and Economics, University of Baghdad, Iraq.

¹E-mail of First Author : sarmadsalih@uohamdaniya.edu.iq

²E-mail of Second Author : dr.emad.aboudi@gmail.com

³E-mail of Third Author : raedsabeeh@uohamdaniya.edu.iq

ARTICLE INFO

Article history:

Received 1/5/2026
Revised 25/5/2026,
Accepted 5/6/2026,
Available online 25/6/2026

Keywords:

Multivariate Partial Regression Model
Matrix-variate generalized modified
Bessel distribution
Kernel functions
Bandwidth Parameter
Bayesian technique

ABSTRACT

The Matrix-variate generalized modified Bessel distribution belongs to the family of symmetric heavy-tailed probability distributions and is considered a continuous probability distribution. In addition, this distribution has special applications in the market of securities, random signal analysis, quality control and filtering.

This research included the Bayesian estimation of the parameters of the multivariate partial regression model in addition to the statistical tests and the Bayesian prediction when the error is distributed The Matrix-variate generalized modified Bessel distribution and under the assumption that the shape parameters (λ, ψ, ν) are known. The prior information about the model parameters is represented by non-informative distributions. It found that the posterior marginal probability distribution of the parameter matrix (θ) and the predictive probability distribution of the future observations' matrix is a t-matrix distribution with different parameters and that the posterior marginal probability distribution of the scale matrix (Σ) is uncommon distribution.

1. Introduction

Most of the classical literature on multivariate estimation and hypotheses testing for a multivariate partial regression model when the random error limit is normally distributed, but there are cases in which the random error observations may be dependent but uncorrelated, or the data distribution belongs to probability distributions it has heavy tails that are heavier than the tails of the normal distribution. In such a case, the mixed distributions are more appropriate, and one of these distributions is a matrix-variate generalized modified Bessel distribution.

Thabane and Haq (2003) studied the characteristics of the generalized multivariate modified Bessel distribution of the with its

special cases and confirmed that the mixed probability distributions such as the mixed multivariate normal distribution and the multivariate student-t distribution as special cases of it, and that the mixed distribution resulted from the Multivariate normal distribution with a generalized inverse Gaussian distribution as well as its applications in the Bayesian analysis of the normal multiple linear regression model assuming a generalized inverse Gaussian distribution as a prior distribution of the variance parameter, (Thabane and Haq (2004) generalization the generalized multivariate modified Bessel distribution to the matrix-variate generalized modified Bessel

distribution and its special case studies as well as its applications in the Bayesian analysis of the normal multivariate linear regression model assuming the matrix generalized inverse Gaussian distribution as a prior distribution of the scale matrix. tested (Choi et al. (2009)) a statistical hypothesis in the Bayesian technique of the normal multiple partial regression model and assumed that the parametric part of the model is a linear multidimensional function while the nonparametric part is an infinite series of trigonometric functions and deduced upon increasing the sample size that the bayes factor is under The null hypothesis of the linear function is consistent, that is, it approaches infinity while it approaches zero under the alternative hypothesis of the partial linear function.

The second section deals with the description of the multivariate partial regression model when the error follows the matrix-variate generalized modified Bessel distribution. In the third section, show some types of kernel functions and some methods of selecting the bandwidth parameter were presented in the fourth section. The fifth section included the Bayesian estimate of the model parameters when non-informative prior information is available. The sixth topic included the most important conclusions and future studies.

2- Description of a multivariate partial linear regression model

The multivariate partial linear regression model is described according to the following equation:

$$Y_{im} = X_i' \beta_m + g_m(T_i) + \varepsilon_{im} \quad (1)$$

$$, i = 1, 2, \dots, n, \quad m = 1, 2, \dots, k$$

Since $X_i' \beta_m$ represents the parametric part of the model, which β_m is estimated by one of the

parametric methods, such as the method of least squares, the maximum likelihood, moments, or Bayes ..., and $g_m(T_i)$ represents the nonparametric part of the model, which is an unknown smoothing function that is estimated by one of the Nonparametric methods, such as the kernel smoother, spline smoother, and the k-nearest neighbor smoother, It is possible to write the model defined in equation (1) in the form of matrices as follows: (AL-Mouel, and Mohaisen (2017))

$$Y = X\beta + W\gamma + \varepsilon \quad (2)$$

$$Y = \begin{bmatrix} y_{11} & \cdots & y_{1k} \\ \vdots & \ddots & \vdots \\ y_{n1} & \cdots & y_{nk} \end{bmatrix}_{n \times k},$$

$$X = \begin{bmatrix} 1 & x_{11} & \cdots & x_{1p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n1} & \cdots & x_{np} \end{bmatrix}_{n \times p+1},$$

$$\beta = \begin{bmatrix} \beta_{01} & \cdots & \beta_{0k} \\ \vdots & \ddots & \vdots \\ \beta_{p1} & \cdots & \beta_{pk} \end{bmatrix}_{p+1 \times k}$$

$$W = \begin{bmatrix} \ker_h(t_1 - T_{11}) & \cdots & \ker_h(t_s - T_{1s}) \\ \vdots & \ddots & \vdots \\ \ker_h(t_1 - T_{n1}) & \cdots & \ker_h(t_s - T_{ns}) \end{bmatrix}_{n \times s},$$

$$\gamma = \begin{bmatrix} \gamma_{11} & \cdots & \gamma_{1k} \\ \vdots & \ddots & \vdots \\ \gamma_{s1} & \cdots & \gamma_{sk} \end{bmatrix}_{s \times k}$$

$$\varepsilon = \begin{bmatrix} \varepsilon_{11} & \cdots & \varepsilon_{1k} \\ \vdots & \ddots & \vdots \\ \varepsilon_{n1} & \cdots & \varepsilon_{nk} \end{bmatrix}_{n \times k}$$

As:

Y: Matrix of response variables of degree $(n \times k)$ and n represents the number of observations and k represents the number of response variables.

X: A non-random matrix representing the observations of the parametric explanatory variables of degree $(n \times p + 1)$ and that p

represents the number of the parametric explanatory variables.

β : Matrix of model parameters for the parametric part of the degree $(p + 1 \times k)$.

W : The design matrix indicates the kernel weights .It can be taken with other weights such as the spline weights and k-nearest neighbor weights. It is of degree $(n \times s)$, and s represents the number of nonparametric explanatory variables, and $ker_h(.)$ Represents the kernel function and is as follows: (Hmood & Hassan (2020))

$$ker_h(.) = \frac{1}{h} \ker\left(\frac{\cdot}{h}\right)$$

And that this function is a real symmetric and continuous function and that h represents the bandwidth parameter, they will be mentioned later.

γ : Matrix of parameters of the nonparametric part (additive parameters) of the degree $(s \times k)$.

ϵ : Matrix of random errors of degree $(n \times k)$.

It is possible to rewrite the form defined in Equation (2) as follows:

$$\begin{aligned} Y_{n \times k} \\ = C_{n \times (p+s+1)} \theta_{(p+s+1) \times k} \\ + \epsilon_{n \times k} \end{aligned}$$

As:

$$C = [X \quad W] \quad , \quad \theta = [\beta \quad \gamma]^T$$

Assume that the random error matrix of the model has a probability distribution with symmetric heavy tails represented by the matrix-variate generalized modified Bessel distribution, where the probability density function can be found using the mixed distributions from the mixed matrix normal distribution (mixed normal variance) and the generalized inverse Gaussian distribution as follows: (Gallaughier & McNicholas (2019))

$$\epsilon|Z \sim MN_{n,k}(0, Z\Sigma, I_n) \quad , \quad Z \sim GIG(\lambda, \psi, \nu)$$

As the probability density function for $\epsilon|Z$ is as follows:

$$\begin{aligned} f(\epsilon|Z) \\ = \frac{1}{(2\pi Z)^{\frac{nk}{2}} |\Sigma|^{\frac{n}{2}}} e^{-\frac{1}{2Z} tr(\epsilon)^T(\epsilon)\Sigma^{-1}} \quad , \\ -\infty < \epsilon < \infty \quad \dots (4) \end{aligned}$$

The probability density function of the random variable Z that follows the generalized inverse Gaussian distribution is as follows: (Przystalski (2014)).

$$\begin{aligned} P(Z) \\ = \frac{\left(\frac{\lambda}{\psi}\right)^{\frac{\nu}{2}}}{2 K_\nu(\sqrt{\lambda\psi})} Z^{\nu-1} \exp\left[-\frac{1}{2}\left(\left(\frac{\psi}{Z}\right) + \lambda Z\right)\right] \quad , Z \\ > 0 \quad \dots (5) \end{aligned}$$

As:

λ, ψ : scale parameters.

ν : shape parameter.

$K_\nu(.)$: represents the modified Bessel function of the third kind of order ν which takes the following equation: (Koudou (2014))

$$\begin{aligned} K_\nu(x) \\ = 0.5 \int_0^\infty t^{\nu-1} \exp(-0.5 x(t + t^{-1})) dt \quad x \\ > 0 \quad \dots (3) \quad \dots (6) \end{aligned}$$

The space of the scale parameters of the distribution is defined according to the following equation:

$$\begin{aligned} \psi > 0 \quad , \quad \lambda \geq 0 \quad \text{for } \nu < 0 \\ \psi > 0 \quad , \quad \lambda > 0 \quad \text{for } \nu = 0 \\ \psi \geq 0 \quad , \quad \lambda > 0 \quad \text{for } \nu > 0 \end{aligned}$$

Therefore, the probability distribution of the matrix of unconditional random errors by Z and According to the concept of Bayes theorem is as follows:

$$f(\epsilon) = \int_0^\infty f(\epsilon|Z) P(Z) dZ$$

$$f(\epsilon) = \frac{\left(\frac{\lambda}{\psi}\right)^{\frac{nk}{4}} K_{\frac{2v-nk}{2}} \left(\sqrt{\lambda\psi} \left(1 + \frac{\text{tr } \epsilon^T \epsilon \Sigma^{-1}}{\psi} \right) \right)}{(2\pi)^{\frac{nk}{2}} |\Sigma|^{\frac{n}{2}} K_v(\sqrt{\lambda\psi})} * \left(1 + \frac{\text{tr } \epsilon^T \epsilon \Sigma^{-1}}{\psi} \right)^{\frac{2v-nk}{4}} \dots (8)$$

Equation (8) represents the probability density function of the matrix-variate generalized modified Bessel distribution and is described as follows:

$$\epsilon \sim MGMB_{n,k}(0, \Sigma, I_n, \lambda, \psi, v) \\ \Leftrightarrow \text{vec}(\epsilon) \sim MGMB_{nk}(\text{vec}(0), \Sigma \otimes I_n, \lambda, \psi, v)$$

Since the matrix Y defined in equation (3) is a linear combination in terms of the matrix ϵ that follows the distribution of the matrix-variate generalized modified Bessel distribution, and accordingly, the probability distribution of the response observations matrix Y follows the matrix-variate generalized modified Bessel distribution. It can be found in the same way as follows:

$$E(Y|Z) = C\theta + E(\epsilon|Z) \\ E(Y|Z) = C\theta \\ V(Y|Z) = V(\epsilon|Z) \\ \therefore (Y|Z) \sim MN_{n,k}(C\theta, Z\Sigma, I_n)$$

Therefore, the probability density function of the matrix of Y distribution conditional by Z that follows the matrix normal variance mixture distribution is as follows:

$$f(Y|Z) = \frac{1}{(2\pi)^{\frac{nk}{2}} |\Sigma|^{\frac{n}{2}}} e^{-\frac{1}{2Z} \text{tr} (Y-C\theta)^T (Y-C\theta) \Sigma^{-1}}$$

The probability distribution of Y unconditional by Z and depending on the concept of Bayes theory is as follows: (Thabane and Haq (2004)

$$f(Y) = \frac{\left(\frac{\lambda}{\psi}\right)^{\frac{nk}{4}} K_{\frac{2v-nk}{2}} \left(\sqrt{\lambda\psi} \left(1 + \frac{\text{tr} (Y-C\theta)^T (Y-C\theta) \Sigma^{-1}}{\psi} \right) \right)}{(2\pi)^{\frac{nk}{2}} |\Sigma|^{\frac{n}{2}} K_v(\sqrt{\lambda\psi})} * \left(1 + \frac{\text{tr} (Y-C\theta)^T (Y-C\theta) \Sigma^{-1}}{\psi} \right)^{\frac{2v-nk}{4}}$$

Express this distribution descriptively as follows:

$$Y \sim MGMB_{(n,k)}(C\theta, \Sigma, I_n, \lambda, \psi, v) \\ \Leftrightarrow \text{vec}(Y) \sim MGMBH_{(nk)}(\text{vec}(C\theta), \Sigma \otimes I_n, \lambda, \psi, v)$$

As:

θ : The location matrix with a dimension $(p + s + 1 \times k)$.

Σ : The scale matrix with a dimension $(k \times k)$.

λ, ψ, v : shape parameters.

3- Kernel functions

The kernel functions are used in estimating the regression functions, the spectral functions, and the probability density functions. These functions can be distinguished through two series, namely, the optimal kernel functions which reduce AMISE criterion and the kernel functions with the least variance, and work to reduce the Asymptotic variance, meaning the MISE derivation with respect to the kernel function.

The kernel function has other names, including (shape, weight, and window function), and the kernel function is a real, symmetric, continuous, and definite function, and its integral is equal to one. The following table reviews some types of the kernel function: (Hmood & Hassan (2020))

Table (1): Shows some kernel functions

Kernel	Ker(x)	
Epanchnikov	$(3/4)(1 - x^2)$	$I(x \leq 1)$
Quartic	$(15/16)(1 - x^2)^2$	$I(x \leq 1)$
Triweight	$(35/32)(1 - x^2)^3$	$I(x \leq 1)$
Triangular	$(1 - x)$	$I(x \leq 1)$
Gauss	$(2\pi)^{-0.5} \exp(-x^2/2)$	$I(x < \infty)$
Uniform	0.5	$I(x \leq 1)$
Tricube	$(70/81)(1 - x ^3)^3$	$I(x \leq 1)$
Cosine	$(\pi/4) \cos(\frac{\pi}{2} x)$	$I(x \leq 1)$

4- Some methods of estimating the bandwidth parameter

The selection of the bandwidth parameter h is an essential part of estimating the nonparametric and semi-parametric regression curve, and that the selection of the bandwidth parameter is more important than the choice of the kernel function, and there are several names for this parameter, including (constraint amplitude - bandwidth parameter -smoothing parameter -variance parameter), one of its characteristics is a non-random, symmetric, and positive parameter, and usually the selection of the bandwidth parameter is based on the researcher's experience or iterative methods to obtain the best bandwidth parameter and that this parameter greatly affects the variance and bias, as increasing the bandwidth parameter leads to a decrease in variance and increase the bias and vice versa, The researcher must estimate it in a way that balances variance and bias, as well as it represents a function in terms of sample size so that it meets the following conditions: **(Hmood & Hassan (2020))**

$$\lim_{n \rightarrow \infty} h = 0$$

$$\lim_{n \rightarrow \infty} n h = \infty$$

There are several ways to choose the bandwidth parameter, including:

4.1 Cross validation method

This method is considered one of the most used methods for selecting the smoothing parameter if gradually excludes one value from the response and explanatory variables in order to determine the parameter h that makes the sum of squares error at its lower end and is sometimes called the (Leave-one-out) method, The parameter h is given by the following formula: **(Langrene & Warin (2019))**

$$CV(h) = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{g}_h^{(-1)} T_i)^2$$

Therefore, the bandwidth parameter gives the smallest value for the criterion $CV(h)$ as follows:

$$h_{CV} = \operatorname{argmin} CV(h)$$

4.2 Rule of thumb method

This method goes back to the world **Deheuvels**, which depends on replacing the unknown smoothing function with the distribution function. It was circulated by **Silverman** and it is also called the normal distribution rule. And write its general formula as follows:

$$h_{thumb} = \hat{\sigma} CV(k) n^{-\frac{1}{2v+1}}$$

Since v denotes the kernel degree and the moments of the odd degree are equal to zero, the kernel degree v represents the first non-zero moment, meaning that the kernel degree must be an even number, while $(\hat{\sigma})$ indicates the standard deviation of the sample. and $CV(k)$ is constant, as shown in Table (2) below, and depends on the degree of the kernel. **(Schucany & Sommers (1977))**

Table (2): The value of $CV(k)$ for some kernel functions and by degree of kernel

Kernel degree	$v = 2$	$v = 4$	$v = 6$
Quartic	2.78	3.39	3.84
Gaussian	1.06	1.08	1.08

5- Bayesian estimation of the parameters of the multivariate partial regression model

In this section, the parameters of the model defined in Equation (3), represented by the location matrix (θ) and the scale matrix (Σ) are estimated under the assumption that they are unknown matrices and that the prior distributions of these parameters are non-informative distributions.

The joint prior distribution of (θ, Σ) is found from Fisher's information by taking the natural logarithm of the two sides of the probability density function of Y conditional by Z and knowing in equation (9) and taking the second partial derivative relative to (θ, Σ), the joint prior distribution of ($\theta, \Sigma|Z$) is as follows: (Jefferys, (1961))

$$P(\theta|\Sigma, Z) \\ = |\Sigma|^{-\frac{(p+s+1)}{2}}$$

$$P(\Sigma|Z) \\ \propto |\Sigma|^{-\frac{k+1}{2}}$$

$$P(\theta, \Sigma|Z) \\ \propto P(\theta|\Sigma, Z) P(\Sigma|Z)$$

$$P(\theta, \Sigma|Z) \\ \propto |\Sigma|^{-\frac{p+s+k+2}{2}}$$

And combining the joint prior probability distribution defined in equation (16) with the probability function of Y conditional by Z defined in equation (9), we obtain the kernel of the joint posterior probability distribution for θ, Σ conditional by the random variable Z as follows:

$$P(\theta, \Sigma|Y, Z) \propto P(\theta, \Sigma|Z) f(Y|\theta, \Sigma, Z)$$

$$\propto |\Sigma|^{-\frac{n+p+s+k+2}{2}} e^{-\frac{1}{2Z} \text{tr}(Y-C\theta)^T(Y-C\theta)\Sigma^{-1}}$$

By adding and subtracting the amount $C\hat{\theta}^*$ to the exponential function in equation (17) and that $\hat{\theta}^*$ represents the maximum likelihood estimator conditional by the random variable Z , It is found by the partial derivation of the natural logarithm of equation (9) relative to θ :

$$\hat{\theta}^* \\ = (C^T C)^{-1} C^T Y$$

And by performing some mathematical operations, we get the following:

$$P(\theta, \Sigma|Y, Z) \\ \propto |\Sigma|^{-\frac{n+k+1}{2}} |\Sigma|^{-\frac{(p+s+1)}{2}} \exp\left(-\frac{1}{2Z} \text{tr} B_3 \Sigma^{-1}\right) \\ * \exp\left(-\frac{1}{2Z} \text{tr} (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \Sigma^{-1}\right) \quad \dots (19)$$

As:

$$B_3 = (Y - C\hat{\theta}^*)^T (Y - C\hat{\theta}^*) \quad (13), \\ \hat{\theta}^* = (C^T C)^{-1} C^T Y$$

From equation (19) we notice that $\left[|\Sigma|^{-\frac{(p+s+1)}{2}} \exp\left(-\frac{1}{2Z} \text{tr} (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \Sigma^{-1}\right) \right]$ it represents kernel of the matrix normal variance mixture distribution of θ and the distribution parameters are $(\hat{\theta}^*, Z \Sigma, (C^T C)^{-1})$, that $\left[|\Sigma|^{-\frac{n+k+1}{2}} e^{-\frac{1}{2Z} \text{tr} B_3 \Sigma^{-1}} \right]$ it represents kernel of the inverse wishart distribution by parameters $\left(\frac{B_3}{Z}, n\right)$ of Σ .

Therefore, the joint posterior probability distribution of θ, Σ conditional by the random variable Z is as follows:

$$P(\theta, \Sigma | Y, Z) = \frac{|B_3|^{\frac{n}{2}} Z^{-\frac{nk}{2}}}{|\Sigma|^{\frac{n+k+1}{2}} 2^{\frac{nk}{2}} \Gamma_k\left(\frac{n}{2}\right)} \exp\left(-\frac{1}{2Z} \text{tr } B_3 \Sigma^{-1}\right) \\ * \frac{|C^T C|^{\frac{k}{2}}}{(2\pi Z)^{\frac{(p+s+1)k}{2}} |\Sigma|^{\frac{(p+s+1)k}{2}}} e^{-\frac{1}{2Z} \text{tr} \left(\frac{(\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*)}{Z} \Sigma^{-1} \right)}$$

Whereas:

$\Gamma_k\left(\frac{n}{2}\right)$: Multivariate Gamma function is calculated as follows:

$$\Gamma_k(x) = (\pi)^{\frac{(k-1)k}{4}} \prod_{j=1}^k \Gamma(x + (1-j)/2)$$

In order to obtain the joint posterior probability distribution of θ, Σ , that is not conditioned by the random variable Z , we integrate equation (20) relative to Z as follows:

$$P(\theta, \Sigma | Y) = \int_0^\infty P(\theta, \Sigma | Y, Z) P(Z) dZ$$

$$P(\theta, \Sigma | Y)$$

$$= \frac{|B_3|^{\frac{n}{2}} |C^T C|^{\frac{k}{2}} \left(\frac{\lambda}{\psi}\right)^{\frac{k(n+(p+s+1))}{4}}}{(2\pi)^{\frac{(p+s+1)k}{2}} |\Sigma|^{\frac{n+(p+s+1)+k+1}{2}} 2^{\frac{nk}{2}} \Gamma_k\left(\frac{n}{2}\right) K_v(\sqrt{\lambda\psi})} \frac{\Gamma_k\left(\frac{n+(p+s+1)}{2}\right)}{(\pi)^{\frac{(p+s+1)k}{2}} \Gamma_k\left(\frac{n}{2}\right)}, \quad r \text{ is degree of freedom}$$

$$* K_{\frac{2v-k(n+(p+s+1))}{2}} \left(\sqrt{\lambda\psi \left(1 + \frac{\text{tr} \left(B_3 + (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \right) \Sigma^{-1}}{\psi} \right)} \right) \frac{Q_{(p+s+1) \times k}}{Q_{(p+s+1) \times k} P_{k \times (p+s+1)}} \quad \text{And using property of Box and Tiao .}$$

$$* \left(1 + \frac{\text{tr} \left(B_3 + (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \right) \Sigma^{-1}}{\psi} \right)^{\frac{2v-k(n+(p+s+1))}{4}} \frac{R(k, (p+s+1), r)}{R(k, (p+s+1), r)} \frac{|B_3|^{-\frac{(p+s+1)}{2}} |C^T C|^{\frac{k}{2}}}{\left| I_{(p+s+1)} + (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \right|^{\frac{n+(p+s+1)}{2}}}$$

To find the posterior marginal probability distribution of the parameter matrix θ conditioned by the random variable Z , we

integrate equation (20) relative to the scale matrix Σ as follows:

$$P(\theta | Y, Z) = \int P(\theta, \Sigma | Y, Z) d\Sigma = \frac{|B_3|^{\frac{n}{2}} |C^T C|^{\frac{k}{2}} \Gamma_k\left(\frac{n+(p+s+1)}{2}\right)}{(\pi)^{\frac{(p+s+1)k}{2}} \Gamma_k\left(\frac{n}{2}\right)} * |B_3|$$

$$+ (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \left| I_k + (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) B_3^{-1} \right|^{-\frac{n+(p+s+1)}{2}} \quad \dots (22)$$

It is possible to rewrite equation (22) as follow:

$$P(\theta | Y, Z) = R(k, (p+s+1), r) * \frac{|B_3|^{-\frac{(p+s+1)}{2}} |C^T C|^{\frac{k}{2}}}{\left| I_k + (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) B_3^{-1} \right|^{\frac{n+(p+s+1)}{2}}} \quad \dots (23)$$

As:

$$R(k, (p+s+1), r)$$

$$\frac{\Gamma_k\left(\frac{n+(p+s+1)}{2}\right)}{(\pi)^{\frac{(p+s+1)k}{2}} \Gamma_k\left(\frac{n}{2}\right)}, \quad r \text{ is degree of freedom}$$

Thus:

$$P(\theta | Y, Z)$$

$$= \frac{R(k, (p+s+1), r)}{R(k, (p+s+1), r)} \frac{|B_3|^{-\frac{(p+s+1)}{2}} |C^T C|^{\frac{k}{2}}}{\left| I_{(p+s+1)} + (\theta - \hat{\theta}^*)^T C^T C (\theta - \hat{\theta}^*) \right|^{\frac{n+(p+s+1)}{2}}}$$

And:

$$R(k, (p+s+1), r) = R((p+s+1), k, r)$$

$$\frac{\Gamma_k(\frac{n+(p+s+1)}{2})}{(\pi)^{\frac{(p+s+1)k}{2}} \Gamma_k(\frac{n}{2})} = \frac{\Gamma_{(p+s+1)}(\frac{n+(p+s+1)}{2})}{(\pi)^{\frac{(p+s+1)k}{2}} \Gamma_{(p+s+1)}(\frac{n-k+(p+s+1)}{2})}$$

Hence, the posterior marginal probability distribution of the parameter matrix θ unconditional of the variable Z is as follows:

$$P(\theta | Y) = \frac{|B_3|^{-\frac{(p+s+1)}{2}} |C^T C|^{\frac{k}{2}} \Gamma_{(p+s+1)}(\frac{n+(p+s+1)}{2})}{(\pi)^{\frac{(p+s+1)k}{2}} \Gamma_{(p+s+1)}(\frac{n-k+(p+s+1)}{2})} * |I_{(p+s+1)} + (\theta - \hat{\theta}^*) B_3^{-1} (\theta - \hat{\theta}^*)^T C^T C|^{-\frac{n+(p+s+1)}{2}} \dots (25)$$

We notice from equation (25) that the posterior marginal probability distribution of θ is matrix-t distribution of the parameters $(\hat{\theta}^*, B_3, (C^T C)^{-1}, r = n - k + 1)$ and is described as follows:

$$\theta \sim Mt_{(p+s+1),k}(\hat{\theta}^*, B_3, (C^T C)^{-1}, r), \text{vec}(\theta) \sim Mt_{(p+s+1),k}(\text{vec}(\hat{\theta}^*), B_3 \otimes (C^T C)^{-1}, r)$$

And the estimate Bayes of θ is:

$$\hat{\theta}_B = \hat{\theta}^* = (C^T C)^{-1} C^T Y$$

This estimate is similar if the model error

follows the distribution of the normal matrix

and the t-matrix.

To find the posterior marginal probability distribution of the scale matrix Σ conditioned by Z , we integrate equation (20) relative to the location matrix θ as follows:

$$P(\Sigma | Y, Z) = \int_{\theta} P(\theta, \Sigma | Y, Z) d\theta = \frac{|B_3|^{\frac{n}{2}} Z^{-\frac{nk}{2}}}{|\Sigma|^{\frac{n+k+1}{2}} 2^{\frac{nk}{2}} \Gamma_k(\frac{n}{2})} \exp\left(-\frac{1}{2Z} \text{tr } B_3 \Sigma\right)$$

Equation (26) represents the inverse wishart distribution of parameters $(\frac{B_3}{Z}, n)$ of the posterior marginal probability distribution of Σ conditional by the random variable Z . Based on Bayes theorem we conclude that the posterior marginal probability distribution of the scale matrix Σ unconditioned by Z is uncommon distribution as follows:

$$P(\Sigma | Y) = \int_0^\infty P(\Sigma | Y, Z) P(Z) dZ = \frac{|B_3|^{\frac{n}{2}} \left(\frac{\lambda}{\psi}\right)^{\frac{nk}{4}}}{|\Sigma|^{\frac{n+k+1}{2}} 2^{\frac{nk}{2}} \Gamma_k(\frac{n}{2}) K_v(\sqrt{\lambda\psi})} K_{\frac{2v-nk}{2}} \left(\sqrt{\lambda\psi \left(1 + \frac{\text{tr } B_3 \Sigma^{-1}}{\psi}\right)} \right) + \frac{\text{tr } B_3 \Sigma^{-1}}{\psi} \right)^{\frac{2v-nk}{4}} \dots (28)$$

And the estimate Bayes of Σ is:

$$\hat{\Sigma}_B = E_Z E_{\Sigma|Z}(\Sigma | Y, Z) = \frac{(Y - C\hat{\theta}^*)^T (Y - C\hat{\theta}^*)}{n - k - 1} \frac{K_{v-1}(\sqrt{\lambda\psi})}{K_v(\sqrt{\lambda\psi})} \left(\frac{\lambda}{\psi}\right)^{0.5}$$

6- Conclusions and Recommendations

In this paper, a multivariable partial regression model was used when the error

follows a generalized Bessel matrix distribution, as an alternative to the model where the error follows a normal matrix distribution, to find Bayesian estimates for the model parameters. It was found that the subsequent marginal probability distribution of the parameter matrix θ follows the t-matrix distribution with parameters $(\hat{\theta}^*, B_3, (C^T C)^{-1}, r = n - k + 1)$, defined in Equation (25). Furthermore, the subsequent marginal probability distribution of the measure matrix Σ is an uncommon but suitable and well-defined probability distribution, as defined in Equation (28). The researchers recommend conducting a practical application of these findings using the core and introductory parameters defined in sections (3) and (4), respectively, under different loss functions, and using this as an alternative to the multivariable partial regression model.

References

1. Thabane, L. & Haq, M. S. (2003) "The generalized multivariate modified Bessel distribution and its Bayesian applications" *Journal of Applied Statistical Science*, vol. 11(3), p.p.225-267.
2. Thabane, L. & Haq, M. S. (2004) "On the matrix-variate generalized hyperbolic distribution and its Bayesian applications" *Journal of Theoretical and Applied Statistical Science*, vol. 38(6), p.p.511-526.
3. Choi, T., Lee, J. & Roy, A. (2009) "A note on the bayes factor in a semiparametric regression model", *Journal of Multivariate Analysis* (100), p.p.1316-1327.
4. AL-Mouel, A. S. and Mohaisen A. J. (2017) "Study on Bayes Semiparametric Regression", *Journal of Advances in Applied Mathematics*, Vol. 2, No. 4, pages 197-207 .
5. Gallagher, M.P. & McNicholas, P.D. (2019) "Three skewed matrix variate distribution," *Statistics & Probability Letters*, vol.145, p.p. 103-109.
6. Koudou, A. E. & Ley, C. (2014) "Characterizations of GIG laws: a survey complemented with two new result", *Proba. Surv.*, vol. 11, p.p. 161-176.
7. Przystalski, M. (2014), "Estimation of the covariance matrix in multivariate partially linear models", *Journal of Multivariate Analysis*, Vol.123, p.p. 380 -385.
8. Jefferys, H. (1961) "Theory of Probability", Clarendon press, Oxford, London U.K.
9. Hmood, M. Y. & Hassan, Y. A. (2020)" Estimate the Partial Linear Model Using Wavelet and Kernel Smoothers", *Journal of Economics and Administrative Sciences*, Volume 26, Issue 119, p.p. 428-443.
10. Langrene, N. & Warin, X. (2019)" Fast and stable multivariate kernel density estimation by fast sum updating", *Journal of Computational and Graphical Statistics*, vol.28 ,p.p. 596-608.
11. Schucany, W. R. & Sommers, J. (1977)" Improved of kernel type density estimators ", *JASA* ,vol.72, no. 353, p.p. 420-423.
12. Box G. P. and Tiao G. C. (1973) "Bayesian Inference in Statistical Analysis" Addison Wesley publishing company, Inc. London, U.K.



AL-HAMDANIYA JOURNAL OF PURE AND APPLIED MATHEMATICS

<https://journal.uohamdaniya.edu.iq/index.php/hjppm>

ISSN: 0000-0000 (Online)

On Banach-Valued Measurable Functions

Wafa Y. Yahya¹, Noori F. Al-Mayahi²

¹Mathematics Department, Education College, Al-Hamdaniya University, Iraq

²Mathematics Department, Science College, Al-Qadisiyah University, Iraq

¹E-mail of First Author : rwafa1993@uohamdaniya.edu.iq

²E-mail of Second Author : nfam60@yahoo.com

ARTICLE INFO

Article history:

Received 1/5/2026
Revised 25/5/2026,
Accepted 5/6/2026,
Available online 25/6/2026

Keywords:

Measurable space
Measurable function
Ordered Banach space
Borel measurable function
Banach valued measurable function

ABSTRACT

This study investigated the concept of Banach-valued measurable functions and established several fundamental results related to measurable mappings in ordered Banach spaces. The research aimed to connect concepts from functional analysis and measure theory through the introduction of new measurable structures associated with Banach spaces. In particular, the study introduced the definitions of Borel Banach-valued measurable sets and Banach-valued measurable functions on measurable spaces. Several important properties and characterizations of these functions were developed and analyzed.

The work examined the equivalence between different conditions for Borel measurability in ordered Banach spaces and proved a collection of theorems concerning measurable functions and measurable sets. In addition, closure properties under algebraic operations such as addition, subtraction, multiplication, maximum, and minimum were established. The measurability of compositions of measurable functions was also studied and proved. Furthermore, the relationship between continuity and measurability was investigated, showing that continuous functions between topological spaces induce measurable mappings. The study also demonstrated that the family of Banach-valued measurable functions forms a linear space under suitable conditions.

The obtained results provided a theoretical framework that strengthened the relationship between measure theory and functional analysis and contributed to the development of measurable structures in ordered Banach spaces.

1. Introduction

The Functional analysis and Measure are two of mathematical analysis branches (so that they are two different subjects and each one of them has its own pathway). This research is a subject about Functional analysis but the tools used are from the Measure subject, and the main concept taken from the Measure subject which will be focused on is the set function. In the past, the Measure theory was studied as a function from a family of the whole bounded open intervals of real number $\mathbb{R}$ in to $\mathbb{R}$. Then, the measure has been studied as a non-negative extended real-valued set function defined on

σ – field Γ which is subsets of non-empty set $\mathfrak{X}$.

This research is a new subject in the set functions of Banach valued measurable function that links up between functional analysis and measure theory. The purpose of this study is to give a definition for Banach valued measurable function and provide some properties and theorems about it.

Many authors have studied measure like: A.N. Kolmogoroff (1933), G. Choquet (1954), Halmos (1970), M. Sugeno (1974), L.A. Zadeh (1978), Capinski and Kopp (1998), Dudley (2004) and B. Liu (2007).

2. Banach Valued Measurable Function

Explaining We introduced new definitions which are Borel Banach valued measurable set and Banach valued measurable function and obtained results related to these definitions. We also provided some basic properties for these concepts.

2.1 Definition [10]

The ordered vector space is a real linear space $\mathbb{X}$ provided with an antisymmetric, reflexive, transitive relation $\leq$ satisfying the upcoming conditions:

- 1) If w, p, ℓ are elements of $\mathbb{X}$ and $w \leq p$, then $w + \ell \leq p + \ell$
- 2) If w, p are elements of $\mathbb{X}$ and α is a positive real number, then $w \leq p$ implies $\alpha w \leq \alpha p$.

The notation $p \geq w$ often will be used in place of $w \leq p$. If $w \leq p$ and $w \neq p$, we shall write $w < p$.

2.2 Definition [2]

Let $(\mathbb{X}, \Gamma)$ be a topological space. σ – field emerged by Γ is named the Borel σ – field and it is meant to be $\mathfrak{B}(\mathbb{X})$, i.e. $\mathfrak{B}(\mathbb{X}) = \sigma(\Gamma)$. The members of $\mathfrak{B}(\mathbb{X})$ are called Borel sets of $\mathbb{X}$.

2.3 Definition [4]

Let $(\mathbb{X}, \Gamma)$ be a measurable space, and $\mathcal{W}$ a Banach space. The function $\mathcal{M}: \Gamma \rightarrow \mathcal{W}$ is called a vector measure if:

- 1) $\mathcal{M}(\Lambda) \geq 0$ for all $\Lambda \in \Gamma$, where 0 is the zero element in $\mathcal{W}$.
- 2) $\mathcal{M}(\cup_{n=1}^{\infty} \Lambda_n) = \sum_{n=1}^{\infty} \mathcal{M}(\Lambda_n)$ whenever $\{\Lambda_n\}$ is a sequence of disjoint sets in Γ , and $\cup_{n=1}^{\infty} \Lambda_n \in \Gamma$,
i.e. $\mathcal{M}$ is a countably additive (σ – additive).

2.4 Definition

Let $(\mathbb{X}, \Gamma)$ and $(\mathbb{X}', \hat{\Gamma})$ be two of measurable spaces. The function $\theta: \mathbb{X} \rightarrow \mathbb{X}'$ is said to be measurable with reference to Γ and $\hat{\Gamma}$, if $\theta^{-1}(\Lambda) \in \Gamma$ for every $\Lambda \in \hat{\Gamma}$.

The notation $\theta: (\mathbb{X}, \Gamma) \rightarrow (\mathbb{X}', \hat{\Gamma})$ will mean that $\theta: \mathbb{X} \rightarrow \mathbb{X}'$ measurable with respect to Γ and $\hat{\Gamma}$. [5]

For example,

- 1) If $\hat{\Gamma} = \{\emptyset, \mathbb{X}'\}$, therefore every function $\theta: \mathbb{X} \rightarrow \mathbb{X}'$ is measurable.
- 2) If $\Gamma = \{\Lambda: \Lambda \subseteq \mathbb{X}\}$, i.e. Γ is the family of the whole subsets of $\mathbb{X}$, therefore every function $\theta: \mathbb{X} \rightarrow \mathbb{X}'$ is measurable.

2.5 Theorem [1]

Let $(\mathbb{X}, \Gamma)$ and $(\mathbb{X}', \hat{\Gamma})$ be two of measurable spaces, and $\mathcal{P}$ be the family of subsets of $\mathbb{X}'$ such that $\sigma(\mathcal{P}) = \hat{\Gamma}$, then $\theta: (\mathbb{X}, \Gamma) \rightarrow (\mathbb{X}', \hat{\Gamma})$ is measurable if $\theta^{-1}(\Lambda) \in \Gamma$ for all $\Lambda \in \mathcal{P}$.

2.6 Definition

Let $(\mathbb{X}, \Gamma)$ be a measurable space and $\mathcal{W}$ be an ordered Banach space. The function $\theta: \mathbb{X} \rightarrow \mathcal{W}$ is named as Borel Banach valued measurable on $(\mathbb{X}, \Gamma)$ if θ is measurable in reference to Γ and $\mathfrak{B}(\mathcal{W})$, that is $\theta^{-1}(B_{red}(w)) = \{\ell \in \mathbb{X}: \|\theta(\ell) - w\| \leq red\} = \{\ell \in \mathbb{X}: \theta(\ell) \in B_{red}(w)\}$, that is $\|\theta(\ell) - w\| \leq red\}$, where $w \in \mathcal{W}$, $red \in \mathbb{R}$ and we write $\theta \in \Gamma$.

2.7 Remark

First:

Let $\mathcal{W}$ be an ordered Banach space, $w \in \mathcal{W}$, non negative $red, red_1, red_2 \in \mathbb{R}$ and $\theta: \mathbb{X} \rightarrow \mathcal{W}$ be a function

- 1) $\{\|\theta - w\| \leq red\} = \{\ell \in \mathbb{X}: \theta(\ell) \in B_{red}(w), \text{ i.e. } \|\theta(\ell) - w\| \leq red\}$
- 2) $\{\|\theta - w\| < red\} = \{\ell \in \mathbb{X}: \theta(\ell) \in B_{red}(w), \text{ i.e. } \|\theta(\ell) - w\| < red\}$
- 3) $\{\|\theta - w\| \geq red\} = \{\ell \in \mathbb{X}: \theta(\ell) \in B_{red}(w), \text{ i.e. } \|\theta(\ell) - w\| \geq red\}$
- 4) $\{\|\theta - w\| > red\} = \{\ell \in \mathbb{X}: \theta(\ell) \in B_{red}(w), \text{ i.e. } \|\theta(\ell) - w\| > red\}$
- 5) $\{\|\theta - w\| \leq red\}^c = \{\|\theta - w\| > red\}$
- 6) $\{\|\theta - w\| < red\}^c = \{\|\theta - w\| \geq red\}$

- 7) $\{\|\theta - w\| \leq red\} = \bigcap_{n=1}^{\infty} \left\{ \|\theta - w\| < red + \frac{1}{n} \right\}$
- 8) $\{\|\theta - w\| < red\} = \bigcup_{n=1}^{\infty} \left\{ \|\theta - w\| \leq red - \frac{1}{n} \right\}$
- 9) $\{\|\theta - w\| \geq red\} = \bigcap_{n=1}^{\infty} \left\{ \|\theta - w\| > red - \frac{1}{n} \right\}$
- 10) $\{\|\theta - w\| > red\} = \bigcup_{n=1}^{\infty} \left\{ \|\theta - w\| \geq red + \frac{1}{n} \right\}$
- 11) $\{\|\theta - w\| = red\} = \{\|\theta - w\| \geq red\} \cap \{\|\theta - w\| \leq red\}$
- 12) $\{red_1 < \|\theta - w\| < red_2\} = \{\|\theta - w\| > red_1\} \cap \{\|\theta - w\| < red_2\}$

Second:

Let $\theta, \eta: \mathfrak{X} \longrightarrow \mathcal{W}$ be functions and w be an element in $\mathcal{W}$

- 1) $\{\theta = \eta\} = \{\theta \leq \eta\} \cap \{\theta \geq \eta\}$
- 2) $\{\theta > \eta\} = \bigcup_{n=1}^{\infty} \left(\bigcup_{m=1}^{\infty} \left\{ \|\theta - w\| > \frac{m}{n} \right\} \cap \left\{ \|\eta - w\| < \frac{m}{n} \right\} \right)$
- 3) $\text{Min}\{\theta, \eta\}(\ell) = \min\{\theta(\ell), \eta(\ell)\}$
- 4) $\text{Max}\{\theta, \eta\}(\ell) = \max\{\theta(\ell), \eta(\ell)\}$
- 5) $\{\|\min\{\theta, \eta\} - w\| < red\} = \{\|\theta - w\| < red\} \cup \{\|\eta - w\| < red\}$
- 6) $\{\|\min\{\theta, \eta\} - w\| > red\} = \{\|\theta - w\| > red\} \cap \{\|\eta - w\| > red\}$
- 7) $\{\|\max\{\theta, \eta\} - w\| < red\} = \{\|\theta - w\| < red\} \cap \{\|\eta - w\| < red\}$
- 8) $\{\|\max\{\theta, \eta\} - w\| > red\} = \{\|\theta - w\| > red\} \cup \{\|\eta - w\| > red\}$

Third:

We express the upcoming classes of subsets of $\mathcal{W}$

- 1) $\mathcal{P}_1 = \{B_{red}(w)\} = \{t \in \mathcal{W}: \|t - w\| < red\}$
- 2) $\mathcal{P}_2 = \{\bar{B}_{red}(w)\} = \{t \in \mathcal{W}: \|t - w\| \leq red\}$
- 3) $\mathcal{P}_3 = \{B^c_{red}(w)\} = \{t \in \mathcal{W}: \|t - w\| \geq red\}$
- 4) $\mathcal{P}_4 = \{\bar{B}^c_{red}(w)\} = \{t \in \mathcal{W}: \|t - w\| > red\}$

It is clear to show that $\sigma(\mathcal{P}_i) = \mathfrak{B}(\mathcal{W})$ for all $i = 1, 2, 3, 4$

2.8 Theorem

Let $(\mathfrak{X}, \Gamma)$ be a measurable space and $\theta: \mathfrak{X} \longrightarrow \mathcal{W}$ a function, where $\mathcal{W}$ is an ordered Banach space. The upcoming statements are compatible:

- 1) f is a Borel measurable on $(\mathfrak{X}, \Gamma)$
- 2) $\{\|\theta - w\| \leq red\} \in \Gamma$ for every $w \in \mathcal{W}$ & for every $red > 0$
- 3) $\{\|\theta - w\| < red\} \in \Gamma$ for every $w \in \mathcal{W}$ & for every $red > 0$
- 4) $\{\|\theta - w\| \geq red\} \in \Gamma$ for every $w \in \mathcal{W}$ & for every $red > 0$
- 5) $\{\|\theta - w\| > red\} \in \Gamma$ for every $w \in \mathcal{W}$ & for every $red > 0$

Proof:

- $1 \Rightarrow 2$ Since θ is measurable and $\{\|\theta - w\| \leq red\} = \theta^{-1}(B_{red}(w))$
Since $B_{red}(w) \in \mathfrak{B}(\mathcal{W}) \Rightarrow \{\|\theta - w\| \leq red\} \in \Gamma$
- $2 \Rightarrow 3$ $\{\|\theta - w\| < red\} = \bigcup_{n=1}^{\infty} \left\{ \|\theta - w\| \leq red - \frac{1}{n} \right\} \in \Gamma$
- $3 \Rightarrow 4$ $\{\|\theta - w\| \geq red\} = \{\|\theta - w\| < red\}^c \in \Gamma$
- $4 \Rightarrow 5$ $\{\|\theta - w\| > red\} = \bigcup_{n=1}^{\infty} \left\{ \|\theta - w\| \geq red + \frac{1}{n} \right\} \in \Gamma$
- $5 \Rightarrow 1$ Since $\sigma(\mathcal{P}_4) = \mathfrak{B}(\mathcal{W})$, by using theorem (2.5), we have θ a Borel measurable on $(\mathfrak{X}, \Gamma)$. ■

2.9 Corollary

Let $(\mathfrak{X}, \Gamma)$ be a measurable space, and $\theta, \eta: \mathfrak{X} \longrightarrow \mathcal{W}$ be Borel measurable functions, where $\mathcal{W}$ be an ordered Banach space.

Therefore:

- 1) $\{\|\theta - w\| = red\} \in \Gamma$ for all $w \in \mathcal{W}$ and for all $red > 0$
- 2) $\{red_1 < \|\theta - w\| < red_2\} \in \Gamma$ for all w and for all $red_1, red_2 > 0$
- 3) The functions $\max\{\theta, \eta\}$ and $\min\{\theta, \eta\}$ are measurable.

2.10 Theorem

Let $\mathcal{W}$ be an ordered Banach space, $(\mathfrak{X}, \Gamma)$ be a measurable space and

$\theta, \eta: \mathfrak{X} \longrightarrow \mathfrak{B}(\mathcal{W})$ two of measurable functions, then

- 1) $\{\|\theta - \eta\| < 0\} = \bigcup_{red \in Q} (\{\|\theta - w\| < red\} \cap \{red < \|\eta - w\|\})$
 $= \bigcup_{red \in Q} \{\|\theta - w\| < red < \|\eta - w\|\}$
- 2) The sets $\{\|\theta - \eta\| < 0\}, \{\|\theta - \eta\| > 0\}, \{\|\theta - \eta\| = 0\}, \{\|\theta - \eta\| \leq 0\}, \{\|\theta - \eta\| \geq 0\}$ belong to the σ -field Γ .

Proof:

- 1) Let $\ell \in \{\|\theta - \eta\| < 0\} \Rightarrow \theta(\ell) < \eta(\ell)$
 there is $w \in \mathcal{W}$ so that $\theta(\ell) < w < \eta(\ell)$.

This pursued by $\ell \in \{\|\theta - w\| < red\} \cap \{red < \|\eta - w\|\}$ for all $red \in Q$.

So $\ell \in \bigcup_{red \in Q} (\{\|\theta - w\| < red\} \cap \{red < \|\eta - w\|\})$, then

$\{\|\theta - \eta\| < 0\} \subset \bigcup_{red \in Q} (\{\|\theta - w\| < red\} \cap \{red < \|\eta - w\|\})$.

The reverse inclusion is clear.

- 2) Since θ and η are measurable; $\{\|\theta - w\| < red\} = \theta^{-1}(B_{red}(w))$ and $\{red < \|\eta - w\|\} = \eta^{-1}(B_{red}^c(w))$ are both elements of Γ for all $B_{red}(w) \in \mathfrak{B}(\mathcal{W})$ and $red \in Q$ due to the use of (1).

Since Q is countable set, this leads to $\{\|\theta - \eta\| \leq 0\} \in \Gamma$. In the same way in $\{\|\eta - \theta\| \leq 0\} \in \Gamma$.

Furthermore, we have $\{\|\theta - \eta\| \leq 0\} = \{\|\eta - \theta\| \leq 0\}^c \in \Gamma$ and $\{\|\eta - \theta\| \leq 0\} = \{\|\theta - \eta\| < 0\}^c \in \Gamma$

Lastly, $\{\|\theta - \eta\| = 0\} = \{\|\theta - \eta\| \leq 0\} \cap \{\|\eta - \theta\| \leq 0\} \in \Gamma$. ■

2.11 Theorem

Let $(\mathfrak{X}, \Gamma)$ be a measurable space, $\theta, \eta: \mathfrak{X} \longrightarrow \mathcal{W}$ be two measurable functions, and $\alpha \in \mathbb{R}$. If $\mathcal{W}$ be an ordered Banach algebra space, then each of the functions $\alpha\theta$, $\theta \pm \eta$, θ^2 , $\theta\eta$, $\frac{1}{\theta}$ which is defined and measurable, where θ is invertable.

Proof:

Consider the function $\theta + \eta$

$\{\|(\theta + \eta) - w < red\|\} = \bigcup_{q \in Q} (\{\|\theta - w\| < q\} \cap \{\|\eta - w\| < red - q\})$ By theorem (2.8) every set on the right side is in Γ , hence the set on the left-hand is also in Γ , and $\theta + \eta$ must be measurable.

Now

$$\{\|\theta^2 - w\| < red\} = \begin{cases} \emptyset & , red < 0 \\ \emptyset & , red = 0 \\ \{\|\theta - w\| < red\} \cup \{\|\theta - w\| \geq red\} & , red > 0 \end{cases}$$

$$\left\{\left\|\frac{1}{\theta} - w\right\| < red\right\} = \begin{cases} \emptyset & , red < 0 \\ \emptyset & , red = 0 \\ \{\|\theta - w\| < red\} \cup \{\|\theta - w\| \geq red\} & , red > 0 \end{cases}$$

Since $red > 0$, then

$$\begin{aligned} \{\|\theta^2 - w\| < red\} &= \{\|\theta - w\| < red\} \cup \{\|\theta - w\| \geq red\} \\ \left\{\left\|\frac{1}{\theta} - w\right\| < red\right\} &= \{\|\theta - w\| < red\} \cup \{\|\theta - w\| \geq red\} \end{aligned}$$

and each of these sets in Γ , so that θ^2 and $\frac{1}{\theta}$ are measurable.

Since $\theta\eta = \frac{1}{2}((\theta + \eta)^2 - \theta^2 - \eta^2)$, then $\theta\eta$ in Γ .

Now we know that the set of Banach valued measurable functions are liner space. ■

2.12 Theorem

Let $(\mathfrak{X}, \Gamma)$, $(\mathcal{W}, \mathfrak{B}(\mathcal{W}))$ and $(\mathcal{V}, \mathfrak{B}(\mathcal{V}))$ be three measurable spaces, and let

$\theta: (\mathfrak{X}, \Gamma) \longrightarrow (\mathcal{W}, \mathfrak{B}(\mathcal{W}))$ and $\eta: (\mathcal{W}, \mathfrak{B}(\mathcal{W})) \longrightarrow (\mathcal{V}, \mathfrak{B}(\mathcal{V}))$ be two measurable functions, then $\eta \circ \theta: (\mathfrak{X}, \Gamma) \longrightarrow (\mathcal{V}, \mathfrak{B}(\mathcal{V}))$ is measurable function. Hence measurable functions composition is measurable.

Proof:

Let $\Lambda \in \mathfrak{B}(\mathcal{V})$. Since η is measurable, then $\eta^{-1}(\Lambda) \in \mathfrak{B}(\mathcal{W})$ and since θ is measurable, then $\theta^{-1}(\eta^{-1}(\Lambda)) \in \Gamma$
 But $\theta^{-1} \circ \eta^{-1} = (\eta \circ \theta)^{-1}$, hence $(\eta \circ \theta)^{-1}(\Lambda) \in \Gamma$, then $\eta \circ \theta$ is measurable. ■

2.13 Theorem

Let $(\mathfrak{X}, \mathfrak{T})$ and $(\mathfrak{X}, \hat{\mathfrak{T}})$ be two topological spaces. If the function $\theta: (\mathfrak{X}, \mathfrak{T}) \longrightarrow (\mathfrak{X}, \hat{\mathfrak{T}})$ is continuous, then $\theta: (\mathfrak{X}, \mathfrak{B}(\mathfrak{X})) \longrightarrow (\mathfrak{X}, \mathfrak{B}(\hat{\mathfrak{X}}))$ is measurable.

Proof:

Let $\Lambda \in \mathfrak{B}(\mathcal{W})$. Since $\mathfrak{B}(\mathcal{W}) = \sigma(\hat{\mathfrak{T}})$ then $\Lambda \in \sigma(\hat{\mathfrak{T}})$, and as long as θ is continuous, then $\theta(\Lambda) \in \mathfrak{T}$. As for every $\mathfrak{T} \subseteq \sigma(\mathfrak{T})$ then $\theta^{-1}(\Lambda) \in \sigma(\mathfrak{T})$. By the use of Theorem (2.5), it is concluded that $\theta: (\mathfrak{X}, \mathfrak{B}(\mathfrak{X})) \longrightarrow (\mathfrak{X}, \mathfrak{B}(\hat{\mathfrak{X}}))$ is measurable. ■

2.14 Corollary

Let $(\mathfrak{X}, \mathfrak{T})$ be a measurable space, $\mathcal{W}$ be a Banach space, $\theta: \mathfrak{X} \longrightarrow \mathcal{W}$ be a measurable function and $\eta: \mathcal{W} \longrightarrow \mathcal{W}$ be a continuous function, then $\eta \circ \theta: \mathfrak{X} \longrightarrow \mathcal{W}$ is a measurable function.

References

- [1] R. B. Ash and C. A. Doleans-Dade, Academic Press (Londyn), Probability and measure theory, Elsevier Science, 2000.
- [2] R. B. Ash and Karreman Mathematics Research Collection, Real analysis and probability, Academic Press, 1972.
- [3] R. de Jong, Ordered banach algebras. Ms C Thesis, Leiden University. Holland, 2010.
- [4] J. Diestel and J. J. Uhl, Vector Measure, Mathematical Surveys and Monographs, American Mathematical Society, 1977.
- [5] N. Dinculeanu, Vector Measures, International Series of Monographs on Pure and Applied Mathematics, 1967.
- [6] R. G. Douglas, Banach algebra techniques in operator theory, Graduate Texts in Mathematics, Springer New York, 1998.
- [7] P. R. Halmos, Measure Theory, Springer New York, 2013.
- [8] L. Nachbin, I. Kluváněk, G. Knowles, Vector Measures and Control Systems, Elsevier Science, 2011.
- [9] N. F. AL-Mayahi, Measure Theory and Applications, Deposit number in the National

Library and Archives in Baghdad (1780). First edition, Iraq (2018).

- [10] A. L. Peressini, Ordered Topological Vector Spaces, Harper and Row, 1967.
- [11] J. R. Retherford, Review: J. Diestel and J. J. Uhl, Jr., Vector Measures, Bulletin of the American Mathematical Society. 1978;84 (4): 681-5,5.
- [12] S. J. Taylor, Introduction to Measure and Integration, Cambridge University Press, 1973.



AL-HAMDANIYA JOURNAL OF PURE AND APPLIED MATHEMATICS

<https://journal.uohamdaniya.edu.iq/index.php/hjppm>

ISSN: 0000-0000 (Online)

Optimal Linear Codes arising from Multiple Cap by Weight Distribution

Thakreen faisal sultan ¹, Hajir Hayder Abdullah², Nada Yassen Kasm Yahya ³

^{1,2} Department of mathematics, College of Education for pure Scienc University of Al-Hamdaniya, Nineveh , Iraq

³ Department of Mathematics, College of Education for Pure Science, University of Mosul, Mosul,

¹E-mail of First Author : thakreenfaisal@uohamdaniya.edu.iq

²E-mail of Second Author : hajarhayder@uohamdaniya.edu.iq

³E-mail of Third Author : drnadaqasim3@uomosul.edu.iq

ARTICLE INFO

Article history:

Received 1/5/2026
Revised 25/5/2026,
Accepted 5/6/2026,
Available online 25/6/2026

Keywords:

Wight distribution,
optimal linear codes,
(k, r) –cap,
coding theory,
[n, k] –code

ABSTRACT

The main objectives of this thesis are obtained by finding a new examples construction optimal linear code with weight distribution of optimal linear code $[n, k]_q$ -code with $k=4$, dimension (4D) in $PG(3, 7)$. we found the minimum distance d and determine type of linear code and apply the results to the 4D optimal linear code the inferred from cap.

1. Introduction

vector space of n –tuples over F_7 , the field of 7elements. The weight of a vector $c \in F_q^n$ denoted by $wt(c)$ is the number of nonzero coordinate positions in C . An $[n, k, d]_q$ –code C is a F_q -linear sub space of F_q^n with dimension k and minimum (Hamming) distance d .

(a) A cap in $\Sigma = PG(k - 1, q)$, is a set of points with no three are collinear. the connection between caps in $PG(n, q)$ and

linear codes (of minimum distance) has been well-studied (examples [3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15,16]).

And we found the weight distribution of a linear code $[n, k, d]_q$ –code over afield F_7 we represent the codes by the complete (k, r) –caps of $PG(3, 7)$. And we show that all the codes presented in this paper are optimal. So, the main of this study is to how to find the new examples in weight distribution for the optimal linear code form the (k, r) –cap and this was prove practically using the Mat lab

computer program , and show that all the codes are optimal.

2. Preliminaries:

2.1 Definition: [1] A generator matrix G of an $[n, k]$ -code C is a $k \times n$ matrix whose rows form a basis for C .

$$G = [g_0 \ g_1 \dots g_{k-1}] = [g_{0,0} \dots g_{k-1,0} \ g_{0,1} \dots g_{k-1,1} \ g_{0,n-1} \dots g_{k-1,n-1}]$$

2.2 Linear codes: [2],[4],[8]

The space is $V(n, q) = (F_q^n, +, \times)$.

A q -ary linear code C of length n and dimension k , or a $[n, k]_q$ code, is a k -dimensional subspace of F_q^n . Every subspace of C is referred to as a subcode of C . The inner product of two vectors $u = (u_1, u_2, \dots, u_n)$ and $v = (v_1, v_2, \dots, v_n)$ from F_q^n is denoted by $uv = u_1v_1 + u_2v_2 + \dots + u_nv_n$.

Two vectors are said to be orthogonal if their inner product is 0. The set of all vectors of F_q^n orthogonal to all code words from C is called the orthogonal code $C^\perp$ to C : $C^\perp = \{x \in F_q^n \mid xy = 0 \text{ for any } y \in C\}$.

By a well-known fact from linear algebra, the $C^\perp$ is a linear $[n, n - k]_q$ code. [8]

2.3 Definition: [3] The (Hamming) weight $wt(x)$ of a vector x in F_q^n

is the number of its nonzero coordinates, such that $F_q^n = \{(a_1, a_2, \dots, a_n) \mid a_1, \dots, a_n \in F_q\}$.

3. Optimal linear codes problem: [6],[8]

Optimize one of the parameters n, k, d for given the other two For an $[n, k, d]_q$ code C , a generator matrix G is a $k \times n$ matrix over F_q whose k rows form a basis of C . We assume G has no all-zero column.

The weight distribution ($w.d.$) of C is the list of numbers $A_i = |\{c \in C \mid wt(c) = i\}|$.

The weight distribution with $(A_0, A_d, \dots, A_i, \dots) = (1, \alpha, \dots, w, \dots)$ is also expressed as $0^1 d^\alpha \dots i^w \cdot \text{An } [n, K, d]_q \text{ code } C$ is

N -optimal if $\nexists [n - 1, K, d]_q$

K -optimal if $\nexists [n, K + 1, d]_q$

D -optimal if $\nexists [n, K, d + 1]_q$.

N -optimal codes are k -optimal and D -optimal.

Problem 1: [6] Find $n_q(k, d)$, the minimum value of n for which an $[n, k, d]_q$ code exists for given K, d, q . An $[n, k, d]_q$ code is called optimal if $n = n_q(k, d)$.

3.2 Geometric method: [4]

The projective subspace of dimension r over F_q denoted by $PG(r, q)$ and j -dimensional projective subspace of $PG(r, q)$ denoted by j -flat

The 0 - flats, 1- flats, 2 - flats and $(r - 1)$ - flats are called points, lines, planes and hyperplanes respectively.

We denote by θ_j the number of Points in a j -flat $\theta_j := |PG(j, q)| = q^j + q^{j-1} + \dots + q + 1$. Assume C has no coordinate which is identically zero.

G : a generator matrix of C the columns of G can be considered as a multiset of n points in $\Sigma = PG(k - 1, q)$ denoted by M_C . Conversely,

M_C gives linear codes which are equivalent to C .

F_j : the set of all j -flats in Σ

$\Sigma \ni P: i$ -point $\Leftrightarrow P$ has multiplicity i in M_C
 $\gamma_0 = \max\{i \mid \exists P: i\text{-point in } \Sigma\}$

$Ci = \{P \in \Sigma \mid P: i\text{-point}\}, 0 \leq i \leq \gamma_0$

$\Delta_1 + \dots + \Delta_s$: the multiset consisting of the S sets

$\Delta_1, \dots, \Delta_s$ in Σ .

$s\Delta = \Delta_1 + \dots + \Delta_s$ when $\Delta_1 = \dots = \Delta_s = \Delta$.

Then, $M_C = C_1 + 2C_2 + \dots + \gamma_0 C_{\gamma_0}$.

For any set S in Σ , $M_C(S)$ is the multiset $\{P \in M_C \mid P \in S\}$.

The multiplicity of S , denoted by $m_C(S)$, is defined as

$$m_C(S) = |M_C(S)| = \sum_{i=1}^{\gamma_0} i |S \cap Ci|.$$

Then it holds that

$$n = m_C(\Sigma), \quad n - d = \max\{m_C(\pi) \mid \pi \in F_{k-2}\}.$$

Theorem (3.3): [7] Let C be a $[n, k, d]_q$ code. Then $n \geq d + k - 1$.

Now, the following are examples of natural geometric structures C' whose corresponding code C are optimal.

For $a = (a_1, \dots, a_n), b = (b_1, \dots, b_n) \in F_q^n$, the (Hamming) distance between a and b is $d(a, b) = |\{i \mid a_i \neq b_i\}|$.

The weight of a is $wt(a) = |\{i \mid a_i \neq 0\}| = d(a, 0)$.

An $[n, k, d]_q$ code C means a k -dimensional subspace of F_q^n with minimum distance d ,

$$d = \min\{d(a, b) \mid a \neq b, a, b \in C\} \\ = \min\{wt(a) \mid a \neq 0\}.$$

The elements of C are called code words.

(2.3.1) *General bounds*: In this section we present some general bounds on the parameters of a linear $[n, K, d]_q$ code.

Theorem (2.3.2): [2],[4] (The sphere-packing bound, or Hamming bound). Let C be an $[n, k, d]_q$ code and let $t = \lfloor (d-1)/2 \rfloor$. Then

$$\sum_{i=0}^t \binom{n}{i} (q-1)^i \leq q^{n-k}.$$

Theorem (2.3.3): [2] (The Varshamov–Gilbert bound). Let n, k, d be positive integers and let q be a prime power. If $\sum_{i=0}^{d-2} \binom{n-1}{i} (q-1)^i \leq q^{n-k}$ then there exists a $[n, k, d]_q$ code.

Theorem (2.3.4): [5] (The Grasmere bound). Let C be a $[n, k, d]_q$ code. Then $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lfloor \frac{d}{q^i} \rfloor$.

A Linear codes attaining the Grasmere bound; i.e. linear codes with parameters $[g_q(k, d), k, d]_q$ are called Griesmer codes.

For $a = (a_1, \dots, a_n), b = (b_1, \dots, b_n) \in F_q^n$, the (Hamming) distance between a and b is $d(a, b) = |\{i \mid a_i \neq b_i\}|$.

The weight of a is $wt(a) = |\{i \mid a_i \neq 0\}| = d(a, 0)$.

An $[n, k, d]_q$ code C means a k -dimensional subspace of F_q^n with minimum distance d ,

$$d = \min\{d(a, b) \mid a \neq b, a, b \in C\} \\ = \min\{wt(a) \mid a \neq 0\}.$$

The elements of C are called code words.

(2.3.1) *General bounds*: In this section we present some general bounds on the parameters of a linear $[n, K, d]_q$ code.

Theorem (2.3.2): [2],[4] (The sphere-packing bound, or Hamming bound). Let C be an $[n, k, d]_q$ code and let $t = \lfloor (d-1)/2 \rfloor$. Then

$$\sum_{i=0}^t \binom{n}{i} (q-1)^i \leq q^{n-k}.$$

Theorem (2.3.3): [2] (The Varshamov–Gilbert bound). Let n, k, d be positive integers and let q be a prime power. If $\sum_{i=0}^{d-2} \binom{n-1}{i} (q-1)^i \leq q^{n-k}$ then there exists a $[n, k, d]_q$ code.

Theorem (2.3.4): [5] (The Grasmere bound). Let C be a $[n, k, d]_q$ code. Then $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lfloor \frac{d}{q^i} \rfloor$.

A Linear codes attaining the Grasmere bound; i.e. linear codes with parameters $[g_q(k, d), k, d]_q$ are called Griesmer codes.

4. New Examples In Weight Distribution For The Optimal Linear Code.

Example(3): Let C be a $[23, 4, 15]_7$ – code with generator matrix

$G =$

$$\begin{bmatrix} 1 & 0 & 0 & 2 & 1 & 3 & 0 & 2 & 1 & 3 & 2 & 1 & 3 & 4 & 6 & 5 & 6 & 0 & 3 & 5 & 5 & 0 & 4 \\ 0 & 1 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 4 & 5 & 1 & 2 & 4 & 6 & 1 & 6 & 4 \\ 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 3 & 3 & 6 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \end{bmatrix}$$

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$ and p_j is defined as follows $\forall_i$:

$$= \{[1, 1, 1, 1], [1, 1, 1, 2], [1, 1, 1, 3], \dots, [0, 0, 0, 5], [0, 0, 0, 6], [0, 0, 0, 0]\}.$$

Therefore, $G = [p_{343}^T, p_{2107}^T, p_{2359}^T, p_{350}^T, p_{56}^T, p_{791}^T, p_{2395}^T, p_{386}^T, p_{92}^T, p_{827}^T, p_{638}^T, p_1^T, p_{736}^T, p_{1128}^T, p_{1863}^T, p_{1569}^T, p_{1723}^T, p_{2115}^T, p_{841}^T, p_{1625}^T, p_{1387}^T, p_{2318}^T, p_{1212}^T]$.

The code words is generated by multiplying each point in Σ to G . Therefore, the code words in $[23, 4, 15]_7$ linear code are defined as follows:

$$CW = \{ [p_i \times G]_{q^{k \times n}} \mid p_i \in \Sigma \} =$$

With **w.d.:**

$$0^1, 15^6, 16^{12}, 17^{150}, 18^{419}, 19^{606}, 20^{432}, 21^{402}, 22^{240}, 23^{138}.$$

Then we have:

$$C_1 = \{1, 2, 9, 18, 24, 33, 58, 67, 73, 82, 109, 115, 124, 132, 141, 147, 169, 170, 187, 203, 217, 247, 384\}$$

Example(4): Let C be a $[54, 4, 41]_7$ -code with generator matrix

$G =$

$$\begin{bmatrix} 1 & 0 & 0 & 2 & 1 & 3 & 0 & 2 & 1 & 3 & 2 & 1 & 3 & 4 & 6 & 5 & 6 & 0 & 3 & 5 & 5 & 0 & 4 & 1 & 1 & 0 & 4 & 5 \\ 0 & 1 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 4 & 5 & 1 & 2 & 4 & 6 & 1 & 6 & 4 & 1 & 0 & 1 & 2 & 3 \\ 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 3 & 3 & 6 & 0 & 1 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 2 & 1 & 3 & 0 & 2 & 1 & 3 & 2 & 1 & 3 & 4 & 6 & 5 & 6 & 0 & 3 & 5 & 5 & 0 & 4 & 1 & 1 & 0 & 4 & 5 \\ 0 & 1 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 4 & 5 & 1 & 2 & 4 & 6 & 1 & 6 & 4 & 1 & 0 & 1 & 2 & 3 \\ 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 3 & 3 & 6 & 0 & 1 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}$$

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$

and p_j is defined as follows $\forall_i: \Sigma = \{[1, 1, 1, 1], \dots, [0, 0, 0, 0]\}$. There fore $G = [p_{343}^T, p_{2107}^T, \dots, p_{1477}^T]$

The code words is generated by multiplying each point in Σ to G . Therefore, the code words in $[54, 4, 41]_7$ linear code are defined as follows:

$$CW=\{ [p_i \times G]_{q^k \times n} \mid p_i \in \Sigma \} = \begin{bmatrix} p_{11} & \cdots & p_{1n} \\ \vdots & \ddots & \vdots \\ p_{m1} & \cdots & p_{mn} \end{bmatrix}.$$

with **w.d.:**

$0^1, 41^{24}, 42^{42}, 43^{168}, 44^{342}, 45^{396}, 46^{378}, 47^{372}, 48^{270}, 49^{210}, 50^{78}, 51^{36}, 52^{36}, 53^{30}, 54^{18}.$

Then we have:

$C_1 = \{1, 2, 3, 9, 10, 16, 18, 24, 27, 33, 35, 46, 55, 58, 59, 65, 67, 73, 76, 82, 84, 104, 109, 110, 115, 118, 121, 124, 130, 132, 141, 147, 157, 158, 166, 169, 170, 172, 185, 187, 203, 211, 216, 217, 221, 224, 228, 240, 246, 247, 273, 336, 348, 369\}.$

Example(5): Let C be a $[87, 4, 68]_7$ –code with generator matrix

G=

$$\begin{bmatrix} 1 & 0 & 0 & 2 & 1 & 3 & 0 & 2 & 1 & 3 & 2 & 1 & 3 & 4 & 6 & 5 & 6 & 0 & 3 & 5 & 5 & 0 & 4 & 1 & 1 & 0 & 4 \\ 0 & 1 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 4 & 5 & 1 & 2 & 4 & 6 & 1 & 6 & 4 & 1 & 0 & 1 & 2 \\ 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 3 & 3 & 6 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 \\ \\ 5 & 2 & 4 & 1 & 0 & 4 & 5 & 4 & 3 & 4 & 0 & 2 & 1 & 2 & 3 & 2 & 1 & 6 & 4 & 2 & 5 & 2 & 0 & 6 & 5 & 5 & 3 & 2 & 2 & 3 \\ 3 & 5 & 6 & 0 & 1 & 2 & 3 & 6 & 0 & 1 & 2 & 3 & 0 & 0 & 1 & 2 & 4 & 0 & 1 & 2 & 2 & 3 & 5 & 5 & 2 & 4 & 2 & 1 & 2 & 1 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 2 & 3 & 3 & 3 & 3 & 3 & 3 & 4 & 5 & 6 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ \\ 3 & 2 & 4 & 0 & 1 & 5 & 4 & 0 & 3 & 0 & 1 & 6 & 1 & 5 & 6 & 0 & 0 & 6 & 1 & 5 & 2 & 2 & 1 & 0 & 0 & 6 & 3 & 5 & \\ 6 & 0 & 1 & 2 & 3 & 4 & 0 & 1 & 3 & 4 & 4 & 0 & 1 & 2 & 3 & 6 & 0 & 1 & 3 & 3 & 5 & 2 & 3 & 5 & 6 & 6 & 1 & 3 & \\ 1 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 2 & 3 & 3 & 3 & 3 & 3 & 4 & 4 & 4 & 4 & 5 & 5 & \\ 0 & 1 & \end{bmatrix} \begin{matrix} 2 \\ 6 \\ 4 \\ 5 \\ 5 \\ 1 \end{matrix}$$

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$

and p_j is defined as follows $\forall_i: \Sigma = \{[1,1,1,1], \dots, [0,0,0,0]\}$. Therefore,

$G = [p_{343}^T, p_{2107}^T, \dots, p_{1940}^T]$.

The code words is generated by multiplying each point in Σ to G.

Therefore, the code words in $[87, 4, 68]_7$ linear code are defined as follows:

$$CW=\{ [p_i \times G]_{q^k \times n} \mid p_i \in \Sigma \} = \begin{bmatrix} p_{11} & \cdots & p_{1n} \\ \vdots & \ddots & \vdots \\ p_{m1} & \cdots & p_{mn} \end{bmatrix}.$$

with **w.d.:**

$0^1, 68^6, 69^{48}, 70^{84}, 71^{198}, 72^{216}, 73^{342}, 74^{336}, 75^{336}, 76^{282}, 77^{198}, 78^{168}, 79^{66}, 80^{30}, 81^{36}, 82^{30}, 84^{12}, 85^6, 86^6.$

Then we have:

$C_1 = \{1, 2, 3, 4, 9, 10, 16, 18, 19, 24, 25, 27, 33, 35, 46, 54, 55, 58, 59, 60, 65, 67, 69, 72, 73, 76, 80, 82, 84, 91, 104, 109, 110, 111, 114, 115, 118, 121, 124, 130, 131, 132, 135, 136, 141, 147, 157, 158, 162, 164, 166, 169, 170, 172, 175, 185, 187, 198, 203, 205, 211, 216, 217, 218, 221, 224, 227, 228, 231, 240, 242, 246, 247, 270, 273, 276, 289, 296, 302, 313, 329, 333, 336, 344, 348, 369\}$

Example(6): Let C be a $[123, 4, 97]_7$ –code with generator matrix

G=

$$\begin{bmatrix} 1 & 0 & 0 & 2 & 1 & 3 & 0 & 2 & 1 & 3 & 2 & 1 & 3 & 4 & 6 & 5 & 6 & 0 & 3 & 5 & 5 & 0 & 4 & 1 & 1 & 0 & 4 & 5 & 2 & 4 & 1 & 0 & 4 & 5 \\ 0 & 1 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 4 & 5 & 1 & 2 & 4 & 6 & 1 & 6 & 4 & 1 & 0 & 1 & 2 & 3 & 5 & 6 & 0 & 1 & 2 & 3 \end{bmatrix}$$

0 0 1 1 1 1 0 0 0 0 1 1 1 1 1 1 2 2 2 2 3 3 6 0 1 1 1 1 1 1 0 0 0 0
 0 0 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 0 0 0 0 1 1 1 1

4 3 4 0 2 1 2 3 2 1 6 4 2 5 2 0 6 5 5 3 2 2 3 3 2 4 0 1 5 4 0 3 0 1 6 1 5
 6 0 1 2 3 0 0 1 2 4 0 1 2 2 3 5 5 2 4 2 1 2 1 6 0 1 2 3 4 0 1 3 4 4 0 1 2
 0 1 1 1 1 2 2 2 2 2 3 3 3 3 3 3 3 4 5 6 0 1 1 1 0 0 0 0 0 1 1 1 1 1 2 2 2
 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1

6 0 0 6 1 5 2 2 1 0 0 6 3 5 2 6 3 3 1 4 3 1 3 0 2 3 4 0 3 1 4 5 6 5
 3 6 0 1 3 3 5 2 3 5 6 6 1 3 4 5 1 2 3 5 0 1 2 3 4 4 4 0 1 2 2 4 5 1
 2 2 3 3 3 3 3 4 4 4 4 4 5 5 5 5 0 1 1 1 0 0 0 0 0 0 0 0 1 1 1 1 1 1 2
 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1

4 4 2 1 1 4 0 1 6 1 4 2 3 6 2 4 0 2
 2 4 5 6 0 0 1 2 3 6 0 1 4 4 5 5 3 6
 2 2 2 2 3 3 3 3 3 3 4 4 4 4 4 4 5 5
 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$
 and p_j is defined as follows $\forall_i: \Sigma = \{[1,1,1,1], \dots, [0,0,0,0]\}$. Therefore, $G = [p_{343}^T, p_{2107}^T, \dots, p_{617}^T]$
 The code words is generated by multiplying each point in Σ to G .

Therefore, the code words in $[123, 4, 97]_7$ linear code are defined as follows:

$$CW = \{ [p_i \times G]_{q^k \times n} \mid p_i \in \Sigma \} = \begin{bmatrix} p_{11} & \cdots & p_{1n} \\ \vdots & \ddots & \vdots \\ p_{m1} & \cdots & p_{mn} \end{bmatrix}.$$

with **w.d.:**

$0^1, 97^{12}, 98^6, 99^{24}, 100^{36}, 101^{96}, 102^{210}, 103^{270}, 104^{336}, 105^{324}, 106^{282}, 107^{270}, 108^{216}, 109^{90}, 110^{60}, 111^{48}, 112^{36},$
 $113^{24}, 114^{24}, 115^{12}, 116^{12}, 117^6, 119^6.$

Then we have:

$C_1 = \{1, 2, 3, 4, 5, 9, 10, 16, 18, 19, 24, 25, 26, 27, 31, 33, 35, 46, 48, 54, 55, 58, 59, 60, 61, 65, 66, 67, 69, 72, 73,$
 $75, 76, 79, 80, 82, 84, 88, 89, 90, 91, 104, 107, 109, 110, 111, 114, 115, 117, 118, 121, 122, 124, 125, 130, 131, 132,$
 $135, 136, 140, 141, 147, 148, 157, 158, 162, 164, 166, 168, 169, 170, 172, 174, 175, 183, 185, 187,$
 $188, 193, 198, 199, 203, 205, 206, 209, 211, 212, 216, 217, 218, 220, 221, 224, 227, 228, 231,$
 $232, 240, 242, 246, 247, 248, 258, 263, 270, 273, 276, 285, 288, 289, 291, 293, 296, 302, 313,$
 $324, 329, 333, 336, 344, 347, 348, 369\}.$

Example(7): Let C be a $[153, 4, 121]_7$ -code with generator matrix

$G =$

1 0 0 2 1 3 0 2 1 3 2 1 3 4 6 5 6 0 3 5 5 0 4 1 1 0 4 5 2 4 1
 0 1 0 1 2 3 0 1 2 3 0 1 2 3 4 5 1 2 4 6 1 6 4 1 0 1 2 3 5 6 0
 0 0 1 1 1 1 0 0 0 0 1 1 1 1 1 1 2 2 2 2 3 3 6 0 1 1 1 1 1 1 0
 0 0 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 0 0 0 1

0 4 5 4 3 4 0 2 1 2 3 2 1 6 4 2 5 2 0 6 5 5 3 2 2 3 3 2 4 0 1 5 4
 1 2 3 6 0 1 2 3 0 0 1 2 4 0 1 2 2 3 5 5 2 4 2 1 2 1 6 0 1 2 3 4 0
 0 0 0 0 1 1 1 1 2 2 2 2 2 3 3 3 3 3 3 3 4 5 6 0 1 1 1 0 0 0 0 0 1
 1 0 0 0 0 1 1 1 1 1 1

0 3 0 1 6 1 5 6 0 0 6 1 5 2 2 1 0 0 6 3 5 2 6 3 3 1 4 3 1 3
 1 3 4 4 0 1 2 3 6 0 1 3 3 5 2 3 5 6 6 1 3 4 5 1 2 3 5 0 1 2
 1 1 1 1 2 2 2 2 2 3 3 3 3 3 4 4 4 4 4 5 5 5 5 0 1 1 1 0 0 0
 1 0 0 0 0 1 1 1

0 2 3 4 0 3 1 4 5 6 5 4 4 2 1 1 4 0 1 6 1 4 2 3 6 2 4 0 2 4 6 5
 3 4 4 4 0 1 2 2 4 5 1 2 4 5 6 0 0 1 2 3 6 0 1 4 4 5 5 3 6 1 1 1
 0 0 0 0 1 1 1 1 1 1 1 2 2 2 2 2 3 3 3 3 3 3 4 4 4 4 4 5 5 0 1 0
 1 0 0 1

4 1 0 2 5 1 6 4 3 5 3 6 3 1 4 6 5 2 4 3 5 4 6 0 1 6 5
 3 4 5 6 2 3 3 5 3 4 5 5 1 4 4 4 5 6 4 5 6 2 2 6 6 6 6
 0 0 0 0 1 1 1 1 2 2 2 2 3 3 3 3 3 3 4 4 4 5 5 5 5 5 6
 1

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$ and p_j is defined as follows $\forall_i: \Sigma = \{[1,1,1,1], \dots, [0,0,0,0]\}$, Therefore, $G = [p_{343}^T, p_{2107}^T, \dots, p_{1653}^T]$

The code words is generated by multiplying each point in Σ to G

Therefore the code words in $[153, 4, 121]_7$ linear code are defined as follows:

$$CW = \{ [p_i \times G]_{q^k \times n} \mid p_i \in \Sigma \} = \begin{bmatrix} p_{11} & \cdots & p_{1n} \\ \vdots & \ddots & \vdots \\ p_{m1} & \cdots & p_{mn} \end{bmatrix}.$$

with **w.d.:**

$0^1, 121^{12}, 122^6, 124^{18}, 125^{30}, 126^{66}, 127^{90}, 128^{210}, 129^{258}, 130^{396}, 131^{354}, 132^{330}, 133^{180}, 134^{90}, 135^{108}, 136^{96}, 137^{42}, 138^{24}, 139^{18}, 140^{30}, 141^{18}, 143^{12}, 144^{12}.$

Then we have:

$C_1 = \{1, 2, 3, 4, 5, 6, 9, 10, 16, 18, 19, 22, 24, 25, 26, 27, 31, 33, 35, 46, 48, 54, 55, 58, 59, 60, 61, 65, 66, 67, 69, 70, 72, 73, 75, 76, 79, 80, 82, 84, 87, 88, 89, 90, 91, 93, 102, 104, 107, 109, 110, 111, 114, 115, 117, 118, 121, 122, 124, 125, 126, 129, 130, 131, 132, 134, 135, 136, 140, 141, 146, 147, 148, 157, 158, 162, 164, 166, 168, 169, 170, 172, 174, 175, 180, 183, 185, 187, 188, 189, 193, 194, 197, 198, 199, 203, 205, 206, 209, 211, 212, 215, 216, 217, 218, 220, 221, 224, 227, 228, 231, 232, 234, 237, 239, 240, 242, 245, 246, 247, 248, 249, 258, 263, 270, 273, 276, 285, 286, 288, 289, 291, 292, 293, 296, 301, 302, 313, 323, 324, 329, 333, 336, 344, 345, 346, 347, 348, 351, 369, 399\}$

Example(8): Let C be a $[182, 4, 143]_7$ -code with generator matrix

$G =$

1 0 0 2 1 3 0 2 1 3 2 1 3 4 6 5 6 0 3 5 5 0 4 1 1 0 4 5 2 4 1 0 4 5 4 3 4 0
 0 1 0 1 2 3 0 1 2 3 0 1 2 3 4 5 1 2 4 6 1 6 4 1 0 1 2 3 5 6 0 1 2 3 6 0 1 2
 0 0 1 1 1 1 0 0 0 0 1 1 1 1 1 1 2 2 2 2 3 3 6 0 1 1 1 1 1 1 0 0 0 0 1 1 1
 0 0 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 0 0 0 0 1 1 1 1 1 1 1

2 1 2 3 2 1 6 4 2 5 2 0 6 5 5 3 2 2 3 3 2 4 0 1 5 4 0 3 0 1 6 1 5 6 0 0 6 1 5
 3 0 0 1 2 4 0 1 2 2 3 5 5 2 4 2 1 2 1 6 0 1 2 3 4 0 1 3 4 4 0 1 2 3 6 0 1 3 3
 1 2 2 2 2 2 3 3 3 3 3 3 3 4 5 6 0 1 1 1 0 0 0 0 0 1 1 1 1 1 2 2 2 2 2 3 3 3 3
 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1

2 2 1 0 0 6 3 5 2 6 3 3 1 4 3 1 3 0 2 3 4 0 3 1 4 5 6 5 4 4 2 1 1 4 0 1 6 1 4
 5 2 3 5 6 6 1 3 4 5 1 2 3 5 0 1 2 3 4 4 4 0 1 2 2 4 5 1 2 4 5 6 0 0 1 2 3 6 0
 3 4 4 4 4 4 5 5 5 5 0 1 1 1 0 0 0 0 0 0 0 0 1 1 1 1 1 1 2 2 2 2 2 3 3 3 3 3 3 4
 1 1 1 1 1 1 1 1 1 1 0 0 0 0 1

2 3 6 2 4 0 2 4 6 5 4 1 0 2 5 1 6 4 3 5 3 6 3 1 4 6 5 2 4 3 5 4 6 0 1 6 5 5 5
 1 4 4 5 5 3 6 1 1 1 3 4 5 6 2 3 3 5 3 4 5 5 1 4 4 4 5 6 4 5 6 2 2 6 6 6 6 1 2
 4 4 4 4 4 5 5 0 1 0 0 0 0 0 1 1 1 1 2 2 2 2 3 3 3 3 3 3 4 4 4 5 5 5 5 5 6 0 1
 1 1 1 1 1 1 1 0 0 1 0 0

1 5 3 5 2 0 2 1 3 4 1 2 5 2 0 3 1 4 5 0 5 4 3 6 6 6 2
 6 6 6 0 2 3 5 6 0 0 2 4 5 1 2 2 5 6 6 4 5 6 3 3 0 1 3
 1 1 0 1 1 1 1 1 2 2 2 2 2 3 3 3 3 3 3 4 4 4 5 5 6 6 6
 0 0 1

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$ and p_j is defined as follows $\forall_i: \Sigma = \{[1,1,1,1], \dots, [0,0,0,0]\}$. Therefore, $G = [p_{343}^T, p_{2107}^T, \dots, p_{477}^T]$

The code words is generated by multiplying each point in Σ to G .

Therefore, the code words in $[182, 4, 143]_7$ linear code are defined as follows:

$$CW = \{ [p_i \times G]_{q^k \times n} \mid p_i \in \Sigma \} = \begin{bmatrix} p_{11} & \cdots & p_{1n} \\ \vdots & \ddots & \vdots \\ p_{m1} & \cdots & p_{mn} \end{bmatrix}.$$

with **w.d.:**

$0^1, 143^6, 145^6, 147^6, 148^{18}, 150^{36}, 151^{66}, 152^{180}, 153^{186}, 154^{258}, 155^{288}, 156^{390}, 157^{324}, 158^{186}, 159^{120}, 160^{72}, 161^{96}, 162^{54}, 163^{36}, 164^{18}, 165^{12}, 166^{18}, 167^{12}, 169^6, 170^6$.

Then we have:

$C_1 = \{1, 2, 3, 4, 5, 6, 7, 9, 10, 16, 18, 19, 22, 24, 25, 26, 27, 28, 31, 33, 35, 46, 48, 52, 54, 55, 56, 58, 59, 60, 61, 65, 66, 67, 69, 70, 72, 73, 75, 76, 79, 80, 82, 84, 87, 88, 89, 90, 91, 93, 102, 103, 104, 107, 109, 110, 111, 112, 114, 115, 117, 118, 121, 122, 123, 124, 125, 126, 128, 129, 130, 131, 132, 134, 135, 136, 140, 141, 144, 146, 147, 148, 150, 157, 158, 159, 160, 162, 164, 166, 168, 169, 170, 171, 172, 174, 175, 180, 183, 185, 186, 187, 188, 189, 193, 194, 196, 197, 198, 199, 203, 205, 206, 209, 211, 212, 214, 215, 216, 217, 218, 219, 220, 221, 222, 224, 227, 228, 231, 232, 234, 237, 239, 240, 241, 242, 245, 246, 247, 248, 249, 251, 252, 258, 263, 270, 273, 276, 282, 285, 286, 288, 289, 291, 292, 293, 294, 296, 300, 301, 302, 313, 321, 323, 324, 327, 329, 330, 333, 336, 344, 345, 346, 347, 348, 351, 358, 365, 369, 375, 399\}$.

Example(9): Let C be a $[222, 4, 179]_7$ -code with generator matrix

$G =$

1 0 0 2 1 3 0 2 1 3 2 1 3 4 6 5 6 0 3 5 5 0 4 1 1 0 4 5 2 4 1 0 4 5 4 3
 0 1 0 1 2 3 0 1 2 3 0 1 2 3 4 5 1 2 4 6 1 6 4 1 0 1 2 3 5 6 0 1 2 3 6 0
 0 0 1 1 1 1 0 0 0 0 1 1 1 1 1 1 2 2 2 2 3 3 6 0 1 1 1 1 1 1 0 0 0 0 0 1
 0 0 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 0 0 0 0 1 1 1 1 1 1

4 0 2 1 2 3 2 1 6 4 2 5 2 0 6 5 5 3 2 2 3 3 2 4 0 1 5 4 0 3 0 1 6 1 5 6 0
 1 2 3 0 0 1 2 4 0 1 2 2 3 5 5 2 4 2 1 2 1 6 0 1 2 3 4 0 1 3 4 4 0 1 2 3 6
 1 1 1 2 2 2 2 2 3 3 3 3 3 3 3 4 5 6 0 1 1 1 0 0 0 0 0 1 1 1 1 1 2 2 2 2
 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1

0 6 1 5 2 2 1 0 0 6 3 5 2 6 3 3 1 4 3 1 3 0 2 3 4 0 3 1 4 5 6 5 4 4 2 1 1 4
 0 1 3 3 5 2 3 5 6 6 1 3 4 5 1 2 3 5 0 1 2 3 4 4 4 0 1 2 2 4 5 1 2 4 5 6 0 0
 3 3 3 3 3 4 4 4 4 4 5 5 5 5 0 1 1 1 0 0 0 0 0 0 0 1 1 1 1 1 1 2 2 2 2 2 3 3
 1 1 1 1 1 1 1 1 1 1 1 1 1 1 0 0 0 0 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1

 0 1 6 1 4 2 3 6 2 4 0 2 4 6 5 4 1 0 2 5 1 6 4 3 5 3 6 3 1 4 6 5 2 4 3 5 4 6 0 1
 1 2 3 6 0 1 4 4 5 5 3 6 1 1 1 3 4 5 6 2 3 3 5 3 4 5 5 1 4 4 4 5 6 4 5 6 2 2 6 6
 3 3 3 3 4 4 4 4 4 4 5 5 0 1 0 0 0 0 0 1 1 1 1 2 2 2 2 3 3 3 3 3 3 4 4 4 5 5 5 5
 1 1 1 1 1 1 1 1 1 1 1 1 0 0 1

 6 5 5 5 1 5 3 5 2 0 2 1 3 4 1 2 5 2 0 3 1 4 5 0 5 4 3 6 6 6 2 6 2 1 0 2 6 2 5 3
 6 6 1 2 6 6 6 0 2 3 5 6 0 0 2 4 5 1 2 2 5 6 6 4 5 6 3 3 0 1 3 1 0 1 2 3 3 4 4 1
 5 6 0 1 1 1 0 1 1 1 1 1 2 2 2 2 2 3 3 3 3 3 3 4 4 4 5 5 6 6 6 0 1 1 1 1 1 1 1 0
 1 1 0 0 0 0 1

 3 5 2 1 2 5 3 1 0 3 0 5 6 5 2 6 4 0 6 3 4 1 6 3 0 4 0 4 5 4 6
 5 2 3 5 1 1 4 5 1 2 3 3 6 0 4 6 1 2 2 3 3 4 5 6 1 3 4 0 4 6 6
 1 0 0 0 1 1 1 1 2 2 2 2 2 3 3 3 4 4 4 4 4 4 4 5 5 5 6 6 6 6
 0 1

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$ and p_j is defined as follows $\forall_i: \Sigma = \{[1,1,1,1], \dots, [0,0,0,0]\}$. Therefore, $G = [p_{343}^T, p_{2107}^T, \dots, p_{1996}^T]$.

The code words is generated by multiplying each point in Σ to G .

Therefore, the code words in $[222, 4, 179]_7$ linear code are defined as follows:

$$CW = \{ [p_i \times G]_{q^{k \times n}} \mid p_i \in \Sigma \} = \begin{bmatrix} p_{11} & \cdots & p_{1n} \\ \vdots & \ddots & \vdots \\ p_{m1} & \cdots & p_{mn} \end{bmatrix}.$$

with **w.d.:**

$0^1, 179^6, 180^6, 181^6, 183^6, 184^{54}, 185^{60}, 186^{102}, 187^{162}, 188^{258}, 189^{408}, 190^{288}, 191^{270}, 192^{252}, 193^{144}, 194^{96}, 195^9, 196^{66}, 197^{60}, 198^{30}, 199^6, 200^6, 201^6, 204^6, 205^{12}.$

Then we have:

$C_1 = \{1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 16, 17, 18, 19, 22, 23, 24, 25, 26, 27, 28, 31, 32, 33, 35, 39, 42, 46, 48, 52, 54, 55, 56, 58, 59, 60, 61, 65, 66, 67, 68, 69, 70, 72, 73, 74, 75, 76, 79, 80, 81, 82, 83, 84, 87, 88, 89, 90, 91, 93, 102, 103, 104, 107, 109, 110, 111, 112, 114, 115, 116, 117, 118, 119, 121, 122, 123, 124, 125, 126, 128, 129, 130, 131, 132, 134, 135, 136, 138, 140, 141, 143, 144, 146, 147, 148, 150, 157, 158, 159, 160, 162, 163, 164, 166, 168, 169, 170, 171, 172, 173, 174, 175, 177, 180, 182, 183, 185, 186, 187, 188, 189, 193, 194, 196, 197, 198, 199, 203, 204, 205, 206, 209, 210, 211, 212, 214, 215, 216, 217, 218, 219, 220, 221, 222, 22A, 227, 228, 231, 232, 234, 237, 239, 240, 241, 242, 245, 246, 247, 248, 249, 251, 252, 253, 258, 263, 265, 268, 270, 273, 274, 276, 278, 279, 282, 283, 285, 286, 288, 289, 291, 292, 293, 294, 295, 296, 299, 300, 301, 302, 310, 313, 321, 323, 324, 327, 328, 329, 330, 331, 333, 336, 344, 345, 346, 347, 348, 351, 356, 358, 365, 369, 375, 385, 398, 399, 400\}.$

Example (10):

We take the complete $(k, 2) - \text{cap} \subset \Sigma = PG(3, q)$

Such that, the number of the complete $(k, 2) - \text{cap}$ in $PG(3, 7)$ is 23, which are:

$[(1,0,0,0), (0,1,0,0), (0,0,1,0), (2,1,1,0), (1,2,1,0), (3,3,1,0), (0,0,0,1), (2,1,0,1), (1,2,0,1), (3,3,0,1), (2,0,1,1)]$,
 ,

$(1,1,1,1),(3,2,1,1),(4,3,1,1),(6,4,1,1),(5,5,1,1),(6,1,2,1),(0,2,2,1),(3,4,2,1),(5,6,2,1),(5,1,3,1),(0,6,3,1),(4,4,6,1)]$.

we take a parameter (n) from the number of $(k, 2)$ –cap, then $n_q(k, d) = 23$, and we get C is a $[15, 4]$ – code, then the generator matrix G defined as follows:

$G = \{ \text{The columns is the transpose of cartesian points of } (23, 2) \text{ –cap in } PG(3, 7) \subset \Sigma = \{ P : P = \text{mod}_{13}(\text{permutations}(7, 4)) \}$, then ; $G_{23 \times 4} =$

$$\begin{bmatrix} 1 & 0 & 0 & 2 & 1 & 3 & 0 & 2 & 1 & 3 & 2 & 1 & 3 & 4 & 6 & 5 & 6 & 0 & 3 & 5 & 5 & 0 & 4 \\ 0 & 1 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 0 & 1 & 2 & 3 & 4 & 5 & 1 & 2 & 4 & 6 & 1 & 6 & 4 \\ 0 & 0 & 1 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 2 & 2 & 2 & 2 & 3 & 3 & 6 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \end{bmatrix}$$

We note that the minimum distance d is 15, then we get C is a $[23, 4, 15]_7$ –code, such that by **(Theorem (3.3))** we see that $\geq d + K - 1, 23 \geq 15 + 4 - 1$.

In this example $\Sigma = \{ p_1, p_2, \dots, p_m \}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{ [1, 1, 1, 1], \dots, [0, 0, 0, 0] \}$. Therefore, $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to. Any permutation of the rows of C or multiplication of a row of C

$M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[23, 4, 15]_7$ linear code is defined as follows:

$$Cw = \{ [p_i \in \Sigma] = G^*_{2401 \times 23}$$

Then C is Grasmere with **w.d.:**

$$0^1, 15^6, 16^{12}, 17^{150}, 18^{419}, 19^{606}, 20^{432}, 21^{402}, 22^{240}, 23^{138}.$$

The code $[23, 4, 15]_7$ is Optimal.

By Theorem (2.6.4) we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[23, 4, 15]_7$ – code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{15}{7^0} \rceil + \lceil \frac{15}{7^1} \rceil + \lceil \frac{15}{7^2} \rceil + \lceil \frac{15}{7^3} \rceil \\ &= \lceil \frac{15}{1} \rceil + \lceil \frac{15}{7} \rceil + \lceil \frac{15}{49} \rceil + \lceil \frac{15}{343} \rceil \\ &= 15 + 3 + 1 + 1 \cong 20. \end{aligned}$$

Then, C is Optimal linear codes.

Example (11):

We take the complete $(k, 3)$ –cap $\subset \Sigma = PG(3, q)$ Such that,

the number of the complete $(k, 3)$ –cap in $PG(3, 7)$ is 54, we take a parameter n from the number of $(k, 3)$ –cap, then we take $n_q(k, d) = 54$, and we get C is a $[41, 4]_7$ – code, with generator matrix G

defined as follows: $G = \{ \text{The columns is the transpose of cartesian points of } (54, 3) \text{ –cap in } PG(3, 7) \subset \Sigma = \{ P : P = \text{mod}_7(\text{permutations}(7, 4)) \}$, then we get $G_{54 \times 4}$.

Now, by computer Program, we found that the minimum distance d is 41, then we get C is a $[54, 4, 41]_7$ – code. In this example

In this example $\Sigma = \{ p_1, p_2, \dots, p_m \}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{ [1, 1, 1, 1], \dots, [0, 0, 0, 0] \}$, Therefore, $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to.

Any permutation of the rows of C or multiplication of a row of C

$M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[54, 4, 41]_7$ linear code is defined as follows:

$$Cw = \{ [p_i \in \Sigma] = G^*_{2401 \times 54}$$

then C is Grasmere with **w.d.:**

$$0^1, 41^{24}, 42^{42}, 43^{168}, 44^{342}, 45^{396}, 46^{378}, 47^{372}, 48^{270}, 49^{210}, 50^{78}, 51^{36}, 52^{36}, 53^{30}, 54^{18}.$$

The code $[54, 4, 41]_7$ is Optimal

By **Theorem (2.6.4)** we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[23, 4, 15]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{41}{7^0} \rceil + \lceil \frac{41}{7^1} \rceil + \lceil \frac{41}{7^2} \rceil + \lceil \frac{41}{7^3} \rceil \\ &= \lceil \frac{41}{1} \rceil + \lceil \frac{41}{7} \rceil + \lceil \frac{41}{49} \rceil + \lceil \frac{41}{343} \rceil \\ &= 41 + 6 + 1 + 1 \cong 49. \end{aligned}$$

Then C is Optimal linear codes.

Example (12):

We take the complete $(k, 4)$ -cap $\subset \Sigma = PG(3, q)$

Such that, the number of the complete $(k, 4)$ -cap in $PG(3, 7)$ is 87, we take a parameter n from the number of $(k, 4)$ -cap then we take $n_q(k, d) = 87$, and we get C is a $[68, 4]_7$ - code, with generator matrix G defined as follows: $G = \{ \text{The columns is the transpose of Cartesian points of } (87, 4) \text{ -cap in } PG(3, 7) \subset \Sigma = \{ P : P = \text{mod}_7(\text{permutations } (7, 4)) \}$, then we get $G_{87 \times 4}$.

Now, by computer Program, we found that the minimum distance d is 68, then we get C is a $[87, 4, 68]_7$ - code.

In this example $\Sigma = \{ p_1, p_2, \dots, p_m \}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{ [1, 1, 1, 1], \dots, [0, 0, 0, 0] \}$, therefore $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to.

Any permutation of the rows of C or multiplication of a row of C

$M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[54, 4, 41]_7$ linear code is defined as follows:

$Cw = \{ [p_i \in \Sigma] = G^*_{2401 \times 86}$ then C is Grasmere with **w.d.:**

$0^1, 68^6, 69^{48}, 70^{84}, 71^{198}, 72^{216}, 73^{342}, 74^{336}, 75^{336}, 76^{282}, 77^{198}, 78^{168}, 79^{66}, 80^{30}, 81^{36}, 82^{30}, 84^{12}, 85^6, 86^6$.

The code $[87, 4, 68]_7$ is Optimal

By **Theorem (2.6.4)** we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[87, 4, 68]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{68}{7^0} \rceil + \lceil \frac{68}{7^1} \rceil + \lceil \frac{68}{7^2} \rceil + \lceil \frac{68}{7^3} \rceil \\ &= \lceil \frac{68}{1} \rceil + \lceil \frac{68}{7} \rceil + \lceil \frac{68}{49} \rceil + \lceil \frac{68}{343} \rceil \\ &= 68 + 10 + 2 + 1 \cong 81 \end{aligned}$$

Then C is Optimal linear codes.

Example (13):

We take the complete $(k, 5)$ -cap $\subset \Sigma = PG(3, q)$ Such that, the number of the complete $(k, 5)$ -cap in $PG(3, 7)$ is 123, we take a parameter n from the number of $(k, 5)$ -cap then we take $n_q(k, d) = 123$, and we get C is a $[97, 4]_7$ -code, with generator matrix G defined as follows: $G = \{ \text{The columns is the transpose of cartesian points of } (123, 5) \text{ -cap in } PG(3, 7) \subset \Sigma = \{ P : P = \text{mod}_7(\text{permutations } (7, 4)) \}$, then we get $G_{123 \times 4}$.

Now, by computer Program, we found that the minimum distance d is 97, then we get C is a $[123, 4, 97]_7$ - code.

In this example $\Sigma = \{ p_1, p_2, \dots, p_m \}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{ [1, 1, 1, 1], \dots, [0, 0, 0, 0] \}$. Therefore, $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to.

Any permutation of the rows of C or multiplication of a row of C

$M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[123, 4, 97]_7$ linear code is defined as follows:

$Cw = \{[p_i \in \Sigma] = G^*_{2401 \times 123}$ then C is Grasmere with **w.d.**:

$0^1, 97^{12}, 98^6, 99^{24}, 100^{36}, 101^{96}, 102^{210}, 103^{270}, 104^{336}, 105^{324}, 106^{282}, 107^{270}, 108^{216}, 109^{90}, 110^{60}, 111^{48}, 112^{36}, 113^{24}, 114^{24}, 115^{12}, 116^{12}, 117^6, 119^6$.

The code $[123, 4, 97]_7$ is Optimal

By **Theorem (2.6.4)** we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[123, 4, 97]_7 - \text{cod}$, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{97}{7^0} \rceil + \lceil \frac{97}{7^1} \rceil + \lceil \frac{97}{7^2} \rceil + \lceil \frac{97}{7^3} \rceil \\ &= \lceil \frac{97}{1} \rceil + \lceil \frac{97}{7} \rceil + \lceil \frac{97}{49} \rceil + \lceil \frac{97}{343} \rceil \\ &= 97 + 14 + 2 + 1 \cong 114. \end{aligned}$$

Then C is Optimal linear codes.

Example (14):

We take the complete $(k, 6) - \text{cap} \subset \Sigma = PG(3, q)$ Such that, the number of the complete $(k, 6) - \text{cap}$ in $PG(3, 7)$ is 153, we take a parameter n from the number of $(k, 6) - \text{cap}$ then we take $n_q(k, d) = 153$, and we get C is a $[121, 4]_7 - \text{code}$, with generator matrix G defined as follows: $G = \{ \text{The columns is the transpose of cartesian points of } (153, 6) - \text{cap in } PG(3, 7) \subset \Sigma = \{P: P = \text{mod}_7(\text{permutations}(7, 4))\}$, then we get $G_{153 \times 4}$.

Now, by computer Program, we found that the minimum distance d is 121,

then we get C is a $[153, 4, 121]_7$ code.

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{[1, 1, 1, 1], \dots, [0, 0, 0, 0]\}$, therefore $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to.

Any permutation of the rows of C or multiplication of a row of C $M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[153, 4, 121]_7$ linear code is defined as follows:

$Cw = \{[p_i \in \Sigma] = G^*_{2401 \times 153}$ then C is Grasmere with **w.d.**:

$0^1, 121^{12}, 122^6, 124^{18}, 125^{30}, 126^{66}, 127^{90}, 128^{210}, 129^{258}, 130^{396}, 131^{354}, 132^{330}, 133^{180}, 134^{90}, 135^{108}, 136^{96}, 137^{42}, 138^{24}, 139^{18}, 140^{30}, 141^{18}, 143^{12}, 144^{12}$.

The code $[153, 4, 121]_7$ is Optimal

By **Theorem (2.6.4)** we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[153, 4, 121]_7 - \text{code}$, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{121}{7^0} \rceil + \lceil \frac{121}{7^1} \rceil + \lceil \frac{121}{7^2} \rceil + \lceil \frac{121}{7^3} \rceil \\ &= \lceil \frac{121}{1} \rceil + \lceil \frac{121}{7} \rceil + \lceil \frac{121}{49} \rceil + \lceil \frac{121}{343} \rceil \\ &= 121 + 18 + 3 + 1 \cong 143 \end{aligned}$$

Then C is Optimal linear codes.

Example (15):

We take the complete $(k, 7) - \text{cap} \subset \Sigma = PG(3, q)$ Such that, the number of the complete $(k, 7) - \text{cap}$ in $PG(3, 7)$ is 182, we take a parameter n from the number of $(k, 7) - \text{cap}$ then we take $n_q(k, d) = 182$, and we get C is a $[143, 4]_7 - \text{code}$, with generator matrix G defined as follows:

$G = \{ \text{The columns is the transpose of Cartesian points of } (182, 7) - \text{cap in } PG(3, 7) \subset \Sigma = \{P: P = \text{mod}_7(\text{permutations}(7, 4))\}$, then we get $G_{182 \times 4}$.

Now, by computer Program, we found that the minimum distance d is 143, then we get C is a $[182, 4, 143]_7$ code.

In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{[1, 1, 1, 1], \dots, [0, 0, 0, 0]\}$, therefore $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to.

Any permutation of the rows of C or multiplication of a row of C

$M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[182, 4, 143]_7$

linear code is defined as follows:

$Cw = \{[p_i \in \Sigma] = G^*_{2401 \times 182}$ then C is Grasmere with **w.d.**:

$0^1, 143^6, 145^6, 147^6, 148^{18}, 150^{36}, 151^{66}, 152^{180}, 153^{186}, 154^{258}, 155^{288}, 156^{390}, 157^{324}, 158^{186}, 159^{120}, 160^{72}, 161^{96}, 162^{54}, 163^{36}, 164^{18}, 165^{12}, 166^{18}, 167^{12}, 169^6, 170^6$.

The code $[182, 4, 143]_7$ is Optimal

By **Theorem (2.6.4)** we get $n \geq g_q(K, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[182, 4, 143]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{143}{7^0} \rceil + \lceil \frac{143}{7^1} \rceil + \lceil \frac{143}{7^2} \rceil + \lceil \frac{143}{7^3} \rceil \\ &= \lceil \frac{143}{1} \rceil + \lceil \frac{143}{7} \rceil + \lceil \frac{143}{49} \rceil + \lceil \frac{143}{343} \rceil \\ &= 143 + 21 + 3 + 1 \cong 168 \end{aligned}$$

Then C is Optimal linear codes.

Example (16):

We take the complete $(k, 8)$ -cap $\subset \Sigma = PG(3, q)$ Such that, the number of the complete $(k, 8)$ -cap in $PG(3, 7)$ is 222, we take a parameter n from the number of $(K, 8)$ -cap then we take

$n_q(k, d) = 222$, and we get C is a $[41, 4]_7$ - code, with generator matrix G defined as follows: $G = \{$

The columns is the transpose of cartesian points of $(222, 8)$ -cap in $PG(3, 7) \subset \Sigma = \{P: P = \text{mod}_7(\text{permutations}(7, 4))\}$, then we get $G_{222 \times 4}$.

Now, by computer Program, we found that the minimum distance d is 179, then we get C is a $[222, 4, 179]_7$ code. In this example $\Sigma = \{p_1, p_2, \dots, p_m\}$ where $m = q^k = (7)^4 = 2401$,

and p_i is defined as follows $\forall_i, \Sigma = \{[1, 1, 1, 1], \dots, [0, 0, 0, 0]\}$, therefore $G = [p_1^T, \dots, p_m^T]$.

The code words is generated by multiplying each point in Σ to.

Any permutation of the rows of C or multiplication of a row of C

$M(C)$ by an element of F_7 gives another matrix of G , there for the code words in $[222, 4, 179]_7$

linear code is defined as follows:

$Cw = \{[p_i \in \Sigma] = G^*_{2401 \times 222}$

then C is Grasmere with **w.d.**:

$0^1, 179^6, 180^6, 181^6, 183^6, 184^{54}, 185^{60}, 186^{102}, 187^{162}, 188^{258}, 189^{408}, 190^{288}, 191^{270}, 192^{252}, 193^{144}, 194^{96}, 195^9, 196^{66}, 197^{60}, 198^{30}, 199^6, 200^6, 201^6, 204^6, 205^{12}$.

The code $[222, 4, 179]_7$ is Optimal

By **Theorem (2.6.4)** we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[222, 4, 179]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{179}{7^0} \rceil + \lceil \frac{179}{7^1} \rceil + \lceil \frac{179}{7^2} \rceil + \lceil \frac{179}{7^3} \rceil \\ &= \lceil \frac{179}{1} \rceil + \lceil \frac{179}{7} \rceil + \lceil \frac{179}{49} \rceil + \lceil \frac{179}{343} \rceil \\ &= 179 + 26 + 4 + 1 \cong 210 \end{aligned}$$

Then C is Optimal linear codes.

3. Results

1. The code $[23, 4, 15]_7$ is Optimal.

we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[23, 4, 15]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{15}{7^0} \rceil + \lceil \frac{15}{7^1} \rceil + \lceil \frac{15}{7^2} \rceil + \lceil \frac{15}{7^3} \rceil \\ &= \lceil \frac{15}{1} \rceil + \lceil \frac{15}{7} \rceil + \lceil \frac{15}{49} \rceil + \lceil \frac{15}{343} \rceil \\ &= 15 + 3 + 1 + 1 \cong 20. \end{aligned}$$

Then, C is Optimal linear codes.

2. The code $[54, 4, 41]_7$ is Optimal

) we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[54, 4, 41]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{41}{7^0} \rceil + \lceil \frac{41}{7^1} \rceil + \lceil \frac{41}{7^2} \rceil + \lceil \frac{41}{7^3} \rceil \\ &= \lceil \frac{41}{1} \rceil + \lceil \frac{41}{7} \rceil + \lceil \frac{41}{49} \rceil + \lceil \frac{41}{343} \rceil \\ &= 41 + 6 + 1 + 1 \cong 49. \end{aligned}$$

Then C is Optimal linear codes.

3. The code $[87, 4, 68]_7$ is Optimal

we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[87, 4, 68]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{68}{7^0} \rceil + \lceil \frac{68}{7^1} \rceil + \lceil \frac{68}{7^2} \rceil + \lceil \frac{68}{7^3} \rceil \\ &= \lceil \frac{68}{1} \rceil + \lceil \frac{68}{7} \rceil + \lceil \frac{68}{49} \rceil + \lceil \frac{68}{343} \rceil \\ &= 68 + 10 + 2 + 1 \cong 81 \end{aligned}$$

Then C is Optimal linear codes.

4. The code $[123, 4, 97]_7$ is Optimal

we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[123, 4, 97]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{97}{7^0} \rceil + \lceil \frac{97}{7^1} \rceil + \lceil \frac{97}{7^2} \rceil + \lceil \frac{97}{7^3} \rceil \\ &= \lceil \frac{97}{1} \rceil + \lceil \frac{97}{7} \rceil + \lceil \frac{97}{49} \rceil + \lceil \frac{97}{343} \rceil \\ &= 97 + 14 + 2 + 1 \cong 114. \end{aligned}$$

Then C is Optimal linear codes.

5. The code $[153, 4, 121]_7$ is Optimal

we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[153, 4, 121]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{121}{7^0} \rceil + \lceil \frac{121}{7^1} \rceil + \lceil \frac{121}{7^2} \rceil \\ &\quad + \lceil \frac{121}{7^3} \rceil \\ &= \lceil \frac{121}{1} \rceil + \lceil \frac{121}{7} \rceil + \lceil \frac{121}{49} \rceil + \lceil \frac{121}{343} \rceil \\ &= 121 + 18 + 3 + 1 \cong 143 \end{aligned}$$

Then C is Optimal linear codes.

6. The code $[182, 4, 143]_7$ is Optimal

we get $n \geq g_q(k, d) = \sum_{i=0}^{k-1} \lceil \frac{d}{q^i} \rceil$ Since C is a $[182, 4, 143]_7$ - code, then:

$$\begin{aligned} n &\geq \sum_{i=0}^3 \lceil \frac{d}{q^i} \rceil = \lceil \frac{143}{7^0} \rceil + \lceil \frac{143}{7^1} \rceil + \lceil \frac{143}{7^2} \rceil \\ &\quad + \lceil \frac{143}{7^3} \rceil \\ &= \lceil \frac{143}{1} \rceil + \lceil \frac{143}{7} \rceil + \lceil \frac{143}{49} \rceil + \lceil \frac{143}{343} \rceil \\ &= 143 + 21 + 3 + 1 \cong 168 \end{aligned}$$

Then C is Optimal linear codes.

4. Conclusions

From field 7, we capes the weights of the writings and found that each weight differed from the others. It was shown that the dimension always equals 4. We applied the theorem to the codes and found that all the codes are perfect.

References

- [1] N.A.M. Al-Seraji, "Generalized of Optimal Codes". Al-Mustansiriyah Journal of Science, 2013.24(6), 101-108.
- [2] K. Shiromoto, and L. Storme, "A Griesmer bound for codes over finite quasi-frobenius rings". Electronic Notes in Discrete Mathematics, 2001.6, 337-345.
- [3] W. Zhang, B. Chen, N. Yu, "Improving various reversible data hiding schemes via optimal codes for binary covers". IEEE transactions on image processing, 2012.21(6), 2991-3003.
- [4] L. Baumert, R. McEliece, "A note on the Griesmer bound (Corresp.)". IEEE Transactions on Information Theory, 1973.19(1), 134-135.
- [5] Y. Kageyama, and T. Maruta, "On the geometric constructions of optimal linear codes". Designs, Codes and Cryptography, 2016.81(3), 469-480.

- [6] N.A.M. Al-seraji, "On optimal Codes. journal of science", Iraq, 2012.vol.23, no.3.
- [7] J. Hirschfeld, "Finite Projective spaces of three dimensions Clarendon Press Oxford"1985.
- [8] J.W.P.Hirschfeld, "Coding Theory".Spring2014.
- [9] Al-Zangana, E. B., & Yahya, N. Y. K. (2022). Subgroups and Orbits by Companion Matrix in Three Dimensional Projective Space. Baghdad Science Journal.Doi:<https://doi.org/10.21123/bsj.2022.19.4.0805>.
- [10] Yahya, N. Y. K. (2021). Applications geometry of space in $PG(3,P)$. Journal of Interdisciplinary Mathematics, 1-13.
- [11] Ali Ahmed A. Abdulla & Kasm Yahya, N. Y.(2021) A Geometric Construction of Surface Complete (k, r) -cap in $PG(3,7)$ In Journal of Physics: Conference Series 1879 (2021) 022112 IOP Publishing doi:10.1088/1742-6596/1879/2/022112.
- [12] Kasm Yahya, N. Y, Emad Bakr Al-Zangana [22] (2021) The Non-existence of $[1864, 3, 1828]_{53}$ Linear Code by Combinatorial Technique in International Journal of Mathematics and Computer Science, 16(2021), no. 4, 1575–1581.