



Paper Type: Original Article

Optimizing Hypergraph for Polynomials Modeling with Redundancy Scheduling

Marzieh Moradi Daleni* 

Department of Mathematics Neyriz Branch, Islamic Azad University, Fars, Iran; moradidalenimarzieh@gmail.com.

Citation:

Received: 17 April 2024

Revised: 21 July 2024

Accepted: 28 September 2024

Moradi Daleni, M. (2025). Optimizing hypergraph for polynomials modeling with redundancy scheduling. *Annals of optimization with applications*, 1 (1), 41-49.

Abstract

In this paper, we consider the minimization of two classes of polynomials over the standard simplex. These polynomials have their variables labeled by the edges of a complete uniform hypergraph, and their coefficients are defined in terms of some cardinality patterns of unions of edges. Data Envelopment Analysis (DEA) is a non-parametric method that aims to use scientific methods to investigate the performance of Decision-Making Units (DMUs). One of the interesting subjects in DEA is the minimization of the empirical error while satisfying some shape constraints, such as convexity and free disposability. In this research, the question is whether these polynomials attain their minimum value at the barycenter of the standard simplex, which corresponds to showing the optimality of the uniform distribution for the underlying queuing problem. The process focuses on the development of an adaptive observer-based Distributed Fault Estimation Observer (DFEO) for multi-agent nonlinear time-delay systems under a directed communication topology. The process involves constructing a fault estimation observer for each agent based on their relative output estimation errors.

Keywords: Polynomials, Data envelopment analysis, Optimization model, Hypergraph, Symmetric.

1 | Introduction

Optimizing hypergraph-based polynomials modeling job occupancy in queuing with redundancy scheduling. In this paper, we consider a question posed in [1] that arises from redundancy scheduling in queuing theory.

Redundancy scheduling is based on the idea that sending the same job to multiple distinct servers can be advantageous if balanced against the risk of wasted capacity. Here, one aims to determine the optimal policy for choosing which subset of servers to send the job copies to, and it is conjectured that a uniform probability distribution is optimal. This can be formulated as saying that a certain highly symmetric polynomial attains its

 Corresponding Author: moradidalenimarzieh@gmail.com

 <https://doi.org/10.48314/anowa.v1i1.38>



Licensee System Analytics. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0>).

minimum at the normalized all-one vector. While we are unable to prove the general case, we establish a similar result for a simplification of the family of polynomials by exploiting its symmetries, as well as some special cases of the original problem.

Indeed, research on multi-agent systems has garnered significant attention within both academic and industrial communities due to their diverse potential applications. Multi-agent systems are deployed in various areas, including uncrewed air vehicles, satellite formation, and sensor fusion. Symmetry is used more generally to give tractable reformulations for the semidefinite bounds arising from the following levels of Lasserre's hierarchy. For more examples and a broad exposition of the use of symmetry in semidefinite programming, we refer to [1] and further references therein.

On the other hand, Data Envelopment Analysis (DEA) is one of the existing techniques for estimating production functions and measuring efficiency [2]. The DEA relies on the construction of a polyhedral technology in the space of inputs and outputs that satisfies certain classical axioms of production theory, such as monotonicity and convexity.

It is a non-parametric, data-driven approach with many advantages from a benchmarking perspective. Additionally, the treatment of the multi-output, multi-input framework is relatively straightforward using DEA compared to other available methods. However, the DEA has been criticized for its non-statistical nature, being labeled as a purely descriptive tool of the data sample at a frontier level with little inferential power. In this paper, our main objective is to use DEA for polynomial optimization.

2 | Mathematical Preliminaries

2.1 | Data Envelopment Analysis

The DEA model, introduced by Charnes et al. [2], can estimate an efficiency frontier by considering the best performance observations (Extreme points), which "envelop" the remaining observations using mathematical programming techniques. The concept of efficiency can be defined as a ratio of produced outputs to the used inputs:

$$\text{Efficiency} = \frac{\text{outputs}}{\text{inputs}}. \quad (1)$$

So that an inefficient unit can become efficient by expanding products (Output), keeping the same level of used resources, reducing the used resources to keep the same production level, or by a combination of both.

Considering $j = 1, 2, 3, \dots, m$ Decision Making Units (DMUs) using $x_i | i = 1, 2, 3, \dots, n$ inputs to produce $y_r | r = 1, 2, 3, \dots, s$ outputs and prices (Multipliers) v_i and u_r associated with those inputs and outputs, we can also formalize the efficiency expression in *Model (1)* as the ratio of weighted outputs to weighted inputs:

$$\text{Efficiency} = \frac{\sum_{r=1}^s u_r y_{rj}}{\sum_{i=1}^n v_i x_{ij}}. \quad (2)$$

$$\max \frac{\sum_{r=1}^s u_r y_{ro}}{\sum_{i=1}^n v_i x_{io}} \quad \sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^n v_i x_{ij} \leq 0, \quad (3)$$

For all i, r, j $v_i, u_r \geq 0$.

This problem is denominated in the CCR constant return to scale input-oriented model, which by duality is equivalent to solving the following linear programming.

$$\begin{aligned} &\text{Min}(\theta), \\ &\sum_{j=1}^m z_j x_{ij} \leq \theta x_{io}, \end{aligned} \quad (4)$$

$$\sum_{j=1}^m z_j y_{rj} \geq y_{ro},$$

$$\sum_{j=1}^m z_j = 1, \quad z_j \geq 0.$$

As a result, we have an efficiency score θ , which varies from 0 to 1, designating the efficiency for each DMU. We can obtain the marginal contribution of each input and output in the multiplier of *Model (3)*, the peers of efficiency and respective weights in the primal (Or envelopment) form of *Model (4)*, and also the potential for improvements and slacks in an extension form of *Model (4)*.

2.2 | Polynomials

We now introduce the classes of polynomials of interest. Given integers $n, L \geq 2$, we set $V = [n] = \{1, \dots, n\}$ and $E = \{e \subseteq V: |e| = L\}$, so that (V, E) can be seen as the complete L -uniform hypergraph on n elements. We set

$$m = |E| = \binom{n}{L},$$

where we omit the explicit dependence on n, L to simplify notation, and we let

$$\Delta_m = \{x = (x_e)_{e \in E} \in \mathbb{R}^m: x_e \geq 0, \sum_{e \in E} x_e = 1\}.$$

denote the standard simplex in $\mathbb{R}^m$. The elements of Δ_m correspond to probability vectors on m items, and the barycenter $x^* = 1/m (1, \dots, 1)$ of Δ_m corresponds to the uniform probability vector.

Given an integer $d \geq 2$, we consider the following m -variate polynomial in the variables $x = (x_e: e \in E)$, which is a main player in the paper:

$$f_d(x) = \sum_{e \in E} \prod_{i=1}^d \frac{x_{e_i}}{|e_1 \cup \dots \cup e_i|}. \quad (5)$$

So, f_d is a homogeneous polynomial with degree d . We are interested in the following optimization problem

$$f_d^* = \min_{x \in \Delta_m} f_d(x),$$

asking to minimize the polynomial f_d over the simplex Δ_m . The main conjecture, which is stated in [5], claims that the minimum is attained at the uniform probability.

Conjecture 1: Given integers $n, d, L \geq 2$, is the polynomial $f_d(x)$ in *Model (5)* attains its minimum over Δ_m at the barycenter x^* of Δ_m .

As explained in [3], the motivation for this conjecture comes from its relevance to a problem in queuing theory, which we will briefly describe in the next section. In this chapter, we are only able to give a partial positive answer to this conjecture, namely, in the case $d = 2$ and the case $d = 3$ and $L = 2$.

As a first step toward understanding the polynomials f_d , we investigate a related, easier-to-analyze class of polynomials. Given an integer $d \geq 2$, we consider the following related class of polynomials.

$$p_d(x) = \sum_{e \in E^d} \frac{1}{|e_1 \cup \dots \cup e_d|}. \quad (6)$$

which are also homogeneous with degree d . For degree $d \geq 3$, the polynomials f_d have a related but more complicated structure than the polynomials p_d . Here, too, we may ask whether the minimum of p_d over the standard simplex Δ_m is attained at the uniform probability vector x^* . For the polynomials p_d , we can provide a positive answer in the general case.

The following is the first main result of the paper. As we will see, the analysis of the polynomials f_d is technically more involved than that of the polynomials p_d , and we have only partial results so far. In both

cases, the key ingredient is showing that the polynomials are convex on the simplex, i.e., that they have positive semidefinite Hessians at any vector in Δ_m .

It turns out that the Hessian of the polynomial pd enters some way as a component of the Hessian of the polynomial fd . So, this forms a natural motivation for the study of the polynomials pd , though they form a natural class of symmetric polynomials that are interesting for their own sake. Exploiting symmetry plays a central role in our proofs. Indeed, the key idea is to show that the polynomials are convex, which, combined with their symmetry properties, implies that the global minimum is attained at the barycenter of the simplex. To demonstrate this, we show that their Hessian matrices are positive semidefinite at each point of the simplex, which we achieve by exploiting their symmetry structure and links to Terwilliger algebras again.

Symmetry is a widely used ingredient in optimization, in particular in semidefinite optimization and algebraic questions involving polynomials. We mention a few landmark examples as background information. Symmetry can indeed be used to formulate equivalent, more compact reformulations for semidefinite programs.

The underlying mathematical fact is the Artin-Wedderburn theory, which shows that matrix $*$ -algebras can be block-diagonalized. An early well-known example is the linear programming reformulation from [2] for the Theta number of Hamming graphs, showing the link to the Delsarte bound and Bose-Mesner algebras of Hamming schemes [4].

3 | Graph

Consider a directed graph $g = (v, E, A)$ with a non-empty finite set of N nodes $v = (v_1, v_2, \dots, v_N)$, a set of edges or arcs $E \subset v \times v$, and the associated adjacency matrix $A = [a_{ij}] \in \mathbb{R}^{N \times N}$. In this paper, the graph is assumed to be time-invariant, i.e., A is constant. An edge rooted at node j and ended at node i is denoted by (v_j, v_i) , which means that information can flow from node j to node i . a_{ij} is the weight of the edge (v_j, v_i) and $a_{ij} = 1$ if $(v_j, v_i) \in E$ otherwise $a_{ij} = 0$. Node j is called a neighbor of node i if $(v_j, v_i) \in E$. The set of neighbors of node i is denoted as $N_i = \{j | (v_j, v_i) \in E\}$.

Define the in-degree matrix as $D = \text{diag}\{d_i\} \in \mathbb{R}^{N \times N}$ with $d_i = \sum_{j \in N_i} a_{ij}$ and the Laplacian matrix as $L = D - A$. The edges in the form of (v_i, v_i) are called loops. $G = \text{diag}\{g_i\} \in \mathbb{R}^{N \times N}$ is denoted as a loop matrix and has at least one diagonal item being *Model (1)*. A graph with loops is called a multi-graph; otherwise, it is a simple graph.

3.1 | System Description

Let us consider a group of N agents modeled under a communication graph, and each faulty agent is described by the following state-space model:

$$\dot{x}_i(t) = A_i x_i(t) + B_i u_i(t) + f_i(x_i, x_i(t - \tau_i)) + H_i \theta_i(t), \quad i = 1, 2, \dots, N. \quad (7)$$

$$y_i(t) = C_i x_i(t), \quad (8)$$

where $x_i(t) \in \mathbb{R}^n$, $u_i(t) \in \mathbb{R}^m$ and $y_i(t) \in \mathbb{R}^p$ are the states, the inputs, and the outputs of the i th agent, respectively. $f_i(\cdot) \in \mathbb{R}^n$ is the nonlinear function. $\theta_i(t) \in \mathbb{R}^r$ represents the system component or actuator fault, $\theta_i(t)$, and $\hat{\theta}_i(t)$ are assumed to be bounded and $\|\hat{\theta}_i(t)\| \leq \bar{\theta}_i$. A_i , B_i , H_i and C_i are constant matrices with appropriate dimensions.

Assumption 1. There exist known positive constants α_i, β_i such that the following Lipschitz inequality holds:

$$\|f_i(\hat{x}_i(t), \hat{x}_i(t - \tau_i)) - f_i(x_i(t), x_i(t - \tau_i))\| \leq \alpha_i \|\hat{x}_i(t) - x_i(t)\| + \beta_i \|\hat{x}_i(t - \tau_i) - x_i(t - \tau_i)\|. \quad (9)$$

Lemma 1. Assume that X and Y are vectors or matrices with appropriate dimensions, then a constant $\alpha > 0$ can be chosen, such that the following inequality always holds:

$$X^T Y + Y^T X < \alpha X^T X + \alpha^{-1} Y^T Y. \quad (10)$$

Lemma 2. If $\lim_{t \rightarrow \infty} \int_0^t f(\tau) d\tau$ exists and is finite and $f(t)$ is a uniformly continuous function, then $\lim_{t \rightarrow \infty} f(t) = 0$.

Assumption 2. There exists a symmetric positive definite matrix P_i , and matrix F_i such that the following equality is satisfied:

$$H_i^T P_i = F_i C_i. \quad (11)$$

Remark 1. *Assumption 2* is a common assumption to design adaptive fault diagnosis observers; obviously, *Condition (11)* is equality and can now be handled by the LMI toolbox directly. By choosing a small positive scalar σ_i , *Condition (11)* can be changed into the following inequality.

Minimize σ_i ,

S. t.

$$\begin{bmatrix} \sigma_i I, & H_i^T P_i - F_i C_i, \\ *, & \sigma_i I, \end{bmatrix} < 0. \quad (12)$$

4 | Motivation

Our motivation for studying the polynomials pd and fd stems from their relevance to a problem in queuing theory. The question of whether they attain their minimum at the uniform probability distribution was posed to us by the authors of [1], who conjecture this to establish a result about the asymptotic behavior of the job occupancy in a parallel-server system with redundancy scheduling in the light-traffic regime (In contrast to the heavy-traffic regime considered in [5]).

In what follows, we provide only a high-level overview of this connection, referring the reader to the paper [1] for a detailed exposition. We also refer to [5] for an extended review of the relevant literature. A crucial mechanism considered to enhance the performance of parallel-server systems in queuing theory is redundancy scheduling.

The key feature of this policy is that several replicas are created for each arriving job, which are then assigned to distinct servers. As soon as the first of these replicas completes service on a server, the remaining ones are stopped. The underlying idea is that sending replicas of the same job to multiple servers will increase the likelihood of shorter queuing times.

This, however, must be weighed against the risk of capacity wastage. An important question is thus to assess the impact of redundancy scheduling policies. While most papers in the literature on redundant scheduling assume that the set of servers to which the replicas are sent is selected uniformly at random, the paper considers the case when the set of servers is selected according to a given probability distribution. It investigates the impact of this probability distribution on the system's performance [3].

It is shown there that while the impact remains relatively limited in the heavy-traffic regime, the system occupancy is much more sensitive to the selected probability distribution in the light-traffic regime.

We will now introduce only a few elements of the model considered in [1] so that we can establish the connection to the polynomials studied in this paper. We keep our presentation high level and refer to [3] for details. The setting is as follows. There are n parallel servers, with average speed μ .

Jobs arrive as a Poisson process of rate $n\lambda$ for some $\lambda > 0$. When a job arrives, L replicas of it are created that are sent v with probability x_e to a subset $e \subseteq [n]$ of L servers. Here, $L \geq 2$ is an integer, and $x = (x_e)_{e \in E}$ is a probability distribution on the set $E = \{e \subseteq [n]: |e| = L\}$ of possible collections of L servers.

As noted in [3], this can be seen as selecting an edge $e \in E$ with probability x_e in the uniform hypergraph $(V=[n], E)$ (With edge size L). A significant performance parameter is the system occupancy at time t , which is represented by a vector $(e_1, \dots, e_M) \in E^M$, where $M = M(t)$ is the total number of jobs present in the system, and $e_i \in E$ is the collection of servers to which the replicas of the i th longest job in the system have been assigned.

We need three modeling assumptions. First, one needs to assume suitable stability conditions. Second, all servers should have the same speed μ , and third, the service requirements of the jobs are assumed to be independent and exponentially distributed with unit mean. Under these assumptions, the stationary distribution of the occupancy of the above edge selection is given by

$$\pi(e_1 \dots e_m) = c \prod_{i=1}^M \frac{n\lambda x_{e_i}}{\mu |e_1 \dots e_i|}. \quad (13)$$

For some constant $C > 0$. let $Q\lambda(x)$ be a random variable with the stationary distribution of the system occupancy when the edge selection is given by the probability vector $x = (x_e)_{e \in E}$. It then follows that, for any integer $d \geq 1$, the probability that d jobs are present in the system is given by

$$P[Q_\lambda(x) = d] = \sum_{e \in E^d} \pi(e_1 \dots e_d). \quad (14)$$

Therefore, $P[Q\lambda(x) = d]$ is the polynomial $f_d(x)$.

It would be worth mentioning that we have two stages in the search for the generalization error bound: the first stage is based on the construction of the class of piecewise linear hypotheses whose elements are polynomials that are located as close as possible to the data sample, and the second stage is based on the construction of the bound of the fat-shattering dimension of the class of hypothesis constructed in the first stage.

The minimization of the bound on the expected error using the bound on the fat-shattering dimension calculated yields the DEA-based Machines (DEAM) model as a method for estimating piecewise linear production functions, which minimizes both the generalization error and the empirical error [6–10].

5 | Simulation Results

In this section, an example is given to verify the effectiveness of the proposed method. Consider the following equations:

$$\begin{aligned} \dot{x}_{i1} &= x_{i2}, \\ \dot{x}_{i2} &= \left(\frac{m_i g r}{J_i} - \frac{k r^2}{4 J_i} \right) \sin(x_{i1}(t - d_i(t))) + \frac{K r (1-b)}{2 J_i} + \frac{u_i}{J_i}, \\ y_i &= x_{i1} + x_{i2}, \end{aligned} \quad (15)$$

where $m_i = 2$ kg, $J_i = 0.5 \frac{N \cdot S}{m^2}$, $k = 100$ N/m, $r = 0.5$ m, $l = 0.5$ m, and $g = 9.81 \frac{m}{s^2}$, and $b = 0.4$ m. The time delay of each subsystem is selected as $d_i(t) = 0.4 + 0.2 \sin(2t)$, $i = 1, \dots, 4$.

We consider a simple multiplicative actuator fault in each subsystem. We let $u_i = \bar{u}_i + \theta_i \bar{u}_i$, where $\bar{u}_i$ is the nominal control input in the no-fault case ($\bar{u}_i = -20y_i$), and $\theta_i \in [-1, 0]$ is the parameter indicating the magnitude of the fault. To simulate a fault, we set $\theta_i = -0.5$, for *Subsystem 1* at $T_a = 5$ sec (when a fault

happens) and $\theta_3 = -0.5$, for *Subsystem 3* at $T_a = 7$ sec (When a fault occurs). Moreover, it is assumed that $\theta_i (i=1, \dots, 4)$ is zero before occurring faults. For *Subsystems 2-4*, the time-varying faults are considered as

$$\theta_2 = 0.4\sin(t)\cos(0.3t), \theta_4 = 0.3\sin(t)\cos(0.3t).$$

Correspondingly. For *Subsystems 2-4*, the faults are considered to occur at $T_a = 6$ and $T_a = 9$, respectively. We consider the directed graph topology for simulation, which is illustrated in *Fig. 1*. Moreover, the first and third nodes contain a self-loop.

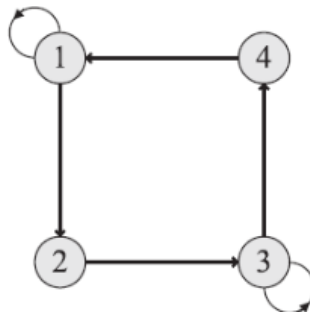


Fig. 1. Communication topology.

For such directed communication topology, we can get the matrix $L + G$ as follows:

$$L + G = \begin{bmatrix} 2 & 0 & 0 & -1 \\ -1 & 1 & 0 & 0 \\ 0 & -1 & 2 & 0 \\ 0 & 0 & -1 & 1 \end{bmatrix}.$$

Each subsystem consists of two states (x_{i1}, x_{i2}). The *Figs. 2-4* display the states and estimations of the first to fourth subsystems. For a more detailed explanation, the first state and its estimation of the first subsystem are demonstrated in *Fig. 2*.

As evident from the simulation results, the estimation of states in each subsystem can converge to their corresponding states, even in the presence of faults. The simulation results validate that the proposed observer accurately estimates the states of the system within each subsystem, even when faults occur.

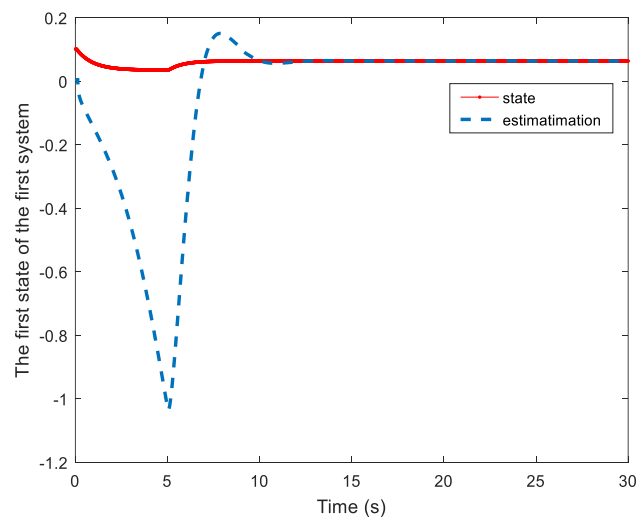


Fig. 2. The first state and its estimation of the first subsystem.

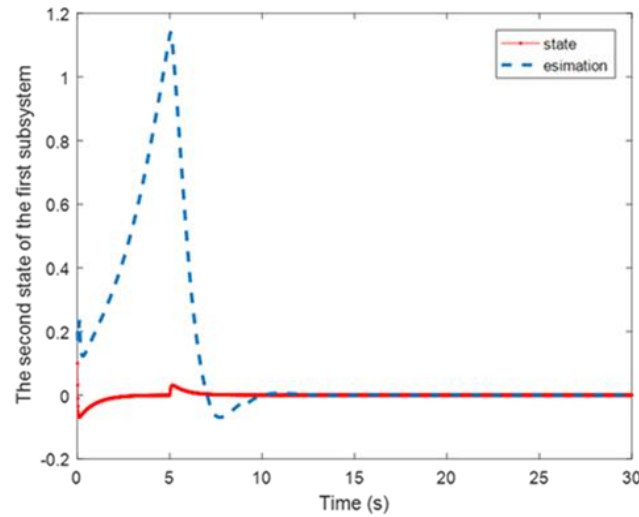


Fig. 3. The second state and its estimation of the first subsystem.

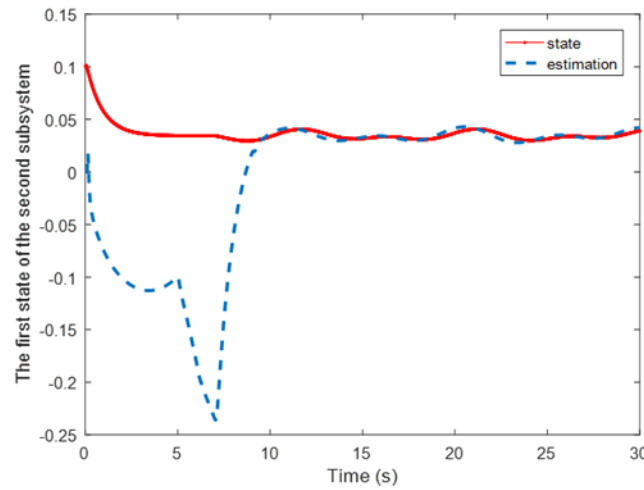


Fig. 4. The first state and its estimation of the second subsystem.

6 | Conclusion

In this paper, an adaptive fault estimation observer approach for nonlinear time-delay multi-agent systems has been investigated. The fault estimator is designed based on the relative output estimation errors of the whole system. Finally, simulation results are presented to confirm that the proposed design technique enables accurate fault estimation of nonlinear time-delay multi-agent systems.

Conflict of Interest

The authors declare no conflict of interest.

Data Availability

All data are included in the text.

Funding

This research received no specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

References

- [1] Cardinaels, E., Borst, S., & van Leeuwen, J. S. H. (2022). Power-of-two sampling in redundancy systems: The impact of assignment constraints. *Operations research letters*, 50(6), 699–706. <https://doi.org/10.1016/j.orl.2022.10.006>
- [2] Schrijver, A. (2005). New code upper bounds from the terwilliger algebra and semidefinite programming. *IEEE transactions on information theory*, 51(8), 2859–2866. <https://doi.org/10.1109/TIT.2005.851748>
- [3] Charnes, A., Cooper, W. W., & Rhodes, E. (1978). Measuring the efficiency of decision making units. *European journal of operational research*, 2(6), 429–444. [https://doi.org/10.1016/0377-2217\(78\)90138-8](https://doi.org/10.1016/0377-2217(78)90138-8)
- [4] Schrijver, A. (2003). A comparison of the Delsarte and Lovász bounds. *IEEE transactions on information theory*, 25(4), 425–429. <https://doi.org/10.1109/TIT.1979.1056072>
- [5] Delsarte, P., & Levenshtein, V. I. (1998). Association schemes and coding theory. *IEEE transactions on information theory*, 44(6), 2477–2504. <https://doi.org/10.1109/18.720545>
- [6] Doyle, J. R., & Green, R. H. (1991). Comparing products using data envelopment analysis. *Omega*, 19(6), 631–638. [https://doi.org/10.1016/0305-0483\(91\)90012-I](https://doi.org/10.1016/0305-0483(91)90012-I)
- [7] Houshyar, E., Davoodi, M. J. S., & Nassiri, S. M. (2010). Energy efficiency for wheat production using data envelopment analysis (DEA) technique., 6(4), 663–672. <http://www.ijat-rmutto.com/>
- [8] Cook, W. D., Kress, M., & Seiford, L. M. (1996). Data envelopment analysis in the presence of both quantitative and qualitative factors. *Journal of the operational research society*, 47(7), 945–953. <https://doi.org/10.1057/jors.1996.120>
- [9] Guerrero, N. M., Aparicio, J., & Valero-Carreras, D. (2022). Combining data envelopment analysis and machine learning. *Mathematics*, 10(6), 909. <https://doi.org/10.3390/math10060909>
- [10] Ucal Sari, I., & Ak, U. (2022). Machine efficiency measurement in industry 4.0 using fuzzy data envelopment analysis. *Journal of fuzzy extension and applications*, 3(2), 177–191. <https://doi.org/10.22105/jfea.2022.326644.1199>