



Core Concepts Behind Quadripartitioned Neutrosophic Soft Block Matrices

S. Ramesh Kumar ^{1*}, A. Stanis Arul Mary ²

¹ Department of Mathematics, Dr. N. G. P. Arts and Science College, Coimbatore, India

² Department of Mathematics, Nirmala College for Women, Coimbatore, India

*Corresponding Author : kamalarames99@gamil.com,

Abstract

We introduce quadripartitioned neutrosophic soft block matrices, extending neutrosophic soft matrices with a four-part structure to model uncertainty, indeterminacy, falsity, and a new component: uncertainty, in decision-making. This enhancement overcomes limitations of traditional neutrosophic soft matrices by comprehensively representing complex situations where all four aspects must be simultaneously considered. Each matrix element is divided into four sub-components for a more nuanced analysis of the decision alternatives and criteria. The paper details fundamental matrix operations (addition, subtraction, multiplication, and inversion) and explores their properties. Case studies in decision support systems and optimization problems illustrate the framework's applicability, particularly where traditional methods are insufficient. The model proves especially valuable for complex decision-making in artificial intelligence, pattern recognition, and fuzzy logic, handling multiple uncertainty layers. Finally, we examine the theoretical basis and practical uses of quadripartitioned neutrosophic soft block matrices, providing a new tool for soft computing and decision science.

Keywords: Quadripartitioned Neutrosophic Soft Block Matrices; Soft Computing; Decision Science.

Introduction:

Neutrosophic soft block matrices offer a novel approach within mathematics and computational intelligence, specifically designed to tackle the complexities and uncertainties inherent in decision-making, data analysis, and other applications where information can be incomplete, indeterminate, or contradictory. Emerging from the fusion of neutrosophic set theory and soft set theory, these matrices provide a structured framework that seamlessly integrates the power of matrices with the adaptability of neutrosophic soft sets.

Neutrosophic sets, pioneered by Smarandache [7], represent a generalization of classical and fuzzy sets, allowing for the representation of truth, indeterminacy, and falsity degrees. These sets prove particularly valuable when dealing with information that is not simply uncertain or imprecise, but also conflicting. Conversely, soft set theory, introduced by Molodtsov [3], offers a parameterized framework for managing uncertainty without the constraints of traditional mathematical tools. Neutrosophic soft Matrix introduced by Deli et al. [10], By integrating these two frameworks into a matrix structure, neutrosophic soft block matrices can effectively capture multifaceted, indeterminate information across various blocks, where each block corresponds to distinct parameters or attributes.

This hybrid structure enhances the ability to model intricate data relationships and facilitates more nuanced analyses, particularly within fields like decision-making systems, data classification, and pattern recognition. Through the manipulation of these matrices, researchers can develop algorithms and procedures for more effective handling of multi-dimensional uncertainty, leading to robust solutions for real-world problems. This study delves into the fundamental principles, properties, and potential applications of neutrosophic soft block matrices [25], contributing to the broader field of neutrosophic mathematics and its practical relevance.

Different kinds of quadripartitioned neutrosophic soft block matrices and the operations that can be carried out on them are the main topics of this article. A concise summary of key terms pertaining to quadripartitioned neutrosophic sets, quadripartitioned neutrosophic soft sets, quadripartitioned neutrosophic soft matrices, quadripartitioned neutrosophic soft block matrices, and other ideas is given in the section that follows. While Section 4 presents certain types of quadripartitioned neutrosophic soft block matrices and talks about their related qualities, Section 3 outlines the operations that apply to quadripartitioned neutrosophic soft matrices.

2. Prelims

Definition 2.1: Neutrosophic sets (Smarandache, 2005)

A *neutrosophic set* is a mathematical framework designed to handle information that may be true, false, or indeterminate, often in complex or uncertain environments. In this structure, $\mathcal{A} = \{ \langle x, \mathcal{T}_{\mathcal{A}}(x), \mathcal{I}_{\mathcal{A}}(x), \mathcal{F}_{\mathcal{A}}(x) \rangle : x \in \mathcal{X} \}$, $0 \leq \mathcal{T}_{\mathcal{A}}(x) + \mathcal{I}_{\mathcal{A}}(x) + \mathcal{F}_{\mathcal{A}}(x) \leq 3$, each element has three independent membership degrees: truth (T), indeterminacy (I), and falsity (F). These values are typically represented in the range $[0, 1]$ but are not constrained to sum to a particular value, allowing for the flexible representation of situations where truth, indeterminacy, and falsity coexist to varying extents.

Definition 2.2: Neutrosophic soft set (Maji, 2013)

. A *soft set* is a mathematical concept used to model uncertainty without the limitations of traditional methods. It is defined as a parameterized family of subsets of a universal set, where each parameter is associated with a subset that describes the elements related to that parameter. Formally, a soft set over a universal set $\mathcal{U}$ and a set of parameters $\mathcal{E}$ is a function $\mathcal{F}$ that maps each parameter $e \in \mathcal{E}$ to a subset $\mathcal{F}(e) \subseteq \mathcal{U}$. Soft sets are particularly useful in decision-making processes because they provide a flexible approach to handle ambiguous and vague data by allowing different subsets to represent various attributes or conditions.

Definition 2.3: Given a universal set $\mathcal{X}$ and a parameter set $\mathcal{R}$, let $\mathcal{A}$ be a non-empty subset of $\mathcal{R}$. Define $\mathcal{R}(\mathcal{X})$ as the set of all quadripartitioned neutrosophic sets in $\mathcal{X}$. The pair $(\mathcal{K}, \mathcal{A})$ constitutes a quadripartitioned neutrosophic soft set (QNSS)[26] over $\mathcal{X}$, where $\mathcal{K}$ is a mapping from $\mathcal{A}$ to the power set of $\mathcal{X}$.

Each element of $\mathcal{A}$ is represented as $\mathcal{A} = \{ \langle x, \mathcal{T}_{\mathcal{A}}(x), \mathcal{C}_{\mathcal{A}}(x), \mathcal{U}_{\mathcal{A}}(x), \mathcal{F}_{\mathcal{A}}(x) \rangle : x \in \mathcal{X} \}$, where x is from $\mathcal{X}$ and $\mathcal{T}_{\mathcal{A}}(x), \mathcal{C}_{\mathcal{A}}(x), \mathcal{U}_{\mathcal{A}}(x), \mathcal{F}_{\mathcal{A}}(x)$ are functions from $\mathcal{X}$ to $[0, 1]$ and $\mathcal{T}_{\mathcal{A}}(x), \mathcal{C}_{\mathcal{A}}(x), \mathcal{U}_{\mathcal{A}}(x), \mathcal{F}_{\mathcal{A}}(x)$ are Truth Contradiction, Ignorance, False membership function respectively.

All membership functions must satisfy $0 \leq \mathcal{T}_{\mathcal{A}}(x), \mathcal{C}_{\mathcal{A}}(x), \mathcal{U}_{\mathcal{A}}(x), \mathcal{F}_{\mathcal{A}}(x) \leq 4$.

3. Quadripartitioned Neutrosophic Soft Matrix

Definition 3.1: A Quadripartitioned Neutrosophic Soft Matrix [QNSM] is a matrix representation that combines quadripartitioned neutrosophic logic with soft set theory, allowing for the handling of uncertainty, indeterminacy, and vagueness in decision-making and problem-solving processes.

Formally, let X be a universal set, and E be a set of parameters. A QNSM over U and E is defined as: $\text{QNSM} = [(T_{ij}, C_{ij}, U_{ij}, F_{ij})]_{m \times n}$

where:

- $T_{ij}, C_{ij}, U_{ij}, F_{ij}$ represent the truth-membership, contradiction-membership, ignorance membership and falsity-membership degrees, respectively, of the element $x_i \in X$ concerning the parameter $e_j \in E$.
- $T_{ij}, C_{ij}, U_{ij}, F_{ij}$ are values in the interval $[0, 1]$, reflecting the degrees of truth, contradiction, ignorance membership, and falsity, respectively, and they allow us to express uncertainty in more complex situations.

A QNSM can be represented in matrix form as follows:

Let $X = \{x_1, x_2, \dots, x_m\}$ be a universal set and $E = \{e_1, e_2, \dots, e_n\}$ be a set of parameters. The quadripartitioned neutrosophic soft matrix QNSM over X and E is defined as:

$$\begin{bmatrix} (T_{11}, C_{11}, U_{11}, F_{11}) & (T_{12}, C_{12}, U_{12}, F_{12}) & \cdots & (T_{1n}, C_{1n}, U_{1n}, F_{1n}) \\ (T_{21}, C_{21}, U_{21}, F_{21}) & (T_{22}, C_{22}, U_{22}, F_{22}) & \cdots & (T_{2n}, C_{2n}, U_{2n}, F_{2n}) \\ \vdots & \vdots & \ddots & \vdots \\ (T_{m1}, C_{m1}, U_{m1}, F_{m1}) & (T_{m2}, C_{m2}, U_{m2}, F_{m2}) & \cdots & (T_{mn}, C_{mn}, U_{mn}, F_{mn}) \end{bmatrix}$$

Example :

$$\begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) & (0.7, 0.5, 0.4, 0.3) & (0.2, 0.3, 0.5, 0.6) \\ (0.4, 0.3, 0.2, 0.6) & (0.2, 0.3, 0.8, 0.3) & (0.8, 0.2, 0.7, 0.5) & (0.4, 0.5, 0.6, 0.7) \\ (0.3, 0.2, 0.6, 0.1) & (0.8, 0.5, 0.2, 0.3) & (0.9, 0.1, 0.4, 0.2) & (0.5, 0.4, 0.2, 0.3) \\ (0.6, 0.7, 0.8, 0.2) & (0.8, 0.6, 0.5, 0.4) & (0.2, 0.3, 0.5, 0.6) & (0.1, 0.6, 0.7, 0.5) \end{bmatrix}$$

Definition 3.2: A triangular QNSM is a specific type of square QNSM..

A square QNSM is classified as a lower triangular matrix if all its entries a_{ij} are equal to $(0, 0, 0, 1)$, for all integers i and j such that $1 \leq i < j \leq n$.

Conversely, If each entry a_{ij} is equal to the vector $(0, 0, 0, 1)$, then it is an upper triangular matrix with for all integers i and j , such that $1 \leq i \leq n$ and $1 \leq j \leq n$, if $i > j$.

Definition 3.3: A zero QNSM consists entirely of entries equal to $(0, 0, 0, 1)$ within a QNSM structure.

Definition 3.4: A Toeplitz QNSM is a square QNSM structured in the following way.

$$\mathcal{K} = \begin{bmatrix} (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) & (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}}) & (T_{13}^{\mathcal{K}}, C_{13}^{\mathcal{K}}, U_{13}^{\mathcal{K}}, F_{13}^{\mathcal{K}}) & (T_{14}^{\mathcal{K}}, C_{14}^{\mathcal{K}}, U_{14}^{\mathcal{K}}, F_{14}^{\mathcal{K}}) \\ (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) & (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}}) & (T_{13}^{\mathcal{K}}, C_{13}^{\mathcal{K}}, U_{13}^{\mathcal{K}}, F_{13}^{\mathcal{K}}) \\ (T_{31}^{\mathcal{K}}, C_{31}^{\mathcal{K}}, U_{31}^{\mathcal{K}}, F_{31}^{\mathcal{K}}) & (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) & (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}}) \\ (T_{41}^{\mathcal{K}}, C_{41}^{\mathcal{K}}, U_{41}^{\mathcal{K}}, F_{41}^{\mathcal{K}}) & (T_{31}^{\mathcal{K}}, C_{31}^{\mathcal{K}}, U_{31}^{\mathcal{K}}, F_{31}^{\mathcal{K}}) & (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) \end{bmatrix}$$

Example :

$$\mathcal{K} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) & (0.7, 0.5, 0.4, 0.3) & (0.2, 0.3, 0.5, 0.6) \\ (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) & (0.7, 0.5, 0.4, 0.3) \\ (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ (0.6, 0.7, 0.8, 0.2) & (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$$

Definition 3.5: Among QNSMs, a QNS tridiagonal matrix is distinguished by its structure, featuring non-zero elements solely along the lower diagonal, main diagonal, and upper diagonal. All remaining entries are populated with the quaternion (0, 0, 0, 1). Therefore, a QNS tridiagonal matrix, represented by A, takes the following form:

$$\begin{bmatrix} A_1 & B_1 & 0 & \cdots & 0 \\ C_1 & A_2 & B_2 & 0 & \vdots \\ 0 & C_2 & A_3 & B_3 & \vdots \\ \vdots & 0 & C_3 & \ddots & B_{n-1} \\ 0 & \cdots & \cdots & C_{n-1} & A_n \end{bmatrix}$$

Definition 3.6: A QNS block matrix is a QNSM divided into smaller sections called blocks or submatrices. Essentially, it's like drawing lines parallel to the rows and columns of the original matrix to visually create these submatrices.

These submatrices can be viewed as the individual components of the larger matrix. Any QNSM can be structured as a block matrix in various ways, depending on how its rows and columns are divided into partitions.

$$\mathcal{K} = \begin{bmatrix} (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) & (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}}) & \vdots & (T_{13}^{\mathcal{K}}, C_{13}^{\mathcal{K}}, U_{13}^{\mathcal{K}}, F_{13}^{\mathcal{K}}) & (T_{14}^{\mathcal{K}}, C_{14}^{\mathcal{K}}, U_{14}^{\mathcal{K}}, F_{14}^{\mathcal{K}}) \\ \cdots & \cdots & \vdots & \cdots & \cdots \\ (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{22}^{\mathcal{K}}, C_{22}^{\mathcal{K}}, U_{22}^{\mathcal{K}}, F_{22}^{\mathcal{K}}) & \vdots & (T_{23}^{\mathcal{K}}, C_{23}^{\mathcal{K}}, U_{23}^{\mathcal{K}}, F_{23}^{\mathcal{K}}) & (T_{24}^{\mathcal{K}}, C_{24}^{\mathcal{K}}, U_{24}^{\mathcal{K}}, F_{24}^{\mathcal{K}}) \\ (T_{31}^{\mathcal{K}}, C_{31}^{\mathcal{K}}, U_{31}^{\mathcal{K}}, F_{31}^{\mathcal{K}}) & (T_{32}^{\mathcal{K}}, C_{32}^{\mathcal{K}}, U_{32}^{\mathcal{K}}, F_{32}^{\mathcal{K}}) & \vdots & (T_{33}^{\mathcal{K}}, C_{33}^{\mathcal{K}}, U_{33}^{\mathcal{K}}, F_{33}^{\mathcal{K}}) & (T_{34}^{\mathcal{K}}, C_{34}^{\mathcal{K}}, U_{34}^{\mathcal{K}}, F_{34}^{\mathcal{K}}) \end{bmatrix}$$

Example :

$$\mathcal{K} = \begin{bmatrix} \mathcal{R}_{11} & \mathcal{R}_{12} \\ \mathcal{R}_{21} & \mathcal{R}_{22} \end{bmatrix}$$

Where $\mathcal{R}_{11} = [(T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) \quad (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}})]$

$$\mathcal{R}_{12} = [(T_{13}^{\mathcal{K}}, C_{13}^{\mathcal{K}}, U_{13}^{\mathcal{K}}, F_{13}^{\mathcal{K}}) \quad (T_{14}^{\mathcal{K}}, C_{14}^{\mathcal{K}}, U_{14}^{\mathcal{K}}, F_{14}^{\mathcal{K}})]$$

$$\mathcal{R}_{21} = \begin{bmatrix} (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{22}^{\mathcal{K}}, C_{22}^{\mathcal{K}}, U_{22}^{\mathcal{K}}, F_{22}^{\mathcal{K}}) \\ (T_{31}^{\mathcal{K}}, C_{31}^{\mathcal{K}}, U_{31}^{\mathcal{K}}, F_{31}^{\mathcal{K}}) & (T_{32}^{\mathcal{K}}, C_{32}^{\mathcal{K}}, U_{32}^{\mathcal{K}}, F_{32}^{\mathcal{K}}) \end{bmatrix}$$

$$\mathcal{R}_{22} = \begin{bmatrix} (T_{23}^{\mathcal{K}}, C_{23}^{\mathcal{K}}, U_{23}^{\mathcal{K}}, F_{23}^{\mathcal{K}}) & (T_{24}^{\mathcal{K}}, C_{24}^{\mathcal{K}}, U_{24}^{\mathcal{K}}, F_{24}^{\mathcal{K}}) \\ (T_{33}^{\mathcal{K}}, C_{33}^{\mathcal{K}}, U_{33}^{\mathcal{K}}, F_{33}^{\mathcal{K}}) & (T_{34}^{\mathcal{K}}, C_{34}^{\mathcal{K}}, U_{34}^{\mathcal{K}}, F_{34}^{\mathcal{K}}) \end{bmatrix}$$

Example :

$$\mathcal{K} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) & \vdots & (0.7, 0.5, 0.4, 0.3) & (0.2, 0.3, 0.5, 0.6) \\ \cdots & \cdots & \vdots & \cdots & \cdots \\ (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) & \vdots & (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ (0.6, 0.7, 0.8, 0.2) & (0.3, 0.2, 0.6, 0.1) & \vdots & (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$$

$$\mathcal{K} = \begin{bmatrix} \mathcal{R}_{11} & \mathcal{R}_{12} \\ \mathcal{R}_{21} & \mathcal{R}_{22} \end{bmatrix}$$

Where $\mathcal{R}_{11} = [(0.6, 0.5, 0.4, 0.3) \quad (0.5, 0.6, 0.7, 0.2)]$

$\mathcal{R}_{12} = [(0.7, 0.5, 0.4, 0.3) \quad (0.2, 0.3, 0.5, 0.6)]$

$\mathcal{R}_{21} = \begin{bmatrix} (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) \\ (0.6, 0.7, 0.8, 0.2) & (0.3, 0.2, 0.6, 0.1) \end{bmatrix}$

$\mathcal{R}_{22} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$

Definition 3.7: A square QNS block matrix is a QNS block matrix where the number of rows and the number of columns is equal.

$$\mathcal{K} = \begin{bmatrix} (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) & (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}}) & \vdots & (T_{13}^{\mathcal{K}}, C_{13}^{\mathcal{K}}, U_{13}^{\mathcal{K}}, F_{13}^{\mathcal{K}}) & (T_{14}^{\mathcal{K}}, C_{14}^{\mathcal{K}}, U_{14}^{\mathcal{K}}, F_{14}^{\mathcal{K}}) \\ (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{22}^{\mathcal{K}}, C_{22}^{\mathcal{K}}, U_{22}^{\mathcal{K}}, F_{22}^{\mathcal{K}}) & \vdots & (T_{23}^{\mathcal{K}}, C_{23}^{\mathcal{K}}, U_{23}^{\mathcal{K}}, F_{23}^{\mathcal{K}}) & (T_{24}^{\mathcal{K}}, C_{24}^{\mathcal{K}}, U_{24}^{\mathcal{K}}, F_{24}^{\mathcal{K}}) \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ (T_{31}^{\mathcal{K}}, C_{31}^{\mathcal{K}}, U_{31}^{\mathcal{K}}, F_{31}^{\mathcal{K}}) & (T_{32}^{\mathcal{K}}, C_{32}^{\mathcal{K}}, U_{32}^{\mathcal{K}}, F_{32}^{\mathcal{K}}) & \vdots & (T_{33}^{\mathcal{K}}, C_{33}^{\mathcal{K}}, U_{33}^{\mathcal{K}}, F_{33}^{\mathcal{K}}) & (T_{34}^{\mathcal{K}}, C_{34}^{\mathcal{K}}, U_{34}^{\mathcal{K}}, F_{34}^{\mathcal{K}}) \\ (T_{41}^{\mathcal{K}}, C_{41}^{\mathcal{K}}, U_{41}^{\mathcal{K}}, F_{41}^{\mathcal{K}}) & (T_{42}^{\mathcal{K}}, C_{42}^{\mathcal{K}}, U_{42}^{\mathcal{K}}, F_{42}^{\mathcal{K}}) & \vdots & (T_{43}^{\mathcal{K}}, C_{43}^{\mathcal{K}}, U_{43}^{\mathcal{K}}, F_{43}^{\mathcal{K}}) & (T_{44}^{\mathcal{K}}, C_{44}^{\mathcal{K}}, U_{44}^{\mathcal{K}}, F_{44}^{\mathcal{K}}) \end{bmatrix}$$

$$\mathcal{K} = \begin{bmatrix} \mathcal{R}_{11} & \mathcal{R}_{12} \\ \mathcal{R}_{21} & \mathcal{R}_{22} \end{bmatrix}$$

It is a square QNS block matrix because each $\mathcal{R}_{ij}$ is a square block.

Example :

$$\mathcal{K} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) & \vdots & (0.7, 0.5, 0.4, 0.3) & (0.2, 0.3, 0.5, 0.6) \\ (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) & \vdots & (0.5, 0.6, 0.7, 0.2) & (0.7, 0.5, 0.4, 0.3) \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) & \vdots & (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ (0.6, 0.7, 0.8, 0.2) & (0.3, 0.2, 0.6, 0.1) & \vdots & (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$$

$$\mathcal{K} = \begin{bmatrix} \mathcal{R}_{11} & \mathcal{R}_{12} \\ \mathcal{R}_{21} & \mathcal{R}_{22} \end{bmatrix}$$

Where $\mathcal{R}_{11} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$

$\mathcal{R}_{12} = \begin{bmatrix} (0.7, 0.5, 0.4, 0.3) & (0.2, 0.3, 0.5, 0.6) \\ (0.5, 0.6, 0.7, 0.2) & (0.7, 0.5, 0.4, 0.3) \end{bmatrix}$

$\mathcal{R}_{21} = \begin{bmatrix} (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) \\ (0.6, 0.7, 0.8, 0.2) & (0.3, 0.2, 0.6, 0.1) \end{bmatrix}$

$\mathcal{R}_{22} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$

Definition 3.8: In a rectangular QNS block matrix, the number of block rows and the number of block columns are not the same.

Example :

$$\mathcal{K} = \begin{bmatrix} (T_{11}^{\mathcal{K}}, C_{11}^{\mathcal{K}}, U_{11}^{\mathcal{K}}, F_{11}^{\mathcal{K}}) & (T_{12}^{\mathcal{K}}, C_{12}^{\mathcal{K}}, U_{12}^{\mathcal{K}}, F_{12}^{\mathcal{K}}) & \vdots & (T_{13}^{\mathcal{K}}, C_{13}^{\mathcal{K}}, U_{13}^{\mathcal{K}}, F_{13}^{\mathcal{K}}) & (T_{14}^{\mathcal{K}}, C_{14}^{\mathcal{K}}, U_{14}^{\mathcal{K}}, F_{14}^{\mathcal{K}}) \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ (T_{21}^{\mathcal{K}}, C_{21}^{\mathcal{K}}, U_{21}^{\mathcal{K}}, F_{21}^{\mathcal{K}}) & (T_{22}^{\mathcal{K}}, C_{22}^{\mathcal{K}}, U_{22}^{\mathcal{K}}, F_{22}^{\mathcal{K}}) & \vdots & (T_{23}^{\mathcal{K}}, C_{23}^{\mathcal{K}}, U_{23}^{\mathcal{K}}, F_{23}^{\mathcal{K}}) & (T_{24}^{\mathcal{K}}, C_{24}^{\mathcal{K}}, U_{24}^{\mathcal{K}}, F_{24}^{\mathcal{K}}) \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ (T_{31}^{\mathcal{K}}, C_{31}^{\mathcal{K}}, U_{31}^{\mathcal{K}}, F_{31}^{\mathcal{K}}) & (T_{32}^{\mathcal{K}}, C_{32}^{\mathcal{K}}, U_{32}^{\mathcal{K}}, F_{32}^{\mathcal{K}}) & \vdots & (T_{33}^{\mathcal{K}}, C_{33}^{\mathcal{K}}, U_{33}^{\mathcal{K}}, F_{33}^{\mathcal{K}}) & (T_{34}^{\mathcal{K}}, C_{34}^{\mathcal{K}}, U_{34}^{\mathcal{K}}, F_{34}^{\mathcal{K}}) \end{bmatrix}$$

$$\mathcal{K} = \begin{bmatrix} (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) & \vdots & (0.7, 0.5, 0.4, 0.3) & (0.2, 0.3, 0.5, 0.6) \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ (0.3, 0.2, 0.6, 0.1) & (0.4, 0.3, 0.2, 0.6) & \vdots & (0.6, 0.5, 0.4, 0.3) & (0.5, 0.6, 0.7, 0.2) \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ (0.6, 0.7, 0.8, 0.2) & (0.3, 0.2, 0.6, 0.1) & \vdots & (0.4, 0.3, 0.2, 0.6) & (0.6, 0.5, 0.4, 0.3) \end{bmatrix}$$

4. Mathematical Operations with QNSM:

Definition 4.1

Let $\mathcal{K}$ and $\mathcal{L}$ be two QNSM, where:

- $\mathcal{K}$ is represented as $(T_{ij}^{\mathcal{K}}, C_{ij}^{\mathcal{K}}, U_{ij}^{\mathcal{K}}, F_{ij}^{\mathcal{K}})$ for all i and j .
- $\mathcal{L}$ is represented as $(T_{ij}^{\mathcal{L}}, C_{ij}^{\mathcal{L}}, U_{ij}^{\mathcal{L}}, F_{ij}^{\mathcal{L}})$ for all i and j .

Then, the max-min product of these two QNSM, $\mathcal{K}$ and $\mathcal{L}$, denoted as $\mathcal{K} \oplus \mathcal{L}$, is defined as follows:

$[\max(T_{ij}^{\mathcal{K}}, T_{ij}^{\mathcal{L}}), \min(C_{ij}^{\mathcal{K}}, C_{ij}^{\mathcal{L}}), \min(U_{ij}^{\mathcal{K}}, U_{ij}^{\mathcal{L}}), \min(F_{ij}^{\mathcal{K}}, F_{ij}^{\mathcal{L}})]$ for all i and j .

Definition 4.2: Max-Min Product of QNSM:

Consider two QNSM, denoted as:

- $\mathcal{K} = (T_{ij}^{\mathcal{K}}, C_{ij}^{\mathcal{K}}, U_{ij}^{\mathcal{K}}, F_{ij}^{\mathcal{K}})$ for all i and j .
- $\mathcal{L} = (T_{ij}^{\mathcal{L}}, C_{ij}^{\mathcal{L}}, U_{ij}^{\mathcal{L}}, F_{ij}^{\mathcal{L}})$ for all i and j .

Definition 4.3: The **max-min product** of these two QNSM, $\mathcal{K}$ and $\mathcal{L}$, is represented as $\mathcal{K}\mathcal{L}$. The following definition describes the product:

$\mathcal{K}\mathcal{L}_{ij} = [\max(\min(T_{ij}^{\mathcal{K}}, T_{ij}^{\mathcal{L}})), \min(\max(C_{ij}^{\mathcal{K}}, C_{ij}^{\mathcal{L}})), \min(\max(U_{ij}^{\mathcal{K}}, U_{ij}^{\mathcal{L}})), \min(\max(F_{ij}^{\mathcal{K}}, F_{ij}^{\mathcal{L}}))]$

where this holds for all values of 'i' and 'j'."

Definition 4.4: Transpose of QNSM :

Consider a QNSM, represented as $\mathcal{K} = [(T_{ij}^{\mathcal{K}}, C_{ij}^{\mathcal{K}}, U_{ij}^{\mathcal{K}}, F_{ij}^{\mathcal{K}})]$. The definition of the transpose of a QNSM, denoted by $\mathcal{K}^T$, is $[(T_{ji}^{\mathcal{K}}, C_{ji}^{\mathcal{K}}, U_{ji}^{\mathcal{K}}, F_{ji}^{\mathcal{K}})]$.

Definition 4.5: Addition of QNS Block Matrices:

Let's consider two QNS soft block matrices, $\mathcal{K}$ and $\mathcal{L}$, where:

$$\mathcal{K} = \begin{bmatrix} \mathcal{R}_{11} & \vdots & \mathcal{R}_{12} \\ \cdots & \cdots & \cdots \\ \mathcal{R}_{21} & \vdots & \mathcal{R}_{22} \end{bmatrix} \text{ and } \mathcal{L} = \begin{bmatrix} \mathcal{Q}_{11} & \vdots & \mathcal{Q}_{12} \\ \cdots & \cdots & \cdots \\ \mathcal{Q}_{21} & \vdots & \mathcal{Q}_{22} \end{bmatrix}$$

When the corresponding blocks within matrices $\mathcal{K}$ and $\mathcal{L}$ are compatible for addition, The following defines the addition operation for QNS block matrices: $\mathcal{K} + \mathcal{L} = \begin{bmatrix} \mathcal{R}_{11} + \mathcal{Q}_{11} & \vdots & \mathcal{R}_{12} + \mathcal{Q}_{12} \\ \cdots & \cdots & \cdots \\ \mathcal{R}_{21} + \mathcal{Q}_{21} & \vdots & \mathcal{R}_{22} + \mathcal{Q}_{22} \end{bmatrix}$

Example :

$$\mathcal{K} = \begin{bmatrix} (0.6,0.5,0.4,0.3) & (0.5,0.6,0.7,0.2) & \vdots & (0.7,0.5,0.4,0.3) & (0.2,0.3,0.5,0.6) \\ (0.4,0.3,0.2,0.6) & (0.6,0.5,0.4,0.3) & \vdots & (0.5,0.6,0.7,0.2) & (0.7,0.5,0.4,0.3) \\ \cdots & \cdots & \vdots & \cdots & \cdots \\ (0.3,0.2,0.6,0.1) & (0.4,0.3,0.2,0.6) & \vdots & (0.6,0.5,0.4,0.3) & (0.5,0.6,0.7,0.2) \\ (0.6,0.7,0.8,0.2) & (0.3,0.2,0.6,0.1) & \vdots & (0.4,0.3,0.2,0.6) & (0.6,0.5,0.4,0.3) \end{bmatrix}$$

$$\mathcal{L} = \begin{bmatrix} (0.8,0.5,0.3,0.4) & (0.8,0.6,0.7,0.2) & \vdots & (0.6,0.8,0.7,0.2) & (0.5,0.6,0.2,0.6) \\ (0.1,0.5,0.6,0.6) & (0.5,0.4,0.3,0.9) & \vdots & (0.2,0.3,0.4,0.5) & (0.7,0.5,0.8,0.3) \\ \cdots & \cdots & \vdots & \cdots & \cdots \\ (0.3,0.4,0.7,0.7) & (0.4,0.3,0.8,0.7) & \vdots & (0.4,0.8,0.5,0.3) & (0.4,0.8,0.1,0.5) \\ (0.5,0.9,0.9,0.6) & (0.3,0.8,0.7,0.6) & \vdots & (0.4,0.5,0.8,0.5) & (0.7,0.5,0.1,0.8) \end{bmatrix}$$

$$\mathcal{K} + \mathcal{L} = \begin{bmatrix} (0.8,0.5,0.3,0.3) & (0.5,0.6,0.7,0.2) & \vdots & (0.6,0.5,0.7,0.2) & (0.2,0.6,0.2,0.6) \\ (0.4,0.3,0.2,0.6) & (1.1,0.9,0.7,1.2) & \vdots & (0.7,0.9,1.1,0.7) & (1.4,1.0,1.2,0.6) \\ \cdots & \cdots & \vdots & \cdots & \cdots \\ (0.3,0.2,0.6,0.1) & (0.4,0.3,0.2,0.7) & \vdots & (0.4,0.5,0.5,0.3) & (0.4,0.8,0.1,0.2) \\ (0.6,0.7,0.8,0.2) & (0.3,0.2,0.6,0.6) & \vdots & (0.4,0.3,0.8,0.5) & (0.6,0.5,0.1,0.3) \end{bmatrix}$$

Definition 4.6: Multiplication of QNS Block Matrices

Consider two QNS block matrices, denoted by $\mathcal{K}$ and $\mathcal{L}$, which can be expressed as follows: $\mathcal{K} =$

$$\begin{bmatrix} \mathcal{R}_{11} & \vdots & \mathcal{R}_{12} \\ \cdots & \cdots & \cdots \\ \mathcal{R}_{21} & \vdots & \mathcal{R}_{22} \end{bmatrix} \text{ and } \mathcal{L} = \begin{bmatrix} \mathcal{Q}_{11} & \vdots & \mathcal{Q}_{12} \\ \cdots & \cdots & \cdots \\ \mathcal{Q}_{21} & \vdots & \mathcal{Q}_{22} \end{bmatrix}, \text{ The product of two QNS block matrices is denoted as } \mathcal{KL},$$

defined by (assuming the mutual compatibility of the blocks for multiplication):

$$\mathcal{KL} = \begin{bmatrix} \mathcal{R}_{11}\mathcal{Q}_{11} + \mathcal{R}_{12}\mathcal{Q}_{21} & \vdots & \mathcal{R}_{11}\mathcal{Q}_{12} + \mathcal{R}_{12}\mathcal{Q}_{22} \\ \cdots & \cdots & \cdots \\ \mathcal{R}_{21}\mathcal{Q}_{11} + \mathcal{R}_{22}\mathcal{Q}_{21} & \vdots & \mathcal{R}_{21}\mathcal{Q}_{12} + \mathcal{R}_{22}\mathcal{Q}_{22} \end{bmatrix}$$

Definition 4.7: Transpose of a QNS Block Matrix

Consider a QNS block matrix $\mathcal{K}$, defined as, $\mathcal{K} = \begin{bmatrix} \mathcal{R}_{11} & \mathcal{R}_{12} \\ \mathcal{R}_{21} & \mathcal{R}_{22} \end{bmatrix}$. The transpose of $\mathcal{K}$, denoted as $\mathcal{K}^T$,

$$\text{is given by } \mathcal{K}^T = \begin{bmatrix} \mathcal{R}_{11}^T & \mathcal{R}_{12}^T \\ \mathcal{R}_{21}^T & \mathcal{R}_{22}^T \end{bmatrix}$$

Definition 4.8: QNS block triangular matrix

A QNS block triangular matrix is a specific kind of square QNSM. It can be categorized into two types: 1. QNS block upper triangular 2. QNS block lower triangular.

Definition 4.8.1: QNS Block upper triangular matrix

Consider the square QNS matrices F, G, and H, Then Next, a QNS matrices with a certain structure becomes known to represent the QNS block upper triangular matrices: This matrix is formatted as follows: $\mathcal{K} = \begin{bmatrix} F & G \\ 0 & H \end{bmatrix}$

Definition 4.8.2 : QNS Block lower triangular matrix

Consider the square QNS matrices F, G, and H, then QNS block lower triangular matrix is defined as a QNS matrix that takes the form: $\mathcal{K} = \begin{bmatrix} F & 0 \\ G & H \end{bmatrix}$

Properties

- Another QNS block upper triangular matrix can be generated through the addition of two QNS block upper triangular matrices of the identical order.
- A QNS block upper triangular matrix can be generated by multiplying two QNS block upper triangular matrices.
- A QNS block lower triangular matrix can be formed through the addition of two QNS block lower triangular matrices.
- Another QNS block lower triangular matrix is created whenever two QNS block lower triangular matrices of the identical order are multiplied

Definition 4.9: QNS block diagonal matrix

A diagonal matrix of the QNS block type is a square matrix with zero QNSM in all off-diagonal blocks and square QNSM in the main diagonal blocks. A QNS block diagonal matrix A has the subsequent form :

$$\mathcal{K} = \begin{bmatrix} \mathcal{K}_{11} & 0 & \cdots & 0 \\ 0 & \mathcal{K}_{22} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \mathcal{K}_{nn} \end{bmatrix}$$

Each $\mathcal{A}_{ij}$ represents a square QNS block matrix, where i and j range from 1 to n.

Definition 4.10: QNS block quasidiagonal matrix :

This matrix is a QNS block. The Square QNSM of various orders constitutes its diagonal blocks. Zero QNS block matrices comprise all of the off-diagonal blocks. Thus,

$$\mathcal{K} = \begin{bmatrix} \mathcal{D}_1 & 0 & \cdots & 0 \\ 0 & \mathcal{D}_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \mathcal{D}_n \end{bmatrix}$$

This is a QNS quasidiagonal matrix composed of diagonal blocks denoted as D_i , where i ranges from 1 to n . Each D_i is a square QNSM with a unique order

Definition 4.11: QNS block tridiagonal matrix

Assume that the square QNS block matrices A , B , and C exist in the lower, main, and upper diagonals, respectively. One particular kind of QNS block matrix is a tridiagonal matrix. It is a square QNSM, like a QNS block diagonal matrix. Square QNSM are arranged along the upper, middle, and lower diagonals to form its structure. The matrix's other blocks are all zero QNS. The form of a QNS block tridiagonal matrix, symbolised by the letter A , is as follows:

$$\begin{bmatrix} A_1 & B_1 & 0 & \cdots & 0 \\ C_1 & A_2 & B_2 & 0 & \vdots \\ 0 & C_2 & A_3 & B_3 & \vdots \\ \vdots & 0 & C_3 & \ddots & B_{n-1} \\ 0 & \cdots & \cdots & C_{n-1} & A_n \end{bmatrix}$$

Properties :

When two QNS block tridiagonal matrices of equal order are added together, the resulting matrix is also a QNS block tridiagonal matrix.

$$\mathcal{K} = \begin{bmatrix} A_1 & B_1 & 0 & \cdots & 0 \\ C_1 & A_2 & B_2 & 0 & \vdots \\ 0 & C_2 & A_3 & B_3 & \vdots \\ \vdots & 0 & C_3 & \ddots & B_{n-1} \\ 0 & \cdots & \cdots & C_{n-1} & A_n \end{bmatrix}, \mathcal{L} = \begin{bmatrix} \mathcal{D}_1 & \mathcal{F}_1 & 0 & \cdots & 0 \\ \mathcal{E}_1 & \mathcal{D}_2 & \mathcal{F}_2 & 0 & \vdots \\ 0 & \mathcal{E}_2 & \mathcal{D}_3 & \mathcal{F}_3 & \vdots \\ \vdots & 0 & \mathcal{E}_3 & \ddots & \mathcal{F}_{n-1} \\ 0 & \cdots & \cdots & \mathcal{E}_{n-1} & \mathcal{D}_n \end{bmatrix}$$

Based on the definition of the addition of two QNS block matrices, it follows that

$$\mathcal{K} + \mathcal{L} = \begin{bmatrix} A_1 + \mathcal{D}_1 & B_1 + \mathcal{F}_1 & 0 & \cdots & 0 \\ C_1 + \mathcal{E}_1 & A_2 + \mathcal{D}_2 & B_2 + \mathcal{F}_2 & 0 & \vdots \\ 0 & C_2 + \mathcal{E}_2 & A_3 + \mathcal{D}_3 & B_3 + \mathcal{F}_3 & \vdots \\ \vdots & 0 & C_3 + \mathcal{E}_3 & \ddots & B_{n-1} + \mathcal{F}_{n-1} \\ 0 & \cdots & \cdots & C_{n-1} + \mathcal{E}_{n-1} & A_n + \mathcal{D}_n \end{bmatrix}$$

Note :

Transpose of QNS block tridiagonal matrix is again a QNS tridiagonal matrix.

Example:

$$\text{If } \mathcal{K} = \begin{bmatrix} A_1 & B_1 & 0 & \cdots & 0 \\ C_1 & A_2 & B_2 & 0 & \vdots \\ 0 & C_2 & A_3 & B_3 & \vdots \\ \vdots & 0 & C_3 & \ddots & B_{n-1} \\ 0 & \cdots & \cdots & C_{n-1} & A_n \end{bmatrix} \text{ then } \mathcal{K}^T = \begin{bmatrix} A_1 & C_1 & 0 & \cdots & 0 \\ B_1 & A_2 & C_2 & 0 & \vdots \\ 0 & B_2 & A_3 & C_3 & \vdots \\ \vdots & 0 & B_3 & \ddots & C_{n-1} \\ 0 & \cdots & \cdots & B_{n-1} & A_n \end{bmatrix}$$

Definition 4.12: QNS block toeplitz matrix

A QNS block Toeplitz matrix is a specific type of QNS block matrix. It features repeating blocks along its diagonals. Furthermore, each individual block element ($\mathcal{K}_{ij}$) within the matrix must also be a Toeplitz matrix. A QNS block Toeplitz matrix, denoted by $\mathcal{K}$, possesses the following format:

$$\mathcal{K} = \begin{bmatrix} \mathcal{K}_{11} & \mathcal{K}_{12} & \mathcal{K}_{13} & \mathcal{K}_{14} \\ \mathcal{K}_{21} & \mathcal{K}_{11} & \mathcal{K}_{12} & \mathcal{K}_{13} \\ \mathcal{K}_{31} & \mathcal{K}_{21} & \mathcal{K}_{11} & \mathcal{K}_{12} \\ \mathcal{K}_{41} & \mathcal{K}_{31} & \mathcal{K}_{21} & \mathcal{K}_{11} \end{bmatrix}$$

In this context, the matrices $\mathcal{K}_{ij}$ are all square QNS block matrices.

Properties

1. QNS block Toeplitz matrices are closed under addition, meaning their sum will always be another QNS block Toeplitz matrix, if and only if the matrices have compatible dimensions.
2. QNS block Toeplitz matrices possess the property that their transpose is also a QNS block Toeplitz matrix.
3. If $\mathcal{K}$ and $\mathcal{L}$ are two QNS soft block Toeplitz matrices, Consequently, the sum of their transposes equals the transpose of their sum., i.e., $(\mathcal{K} + \mathcal{L})^T = \mathcal{K}^T + \mathcal{L}^T$.

Let

$$\mathcal{K} = \begin{bmatrix} \mathcal{K}_{11} & \mathcal{K}_{12} & \mathcal{K}_{13} & \mathcal{K}_{14} \\ \mathcal{K}_{21} & \mathcal{K}_{11} & \mathcal{K}_{12} & \mathcal{K}_{13} \\ \mathcal{K}_{31} & \mathcal{K}_{21} & \mathcal{K}_{11} & \mathcal{K}_{12} \\ \mathcal{K}_{41} & \mathcal{K}_{31} & \mathcal{K}_{21} & \mathcal{K}_{11} \end{bmatrix} \text{ and } \mathcal{L} = \begin{bmatrix} \mathcal{L}_{11} & \mathcal{L}_{12} & \mathcal{L}_{13} & \mathcal{L}_{14} \\ \mathcal{L}_{21} & \mathcal{L}_{11} & \mathcal{L}_{12} & \mathcal{L}_{13} \\ \mathcal{L}_{31} & \mathcal{L}_{21} & \mathcal{L}_{11} & \mathcal{L}_{12} \\ \mathcal{L}_{41} & \mathcal{L}_{31} & \mathcal{L}_{21} & \mathcal{L}_{11} \end{bmatrix}$$

$$\mathcal{K} + \mathcal{L} = \begin{bmatrix} \mathcal{K}_{11} + \mathcal{L}_{11} & \mathcal{K}_{12} + \mathcal{L}_{12} & \mathcal{K}_{13} + \mathcal{L}_{13} & \mathcal{K}_{14} + \mathcal{L}_{14} \\ \mathcal{K}_{21} + \mathcal{L}_{21} & \mathcal{K}_{11} + \mathcal{L}_{11} & \mathcal{K}_{12} + \mathcal{L}_{12} & \mathcal{K}_{13} + \mathcal{L}_{13} \\ \mathcal{K}_{31} + \mathcal{L}_{31} & \mathcal{K}_{21} + \mathcal{L}_{21} & \mathcal{K}_{11} + \mathcal{L}_{11} & \mathcal{K}_{12} + \mathcal{L}_{12} \\ \mathcal{K}_{41} + \mathcal{L}_{41} & \mathcal{K}_{31} + \mathcal{L}_{31} & \mathcal{K}_{21} + \mathcal{L}_{21} & \mathcal{K}_{11} + \mathcal{L}_{11} \end{bmatrix}$$

$$\mathcal{K}^T = \begin{bmatrix} \mathcal{K}_{11} & \mathcal{K}_{21} & \mathcal{K}_{31} & \mathcal{K}_{41} \\ \mathcal{K}_{12} & \mathcal{K}_{11} & \mathcal{K}_{21} & \mathcal{K}_{31} \\ \mathcal{K}_{13} & \mathcal{K}_{12} & \mathcal{K}_{11} & \mathcal{K}_{21} \\ \mathcal{K}_{14} & \mathcal{K}_{13} & \mathcal{K}_{12} & \mathcal{K}_{11} \end{bmatrix}, \text{ and } \mathcal{L}^T = \begin{bmatrix} \mathcal{L}_{11} & \mathcal{L}_{21} & \mathcal{L}_{31} & \mathcal{L}_{41} \\ \mathcal{L}_{12} & \mathcal{L}_{11} & \mathcal{L}_{21} & \mathcal{L}_{31} \\ \mathcal{L}_{13} & \mathcal{L}_{12} & \mathcal{L}_{11} & \mathcal{L}_{21} \\ \mathcal{L}_{14} & \mathcal{L}_{13} & \mathcal{L}_{12} & \mathcal{L}_{11} \end{bmatrix}$$

$$(\mathcal{K} + \mathcal{L})^T = \begin{bmatrix} \mathcal{K}_{11} + \mathcal{L}_{11} & \mathcal{K}_{21} + \mathcal{L}_{21} & \mathcal{K}_{31} + \mathcal{L}_{31} & \mathcal{K}_{41} + \mathcal{L}_{41} \\ \mathcal{K}_{12} + \mathcal{L}_{12} & \mathcal{K}_{11} + \mathcal{L}_{11} & \mathcal{K}_{21} + \mathcal{L}_{21} & \mathcal{K}_{31} + \mathcal{L}_{31} \\ \mathcal{K}_{13} + \mathcal{L}_{13} & \mathcal{K}_{12} + \mathcal{L}_{12} & \mathcal{K}_{11} + \mathcal{L}_{11} & \mathcal{K}_{21} + \mathcal{L}_{21} \\ \mathcal{K}_{14} + \mathcal{L}_{14} & \mathcal{K}_{13} + \mathcal{L}_{13} & \mathcal{K}_{12} + \mathcal{L}_{12} & \mathcal{K}_{11} + \mathcal{L}_{11} \end{bmatrix} = \mathcal{K}^T + \mathcal{L}^T$$

4. Let $\mathcal{A}_i$ represent an arbitrary collection of $n \times n$ matrices. One kind of QNS block matrix that has a specific arrangement defined by its circulant form is called a QNS block circulant matrix.

$$c\mathcal{K} = \begin{bmatrix} \mathcal{K}_0 & \mathcal{K}_1 & \dots & \mathcal{K}_{n-1} \\ \mathcal{K}_{n-1} & \mathcal{K}_0 & \dots & \mathcal{K}_{n-2} \\ \vdots & \vdots & \ddots & \vdots \\ \mathcal{K}_1 & \mathcal{K}_2 & \dots & \mathcal{K}_0 \end{bmatrix}$$

Note :

Both the product (AB) and the sum (A+B) of two QNS block circulant matrices, A and B, are likewise QNS block circulant matrices. Furthermore, for block circulant matrices in general, the product is commutative, meaning $AB = BA$.

5. Suppose we have square QNS block matrices denoted by $\mathcal{K}_{11}, \mathcal{K}_{22}, \mathcal{K}_{33}, \dots, \mathcal{K}_{qq}$, with orders $n_1, n_2, n_3, \dots, n_q$, respectively.

Then, the matrix formed as follows:

$$\text{diag}(\mathcal{K}_{11}, \mathcal{K}_{22}, \mathcal{K}_{33}, \dots, \mathcal{K}_{qq}) = \begin{bmatrix} \mathcal{K}_{11} & 0 & \dots & 0 \\ 0 & \mathcal{K}_{22} & \dots & 0 \\ \vdots & \vdots & \dots & \vdots \\ 0 & 0 & \dots & \mathcal{K}_{qq} \end{bmatrix} \text{ where the zeros represent zero matrices}$$

of appropriate sizes, is referred to as the square QNS block matrices' direct sum

$\mathcal{K}_{11}, \mathcal{K}_{22}, \mathcal{K}_{33}, \dots, \mathcal{K}_{qq}$. This direct sum is expressed as:

$\mathcal{K}_{11} \oplus \mathcal{K}_{22} \oplus \mathcal{K}_{33} \oplus \dots \oplus \mathcal{K}_{qq}$ and has the order $(n_1 + n_2 + n_3 + \dots + n_q)$.

QNS block matrices satisfy the following algebraic properties:

6. Commutativity:

Let $\mathcal{K}$ and $\mathcal{L}$ be two diagonal QNS block matrices.

The direct sum of $\mathcal{K}$ and $\mathcal{L}$ is represented as: $\mathcal{K} \oplus \mathcal{L} = \begin{bmatrix} \mathcal{K} & \vdots & 0 \\ \dots & \dots & \dots \\ 0 & \vdots & \mathcal{L} \end{bmatrix}$

and the direct sum of $\mathcal{L}$ and $\mathcal{K}$ is represented as: $\mathcal{L} \oplus \mathcal{K} = \begin{bmatrix} \mathcal{L} & \vdots & 0 \\ \dots & \dots & \dots \\ 0 & \vdots & \mathcal{K} \end{bmatrix}$

Thus, $\mathcal{K} \oplus \mathcal{L} \neq \mathcal{L} \oplus \mathcal{K}$.

Therefore, we can conclude that, in QNS block matrices, the direct sum is not commutative.

7. Associativity:

Consider three square QNS block matrices, $\mathcal{K}$, $\mathcal{L}$ and $\mathcal{M}$. As previously determined, $\mathcal{N}$ can be used to represent the direct sum of $\mathcal{K}$ and $\mathcal{L}$:

$$\mathcal{K} \oplus \mathcal{L} = \begin{bmatrix} \mathcal{K} & \vdots & 0 \\ \dots & \dots & \dots \\ 0 & \vdots & \mathcal{L} \end{bmatrix} = \mathcal{N}(\text{say})$$

Therefore,

$$(\mathcal{K} \oplus \mathcal{L}) \oplus \mathcal{M} = \mathcal{N} \oplus \mathcal{M} = \begin{bmatrix} \mathcal{N} & \vdots & 0 \\ \dots & \dots & \dots \\ 0 & \vdots & \mathcal{M} \end{bmatrix} = \begin{bmatrix} \mathcal{K} & 0 & 0 \\ 0 & \mathcal{L} & 0 \\ 0 & 0 & \mathcal{M} \end{bmatrix}$$

where $\mathcal{K} \oplus \mathcal{L} = \mathcal{N}$.

Again,

$$\mathcal{K} \oplus (\mathcal{L} \oplus \mathcal{M}) = \mathcal{K} \oplus \mathcal{P} = \begin{bmatrix} \mathcal{K} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{P} \end{bmatrix} = \begin{bmatrix} \mathcal{K} & 0 & 0 \\ 0 & \mathcal{L} & 0 \\ 0 & 0 & \mathcal{M} \end{bmatrix}$$

where $\mathcal{L} \oplus \mathcal{M} = \mathcal{P}$.

Hence, the associative law holds for QNS block matrices.

8. Mixed sum of QNS block matrices

Consider $\mathcal{K}$, $\mathcal{L}$, $\mathcal{M}$, and $\mathcal{N}$ be four addition-compatible QNS block matrices. The following equations can be obtained using the definitions of addition and the direct sum of QNS block matrices:

$$\text{The direct sum of matrices } \mathcal{K} \text{ and } \mathcal{M} \text{ is represented as: } \mathcal{K} \oplus \mathcal{M} = \begin{bmatrix} \mathcal{K} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{M} \end{bmatrix}$$

$$\text{The direct sum of matrices } \mathcal{L} \text{ and } \mathcal{N} \text{ is represented as: } \mathcal{L} \oplus \mathcal{N} = \begin{bmatrix} \mathcal{L} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{N} \end{bmatrix}$$

The sum of the direct sum of ($\mathcal{K}$ and $\mathcal{M}$) and the direct sum of ($\mathcal{L}$ and $\mathcal{N}$) is equivalent to:

$$(\mathcal{K} \oplus \mathcal{M}) \oplus (\mathcal{L} \oplus \mathcal{N}) = \begin{bmatrix} \mathcal{K} + \mathcal{L} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{M} + \mathcal{N} \end{bmatrix}$$

Consequently, the following result holds true:

$$(\mathcal{K} \oplus \mathcal{M}) \oplus (\mathcal{L} \oplus \mathcal{N}) = (\mathcal{K} \oplus \mathcal{L}) \oplus (\mathcal{M} \oplus \mathcal{N}).$$

9. Multiplication of direct sum of QNS block matrices

The mixed multiplication of the direct sum of four QNS block matrices, $\mathcal{K}$, $\mathcal{L}$, $\mathcal{M}$, and $\mathcal{N}$ that are compatible with addition and multiplication is provided by

$$(\mathcal{K} \oplus \mathcal{L})(\mathcal{M} \oplus \mathcal{N}) = (\mathcal{KL} \oplus \mathcal{MN}).$$

Based on the definition of the direct sum and the multiplication of QNS block matrices, we have the following relationship:

$$(\mathcal{K} \oplus \mathcal{L})(\mathcal{M} \oplus \mathcal{N}) = \begin{bmatrix} \mathcal{K} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{L} \end{bmatrix} \begin{bmatrix} \mathcal{M} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{N} \end{bmatrix} = \begin{bmatrix} \mathcal{KM} & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{LN} \end{bmatrix} = (\mathcal{KM} \oplus \mathcal{LN})$$

If A and B are two QNS block matrices, then the transpose of the direct sum of A and B is

$$(\mathcal{K} \oplus \mathcal{M})^T = \begin{bmatrix} \mathcal{K}^T & \vdots & 0 \\ \cdots & \cdots & \cdots \\ 0 & \vdots & \mathcal{M}^T \end{bmatrix} = (\mathcal{K}^T \oplus \mathcal{M}^T)$$

5. Conclusion :

In this study, we analyzed several classes of Quaternionic Near Semi- (QNS) block matrices, such as triangular, tridiagonal, quasideagonal, circulant, and Toeplitz types. We investigated how

various matrix operations apply to these forms and found that many of the resulting properties align closely with those of traditional matrix structures. These results enhance the theoretical framework surrounding QNS matrices and suggest a strong foundation for further mathematical development. As a next step, we intend to explore determinant calculations for QNS block matrices, which may reveal deeper structural insights and broaden their potential applications in both pure and applied mathematics.

References

1. **Zadeh, L.A.** Fuzzy Sets, *Information and Control*, 1965; 8, 338-353.
2. **Atanassov, K.T.** Intuitionistic Fuzzy Sets, *Fuzzy Sets and Systems*, 1986; 20(1), 87-96.
3. **Molodtsov, D.A.** Soft Set Theory First Results, *Computers and Mathematics with Applications*, 1999; 37, 19-31.
4. **Maji, P.K.; Biswas, R.; Roy, A.R.** Fuzzy Soft Sets, *Journal of Fuzzy Mathematics*, 2001; 9(3), 589-602.
5. **Maji, P.K.; Biswas, R.; Roy, A.R.** Soft Set Theory, *Computers and Mathematics with Applications*, 2003; 45, 555-562.
6. **Maji, P.K.; Biswas, R.; Roy, A.R.** Intuitionistic Fuzzy Soft Sets, *Journal of Fuzzy Mathematics*, 2004; 12, 669-683.
7. **Smarandache, F.** Neutrosophic Set – A Generalisation of the Intuitionistic Fuzzy Sets, *International Journal of Pure and Applied Mathematics*, 2005; 24, 287-297.
8. **Maji, P.K.** Neutrosophic Soft Set, *Annals of Fuzzy Mathematics and Informatics*, 2013; 5(1), 157-168.
9. **Maji, P.K.** A Neutrosophic Soft Set Approach to Decision Making Problem, *Annals of Fuzzy Mathematics and Informatics*, 2013; 3(2), 313–319.
10. **Deli, I.; Broumi, S.** Neutrosophic Soft Matrices and NSM-Decision Making, *Journal of Intelligent and Fuzzy Systems*, 2015; 28(5), 2233-2241.
11. **Deli, I.** Interval-Valued Neutrosophic Soft Sets and Its Decision Making, *International Journal of Machine Learning and Cybernetics*, 2015; DOI: 10.1007/s13042-015-0461-3.
12. **Deli, I.; Broumi, S.** Neutrosophic Soft Relations and Some Properties, *Annals of Fuzzy Mathematics and Informatics*, 2015; 9(1), 169-182.
13. **Broumi, S.; Smarandache, F.** Intuitionistic Neutrosophic Soft Set, *Journal of Information and Computing Science*, 2013; 8(2), 130-140.
14. **Bera, T.; Mahapatra, N.K.** On Neutrosophic Soft Function, *Annals of Fuzzy Mathematics and Informatics*, 2016; 12(1), 101-119.
15. **Das, S.; Kumar, S.; Kar, S.; Pal, T.** Group Decision Making Using Neutrosophic Soft Matrix: An Algorithmic Approach, *Journal of King Saud University - Computer and Information Sciences*, 2017; <https://doi.org/10.1016/j.jksuci.2017.05.001>.

16. **Deli, I.** Npn-Soft Sets Theory and Applications, *Annals of Fuzzy Mathematics and Informatics*, 2015; 10(6), 847-862.
17. **Deli, I.** Linear Optimization Method on Single Valued Neutrosophic Set and Its Sensitivity Analysis, *TWMS Journal of Applied and Engineering Mathematics*, 2020; 10(1), 128-137.
18. **Deli, I.; Eraslan, S.; Çağman, N.** Ivnpiv-Neutrosophic Soft Sets and Their Decision Making Based on Similarity Measure, *Neural Computing and Applications*, 2018; 29(1), 187–203. DOI: 10.1007/s00521-016-2428-z.
19. **Broumi, S.; Deli, I.; Smarandache, F.** Distance and Similarity Measures of Interval Neutrosophic Soft Sets, *Critical Review*, Center for Mathematics of Uncertainty, Creighton University, 2014; 8, 14-31.
20. **Deli, I.; Toktaş, Y.; Broumi, S.** Neutrosophic Parameterized Soft Relations and Their Applications, *Neutrosophic Sets and Systems*, 2014; 4, 25-34.
21. **Broumi, S.; Deli, I.; Smarandache, F.** Interval Valued Neutrosophic Parameterized Soft Set Theory and Its Decision Making, *Journal of New Results in Science*, 2014; 7, 58-71.
22. **Deli, I.; Broumi, S.; Ali, M.** Neutrosophic Soft Multi-Set Theory and Its Decision Making, *Neutrosophic Sets and Systems*, 2014; 5, 65-76.
23. **Uma, R.; Murugdas, P.; Sriram, S.** Fuzzy Neutrosophic Soft Block Matrices and Its Some Properties, *International Journal of Fuzzy Mathematical Archive*, 2017; 14(2), 287-303.
24. **Dhar, M.** Neutrosophic Soft Block Matrices and Some of Its Properties, *International Journal of Neutrosophic Science*, 2020; 12(1), 19-49.
25. **Mamoni Dhar.** Some Basic Concepts of Neutrosophic Soft Block Matrices, *Neutrosophic Sets and Systems*, 2022; Vol. 51.
26. **Ramesh Kumar S.; Stanis Arul Mary A.** Quadri Partitioned Neutrosophic Soft Set, *IRJASH*, 2021; 3.
27. **Ramesh Kumar S.; Stanis Arul Mary A.** Quadri Partitioned Neutrosophic Soft Topological Space, *IRJASH*, 2021; 2(6).

Received: Dec. 8, 2024. Accepted: June 26, 2025