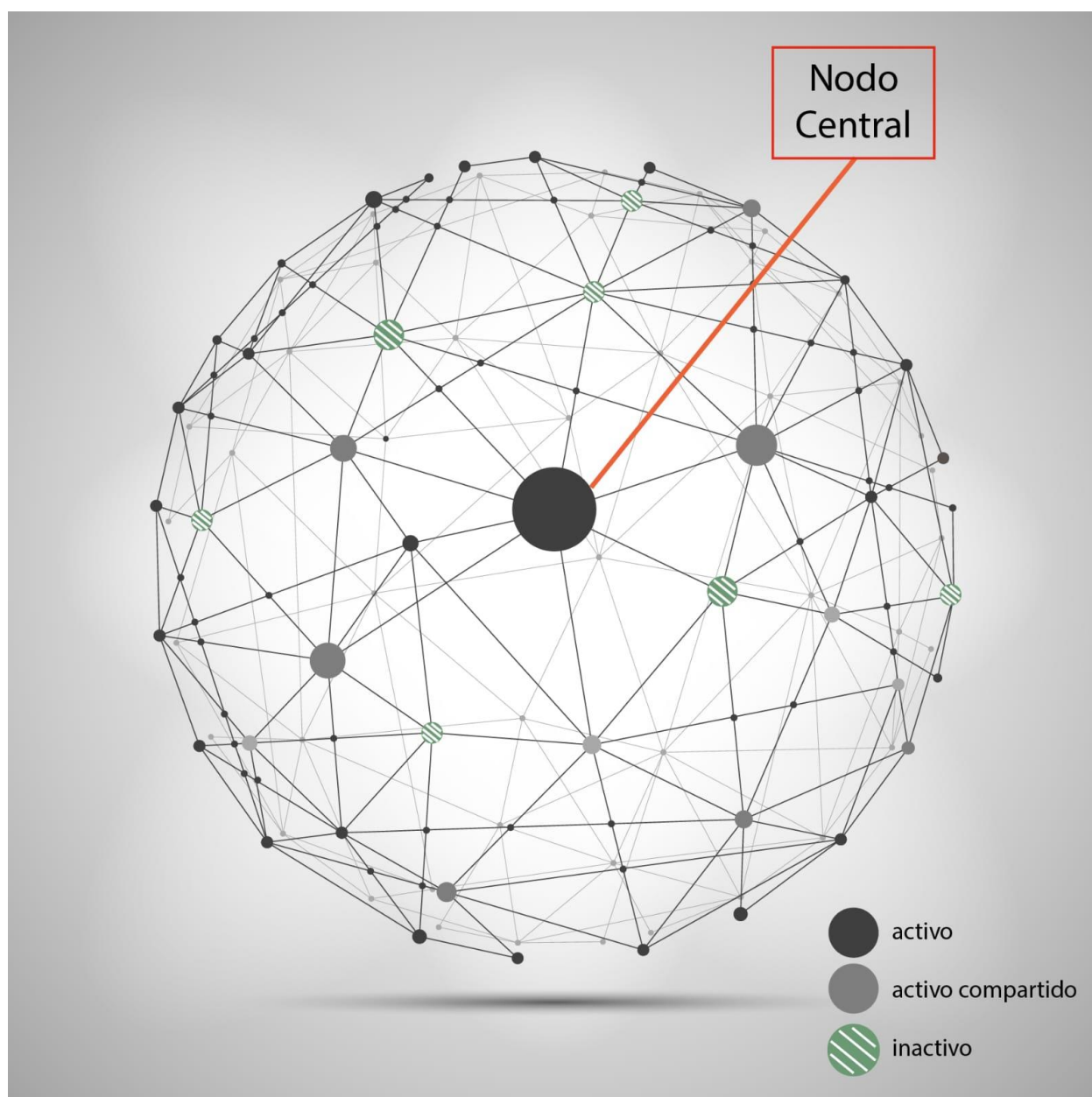


UNIVERSIDAD NACIONAL DE PILAR – FACULTAD DE CIENCIAS APLICADAS
GRUPO DE INVESTIGACION DE SISTEMAS DISTRIBUIDOS

Tesis de Maestría en Informática y Computación



UNIVERSIDAD NACIONAL DE PILAR

FACULTAD DE CIENCIAS APLICADAS

Prof. Mgtr. Jorge Forneron Martínez, Decano

Prof. Mgtr. Miguel Angel Delpino Aguayo, Vicedecano

Miembro del Consejo Directivo

Dra. Luisa Gamarra, consejera Docente

Lic. Luis Alberto Fossati Dávalos, consejero Docente

Ing. Rubén Darío Fornerón Portillo, consejero Docente

Lic. Juan Mora, consejero Docente

MSc. Miguel Ángel Benítez, consejero Docente

Lic. Sonia Escobar, consejera no Docente

Est. Rolando Concepción Vázquez Cuenca, representante Estudiantil

Est. Malena del Pilar Gamarra Gaona, representante Estudiantil

Editor Responsable

Dra. Nélida Soria Rey

Comité Editorial

Dra. Lourdez Sánchez de Velaustegui

MSc. María Elena López de Silva

Mgtr. Miguel Ángel Delpino Aguayo



Presentación

Continuamos el proceso de publicación de los trabajos de las tesis aprobadas en el contexto del programa de Maestría en Informática y Computación, presentadas por integrantes que conforman el Grupo de Investigación en Sistemas Distribuidos de la Facultad de Ciencias Aplicadas.

En el presente número especial, que es el segundo de la serie, se publica la tesis del Magíster **Diego David Duré Attis**, referido al tema “Imputación de datos en la gestión de recursos compartidos en sistemas distribuidos”, desarrollado con el propósito de contribuir al conocimiento y/o aportar soluciones innovadoras al problema planteado.

A handwritten signature in blue ink, reading "Nélide Soria", with a horizontal line underneath the name.

Dra. Nélide Soria Rey

Editora

Editorial

Los sistemas distribuidos, compuestos por múltiples nodos independientes interconectados, plantean el desafío fundamental de asegurar una asignación de recursos que cumpla estrictamente con la modalidad de exclusión mutua, garantizando la disponibilidad de los recursos y manteniendo la integridad de las prioridades de los procesos. En este contexto, la tesis de maestría del Mgter. Diego David Duré Attis, titulada "Imputación de datos en la gestión de recursos compartidos en sistemas distribuidos," se enfoca en abordar estos desafíos.

La propuesta de Duré Attis se centra en el desarrollo de un modelo de decisión dentro de los sistemas distribuidos que incorpora técnicas de imputación de datos para la gestión de recursos compartidos. Esta investigación destaca la importancia crítica de la toma de decisiones en sistemas distribuidos, particularmente en lo que respecta al acceso a recursos compartidos.

Se centra especialmente en situaciones donde, durante los ciclos de recolección de información de gestión, el nodo central, encargado de mantener la información de control de todos los nodos, recibe datos incompletos, como los criterios de evaluación de carga nodal y las prioridades de los procesos. Estos datos faltantes representan un obstáculo significativo en la gestión de recursos. Para superar este desafío, se aplican técnicas de imputación de datos que permiten estimar los valores faltantes utilizando el método de imputación K-Means, reconocido por su fiabilidad y eficiencia computacional.

En resumen, la investigación de Diego David Duré Attis ofrece una solución innovadora y efectiva para abordar los desafíos en la gestión de recursos en sistemas distribuidos. Este trabajo tiene el potencial de inspirar a futuros investigadores y contribuir al progreso de la informática distribuida, en especial a aquellos que se dedican a este tipo de estudios en Paraguay.

Corrientes, 07 de octubre de 2023.



Mgter. Federico Agostini

UNIVERSIDAD NACIONAL DE PILAR

FACULTAD DE CIENCIAS APLICADAS



**Imputación de datos en la gestión de recursos compartidos en sistemas
distribuidos**

Diego David Duré Attis

Orientador: **Prof. Mgter. Federico Agostini**

Coorientador: **Prof. Dr. David Luis la Red Martínez**

Tesis presentada a la Facultad de Ciencias Aplicadas de la Universidad Nacional de Pilar, para la obtención del Grado de Magíster en Informática y Computación, correspondiente al Programa de Maestría en Informática y Computación.

PILAR – PARAGUAY

DICIEMBRE 2021



**Imputación de datos en la gestión de recursos compartidos en sistemas
distribuidos**

Diego David Duré Attis

Orientador: **Prof. Mgter. Federico Agostini**

Coorientador: **Prof. Dr. David Luis la Red Martínez**

Tesis presentada a la Facultad de Ciencias Aplicadas de la Universidad Nacional de Pilar, para la obtención del Grado de Magíster en Informática y Computación, correspondiente al Programa de Maestría en Informática y Computación.

PILAR – PARAGUAY

DICIEMBRE 2021

Duré Attis, Diego David. (2021). **Imputación de datos en la gestión de recursos compartidos en sistemas distribuidos**. Diego David Duré Attis. ____ p.: ____

Orientador: Prof. Mgter. Federico Agostini

Co-orientador: Prof. Dr. David Luis la Red Martínez

Tesis académica de Magíster en Informática y Computación - Facultad de Ciencias Aplicadas -
Universidad Nacional de Pilar



IMPUTACIÓN DE DATOS EN LA GESTIÓN DE RECURSOS
COMPARTIDOS EN SISTEMAS DISTRIBUIDOS

DIEGO DAVID DURÉ ATTIS

Aprobado en fecha: __/__/2021 (Res. N° __/2021 CD-FCA).

Tribunal Examinador: -----

Prof. Mgter. FEDERICO AGOSTINI

Orientador

Prof. Dr. DAVID LUIS LA RED MARTÍNEZ

Coorientador

Sello

Prof. Mgter. JORGE TOMÁS FORNERON MARTÍNEZ

Decano

Dedicatoria

- Dedico esta Tesis a mi familia, que es la fuente de motivación para la conclusión de esta maestría.
- A mi señora Luz Marisol, gracias por la paciencia y el apoyo incondicional.
- A mis hijos Luz Abigail y Thiago Enmanuel.
- A mis padres quienes me dieron la vida, la posibilidad a la educación, y me brindaron su apoyo incondicional en todo momento.

Agradecimiento

- ❖ Primeramente, agradecer al Prof. Mgter. Federico Agostini (UNNE), Orientador de Tesis, por la confianza hacia mi persona, el esfuerzo y la dedicación incondicional en el horario que fuera, desde el momento de asumir el desafío en realizar esta investigación, hasta culminarla.
- ❖ Agradecimiento muy especial al Prof. Dr. David Luis la Red Martínez (UNNE), Coorientador de Tesis y Director del Programa Maestría en Informática y Computación de la Facultad de Ciencias Aplicadas de la UNP, por sus dedicaciones, por compartir sus experiencias, conocimientos, sus consejos e indicaciones y por la gran predisposición de trabajar sin importar el horario para poder realizar este trabajo de investigación.
- ❖ A la Facultad de Ciencias Aplicadas de la Universidad Nacional de Pilar por la oportunidad y el espacio educativo brindado.
- ❖ A todos los Docentes que integran el programa de la Maestría en Informática y Computación de la Facultad de Ciencias Aplicadas de la UNP.
- ❖ A todas las personas que de alguna u otra manera hicieron posible lograr este objetivo.

Índice

Dedicatoria	v
Agradecimiento	vi
Listado de tablas	ix
Lista de figuras.....	xvi
Resumen.....	xvii
Capítulo I - Introducción.....	1
1.1. Introducción.....	1
1.2. El problema.....	4
1.3. Objetivos.....	8
1.4. Justificación	8
1.5. Hipótesis y objeto	9
1.6. Antecedentes	10
1.7. Encuadre o Marco Teórico	11
1.8. Relevancia.....	22
1.9. Metodología.....	23
1.10. Estructura de la Tesis.....	32
1.11. Discusiones y Comentarios.....	33
Capítulo II – Capa de imputación/asignación para el modelo de decisión	36
2.1. Trabajos previos.....	36
2.2. Propuesta de Solución.....	36
2.3. Descripción de la capa de imputación/asignación (E1)	40
2.3.1. Identificar el propósito del problema	41
2.3.2. Identificar las alternativas	42
2.3.3. Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa	43
2.4. Ejemplo.....	64
2.5. Evaluación	109
2.6. Discusiones y comentarios.....	109
Capítulo III - Capa de imputación/asignación para datos inexistentes - escenario 2 (E2).....	111
3.1. Trabajos previos.....	111
3.2. Propuesta de Solución.....	111
3.3. Descripción de la capa de imputación/asignación (E2)	114
3.3.1. Identificar el propósito del problema	114
3.3.2. Identificar alternativas	115
3.3.3. Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa	116
3.4. Ejemplo.....	121

3.5.	Evaluación	153
3.6.	Discusiones y comentarios.....	153
Capítulo IV – Modelo de decisión.....		155
4.1.	Introducción a los modelos de decisión.....	155
4.2.	Escenario propuesto	155
4.3.	Operador de agregación utilizado por el Modelo de Decisión.....	156
4.4.	Modelo de decisión propuesto	163
4.5.	Pautas para aplicar el método de imputación cuando se reciba de algún nodo información de control incompleta o inexistente	165
4.6.	Consideraciones del modelo de decisión propuesto	168
4.7.	Ejemplo de modelo de decisión aplicado a alguno de los algoritmos tradicionales	169
4.8.	Comparación de los resultados obtenidos con los métodos propuestos y los métodos tradicionales	198
4.9.	Consideraciones acerca de las operaciones de agregación	205
4.10.	Discusiones y comentarios.....	206
Capítulo V - Conclusiones y futuras líneas de investigación.....		208
5.1.	Conclusiones.....	208
5.2.	Futuras líneas de investigación.....	210
Referencias		211
Apéndice de publicaciones		219

Listado de tablas

Tabla 1. Valores de las características referentes a las prestaciones de cada nodo.	43
Tabla 2. Valoración de criterios para calcular la carga computacional en cada nodo.	44
Tabla 3. Categorías para medir la carga computacional en cada nodo.	48
Tabla 4. Pesos asignados a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo.	50
Tabla 5. Valoraciones asignadas a los criterios según la clasificación, para calcular las prioridades o preferencias de los procesos.	51
Tabla 6. Clasificación 1 - Variables de estados principales del nodo.	53
Tabla 7. Clasificación s - Valores del sistema operativo asignados por defecto, cuando se genera el proceso.	55
Tabla 8. Clasificación t – Valores que indican cuánto costaría al nodo hacer la asignación, el impacto esperado de atender el requerimiento (se tiene en cuenta la relación de recurso-proceso).	58
Tabla 9. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo.	60
Tabla 10. Prioridades nodales de los procesos para acceder a cada recurso.	61
Tabla 11. Pesos asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos.	62
Tabla 12. Pesos finales normalizados asignados a los procesos.	62
Tabla 13. Orden o prioridad final de asignación de los recursos a procesos.	63
Tabla 14. Procesos ejecutados en cada nodo.	64
Tabla 15. Procesos ejecutados en cada grupo.	65
Tabla 16. Recursos compartidos disponibles en cada nodo.	65
Tabla 17. Recursos solicitados por los procesos.	66
Tabla 18. Valores de las características referentes a las prestaciones de cada nodo.	68
Tabla 19. Valores de los criterios de carga computacional con valor faltante de los criterios en el nodo 1, nodo 2 y nodo 3 – (E1).	69
Tabla 20. Clasificación de valores faltantes de los criterios de carga computacional de los nodos, de acuerdo a las pautas establecidas – (E1).	69
Tabla 21. Valores históricos de los criterios de carga computacional del Nodo 1 . Primeros cinco registros de mil – (E1).	70
Tabla 22. Valores de los criterios de carga computacional con valor imputado del criterio % Uso de CPU – Nodo 1 , aplicando la pauta uno - (E1).	70
Tabla 23. Valores de los criterios de carga computacional con valor asignado del % Uso de Memoria – Nodo 2 , aplicando la pauta dos – (E1).	72
Tabla 24. Valores históricos del criterio de carga computacional del Nodo 3 . Primeros cinco registros de mil – (E1).	72
Tabla 25. Valores de los criterios de carga computacional con valor imputado del criterio % Uso de Oper. E/S – Nodo 3 , aplicando la pauta uno – (E1).	73
Tabla 26. Valores de los criterios de carga computacional en cada nodo, con valoraciones imputados/asignado – (E1).	73
Tabla 27. Valores de las categorías para medir la carga computacional en cada nodo.	74
Tabla 28. Pesos asignados a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo.	75
Tabla 29. Valores asignados a los criterios para calcular la prioridad de proceso, amputados con el mecanismo MCAR al 10% - (E1).	76
Tabla 30. Valores faltantes para la clasificación 1 – Variables de estados principales del nodo, que son las medidas que toma el sistema operativo cada vez que le llega un requerimiento - (E1).	78
Tabla 31. Clasificación de valores faltantes para la clasificación 1, de acuerdo a las pautas establecidas – (E1).	79
Tabla 32. Valores históricos de la relación proceso-recurso p₁₁r₂₃ . Primeros cinco registros de mil – (E1).	

.....	79
Tabla 33. Valores de la relación proceso-recurso $p_{11}r_{23}$ con valor imputado del criterio N° Proc., aplicando la pauta uno – (E1).	80
Tabla 34. Valores históricos de la relación proceso con otros recursos de $p_{22}r_{22}$. Primeros cinco registros de tres mil – (E1).....	81
Tabla 35. Valores de la relación proceso-recurso $p_{22}r_{22}$ con valor imputado del criterio N° Proc., aplicando la pauta dos – (E1).....	81
Tabla 36. Valores de la relación proceso-recurso $p_{25}r_{21}$ con valor asignado del criterio % Mem., aplicando la pauta tres - (E1).	82
Tabla 37. Valores históricos de la relación proceso con otros recursos de $p_{34}r_{33}$. Primeros cinco registros de siete mil - (E1).....	83
Tabla 38. Valores de la relación proceso-recurso $p_{22}r_{22}$ con valor imputado del criterio % MV, aplicando la pauta dos - (E1).	84
Tabla 39. Valores históricos de la relación proceso-recurso $p_{35}r_{24}$. Primeros cinco registros de mil - (E1).84	
Tabla 40. Valores de la relación proceso-recurso $p_{35}r_{24}$ con valor imputado del criterio % CPU, aplicando la pauta uno - (E1).	85
Tabla 41. Valor faltante para la clasificación 2 - Prioridad Proceso, son valores asignados por el sistema operativo por defecto cuando se genera el proceso - (E1).....	85
Tabla 42. Valores históricos de la relación proceso-recurso $p_{24}r_{12}$. Primeros cinco registros de mil - (E1).86	
Tabla 43. Valores de la relación proceso-recurso $p_{24}r_{12}$ con valor imputado del criterio Prioridad Proc. , aplicando la pauta uno - (E1).....	86
Tabla 44. Valores faltantes para la clasificación 3 – Valores que indican cuánto costaría al nodo hacer la asignación (se tiene en cuenta la relación de proceso-recurso) - (E1).....	87
Tabla 45. Clasificación de valores amputados para la clasificación 3 de acuerdo a las pautas establecidas - (E1).....	87
Tabla 46. Valores históricos de la relación proceso-recurso $p_{12}r_{21}$. Primeros cinco registros de mil - (E1).88	
Tabla 47. Valores de la relación proceso-recurso $p_{12}r_{21}$ con valor imputado del criterio Sobrec. Proc. , aplicando la pauta uno - (E1).....	88
Tabla 48. Valores históricos de la relación proceso-recurso $p_{13}r_{11}$. Primeros cinco registros de mil - (E1).89	
Tabla 49. Valores de la relación proceso-recurso $p_{13}r_{11}$ con valor imputado del criterio Sobrec. Proc. , aplicando la pauta uno - (E1).....	89
Tabla 50. Valores históricos de la relación proceso-recurso $p_{23}r_{32}$. Primeros cinco registros de mil - (E1).	90
Tabla 51. Valores de la relación proceso-recurso $p_{23}r_{32}$ con valor imputado del criterio Sobrec. E/S , aplicando la pauta uno - (E1).....	90
Tabla 52. Valores históricos de la relación proceso-recurso $p_{33}r_{13}$. Primeros cinco registros de mil - (E1).91	
Tabla 53. Valores de la relación proceso-recurso $p_{33}r_{13}$ con valor imputado del criterio Sobrec. E/S , aplicando la pauta uno - (E1).....	91
Tabla 54. Valores históricos de la relación proceso-recurso $p_{33}r_{22}$. Primeros cinco registros de mil - (E1).92	
Tabla 55. Valores de la relación proceso-recurso $p_{33}r_{22}$ con valor imputado del criterio Sobrec. Mem. , aplicando la pauta uno - (E1).....	92
Tabla 56. Valores de la relación proceso-recurso $p_{35}r_{12}$ con valor imputado del criterio Sobrec. E/S , aplicando la pauta dos - (E1).	94
Tabla 57. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia de procesos para cada nodo, con valores imputados/asignados - (E1).	94
Tabla 58. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia nodal que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo - (E1).....	97
Tabla 59. Prioridades nodales de los procesos para acceder a cada recurso p_{11} al p_{25} - (E1).	101
Tabla 60. Prioridades nodales de los procesos para acceder a cada recurso p_{31} al p_{37} - (E1).	102
Tabla 61. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, desde p_{11} al p_{25} - (E1).	102

Tabla 62. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, desde p_{31} al p_{37} - (E1).	103
Tabla 63. Prioridades globales finales de los procesos para acceder a cada recurso, desde p_{11} al p_{25} - (E1).	103
Tabla 64. Prioridades globales finales de los procesos para acceder a cada recurso, desde p_{31} al p_{37} - (E1).	104
Tabla 65. Prioridades globales finales para asignar los recursos, constituye la Función de Asignación para Sistemas Distribuidos (FASD) – (E1).	105
Tabla 66. Prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso, constituye la Función de Asignación para Sistemas Distribuidos Ordenada (FASDO) – (E1).	105
Tabla 67. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (segunda iteración – E1).	106
Tabla 68. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (tercera iteración – E1).	106
Tabla 69. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (cuarta iteración – E1).	106
Tabla 70. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (quinta iteración – E1).	107
Tabla 71. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (sexta iteración – E1).	107
Tabla 72. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (séptima iteración – E1).	107
Tabla 73. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (octava iteración – E1).	108
Tabla 74. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (novena iteración – E1).	108
Tabla 75. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décima iteración – E1).	108
Tabla 76. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo primera iteración – E1).	108
Tabla 77. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo segunda iteración – E1).	108
Tabla 78. Coeficientes de ponderación para diez registros históricos.	114
Tabla 79. Valoraciones de los criterios de carga computacional, con valores faltantes en su totalidad en el Nodo 1 - (E2).	123
Tabla 80. Valores históricos de los criterios de carga computacional Nodo 1, últimos cinco registros de diez - (E2).	124
Tabla 81. Valores de los criterios de carga computacional con valores imputados del Nodo 1, aplicando la pauta una - (E2).	125
Tabla 82. Valores de los criterios de carga computacional con valores asignado según las prestaciones del Nodo 1, aplicando la pauta 2 - (E2).	125
Tabla 83. Valores de los criterios de carga computacional en cada nodo, y el cálculo del promedio de cada nodo, con valores imputados en el Nodo 1, con la pauta una - (E2).	126
Tabla 84. Valores asignados a los criterios para calcular la prioridad nodal, faltando la totalidad de información de control en algunas relaciones de proceso-recurso - (E2).	127
Tabla 85. Clasificación de valores faltantes de acuerdo a las pautas establecidas, para calcular la prioridad o preferencia de los procesos - E2.	130
Tabla 86. Valores históricos de la relación proceso-recurso $p_{11}r_{23}$, últimos cinco registros de diez (E2).	130
Tabla 87. Valores imputados con la media ponderada en la relación proceso-recurso $p_{11}r_{23}$, aplicando la pauta uno (E2).	131
Tabla 88. Valores históricos de la relación proceso-recurso $p_{12}r_{21}$, últimos cinco registros de seis - (E2).	131

Tabla 89. Valores imputados con la media ponderada en la relación proceso-recurso $p_{12}r_{21}$, aplicando la pauta uno - (E2).	131
Tabla 90. Valores históricos de la relación proceso-recurso $p_{13}r_{11}$, últimos cinco registros de nueve - (E2).	132
Tabla 91. Valores imputados con la media ponderada en la relación proceso-recurso $p_{13}r_{11}$, aplicando la pauta uno - (E2).	132
Tabla 92. Valores históricos de la relación proceso-recurso $p_{22}r_{22}$, tres únicos registros - (E2).	133
Tabla 93. Valores imputados con la media ponderada en la relación proceso-recurso $p_{22}r_{22}$, aplicando la pauta uno - (E2).	133
Tabla 94. Valores históricos de la relación proceso-recurso $p_{23}r_{32}$, últimos cinco registros de ocho - (E2).	133
Tabla 95. Valores imputados con la media ponderada en la relación proceso-recurso $p_{23}r_{32}$, aplicando la pauta uno - (E2).	134
Tabla 96. Valores históricos de la relación proceso-recurso $p_{24}r_{12}$, últimos cinco registros de diez - (E2).	134
Tabla 97. Valores imputados con la media ponderada en la relación proceso-recurso $p_{24}r_{12}$, aplicando la pauta uno - (E2).	135
Tabla 98. Valores históricos de la relación proceso-recurso $p_{25}r_{21}$, dos únicos registros - (E2).	135
Tabla 99. Valores imputados con la media ponderada en la relación proceso-recurso $p_{25}r_{21}$, aplicando la pauta uno - (E2).	135
Tabla 100. Valores históricos de la relación proceso-recurso $p_{33}r_{13}$, los cinco últimos registros de siete - (E2).	136
Tabla 101. Valores imputados con la media ponderada en la relación proceso-recurso $p_{33}r_{13}$, aplicando la pauta uno - (E2).	136
Tabla 102. Valores históricos de la relación proceso-recurso $p_{33}r_{22}$, los cinco únicos registros - (E2).	137
Tabla 103. Valores imputados con la media ponderada en la relación proceso-recurso $p_{33}r_{22}$, aplicando la pauta uno - (E2).	137
Tabla 104. Valores históricos de la relación proceso-recurso $p_{34}r_{33}$, los cuatro únicos registros - (E2).	137
Tabla 105. Valores imputados con la media ponderada en la relación proceso-recurso $p_{34}r_{33}$, aplicando la pauta uno - (E2).	138
Tabla 106. Valores históricos de la relación proceso-recurso $p_{35}r_{12}$, registro único (E2).	138
Tabla 107. Valores imputados con la media ponderada en la relación proceso-recurso $p_{35}r_{12}$, aplicando la pauta uno - (E2).	138
Tabla 108. Valores asignados según las prestaciones del nodo para los criterios de $p_{35}r_{24}$, aplicando la pauta dos - (E2).	139
Tabla 109. Valoraciones asignadas a los criterios para calcular la prioridad nodal, faltando la totalidad de información de control en algunas relaciones de proceso-recurso - (E2).	139
Tabla 110. Valores asignados a los criterios para calcular la prioridad o preferencia nodal que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo, con valores imputados/asignados correspondientes al escenario E2.	142
Tabla 111. Prioridades nodales de los procesos para acceder a cada recurso p_{11} al p_{25} - (E2).	146
Tabla 112. Prioridades nodales de los procesos para acceder a cada recurso p_{31} al p_{37} - (E2).	146
Tabla 113. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos p_{11} al p_{25} - (E2).	147
Tabla 114. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos p_{31} al p_{37} - (E2).	147
Tabla 115. Prioridades globales finales de los procesos para acceder a cada recurso p_{11} al p_{25} - (E2).	148
Tabla 116. Prioridades globales finales de los procesos para acceder a cada recurso p_{31} al p_{37} - (E2).	148
Tabla 117. Prioridades globales finales para asignar los recursos, constituye la Función de Asignación para Sistemas Distribuidos (FASD) - (E2).	149
Tabla 118. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso, constituye la Función de Asignación para Sistemas Distribuidos Ordenada (FASDO) - (E2).	149

Tabla 119. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (segunda iteración) - (E2).....	150
Tabla 120. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (tercera iteración) - (E2).	150
Tabla 121. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (cuarta iteración) - (E2).	151
Tabla 122. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (quinta iteración) - (E2).	151
Tabla 123. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (sexta iteración) - (E2).....	151
Tabla 124. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (séptima iteración) - (E2).	152
Tabla 125. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (octava iteración) - (E2).	152
Tabla 126. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (novena iteración) - (E2).	152
Tabla 127. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décima iteración) - (E2).	152
Tabla 128. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo primera iteración) - (E2).....	153
Tabla 129. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo segunda iteración) - (E2).	153
Tabla 130. Descripción de las pautas establecidas para los diferentes escenarios	167
Tabla 131. Pesos asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos (método tradicional).....	169
Tabla 132. Pesos asignados a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo - (método tradicional).	170
Tabla 133. Valoraciones asignadas a los criterios para calcular la prioridad, amputados con el mecanismo MCAR al 10% - (método tradicional).	172
Tabla 134. Clasificación de valores faltantes de acuerdo a las pautas establecidas, para calcular la prioridad o preferencia de los procesos - (método tradicional).....	174
Tabla 135. Valores históricos de la relación proceso-recurso $p_{11}r_{23}$, últimos cinco registros de diez - (método tradicional).....	175
Tabla 136. Valor imputado con la media ponderada en la relación proceso-recurso $p_{11}r_{23}$, aplicando la pauta uno - (método tradicional).	175
Tabla 137. Valores históricos de la relación proceso-recurso $p_{12}r_{21}$, últimos cinco registros de seis (método tradicional).....	176
Tabla 138. Valor imputado con la media ponderada en la relación proceso-recurso $p_{12}r_{21}$, aplicando la pauta uno - (método tradicional).	176
Tabla 139. Valores históricos de la relación proceso-recurso $p_{13}r_{11}$, últimos cinco registros de nueve - (método tradicional).	177
Tabla 140. Valor imputado con la media ponderada en la relación proceso-recurso $p_{13}r_{11}$, aplicando la pauta uno - (método tradicional).	177
Tabla 141. Valores históricos de la relación proceso-recurso $p_{22}r_{22}$, tres únicos registros - (método tradicional).....	177
Tabla 142. Valor imputado con la media ponderada en la relación proceso-recurso $p_{22}r_{22}$, aplicando la pauta uno - (método tradicional).	178
Tabla 143. Valores históricos de la relación proceso-recurso $p_{23}r_{32}$, últimos cinco registros de ocho - (método tradicional).....	178
Tabla 144. Valor imputado con la media ponderada en la relación proceso-recurso $p_{23}r_{32}$, aplicando la pauta uno - (método tradicional).	178

Tabla 145. Valores históricos de la relación proceso-recurso $p_{24}r_{12}$, últimos cinco registros de diez - (método tradicional).....	179
Tabla 146. Valor imputado con la media ponderada en la relación proceso-recurso $p_{24}r_{12}$, aplicando la pauta uno - (método tradicional).....	179
Tabla 147. Valores históricos de la relación proceso-recurso $p_{25}r_{21}$, dos únicos registros - (método tradicional).....	180
Tabla 148. Valor imputado con la media ponderada en la relación proceso-recurso $p_{25}r_{21}$, aplicando la pauta uno - (método tradicional).....	180
Tabla 149. Valores históricos de la relación proceso-recurso $p_{33}r_{13}$, los cinco últimos registros de siete - (método tradicional).	180
Tabla 150. Valor imputado con la media ponderada en la relación proceso-recurso $p_{33}r_{13}$, aplicando la pauta uno - (método tradicional).....	181
Tabla 151. Valores históricos de la relación proceso-recurso $p_{33}r_{22}$, los cinco únicos registros - (método tradicional).....	181
Tabla 152. Valor imputado con la media ponderada en la relación proceso-recurso $p_{33}r_{22}$, aplicando la pauta uno - (método tradicional).....	181
Tabla 153. Valores históricos de la relación proceso-recurso $p_{34}r_{33}$, los cuatro únicos registros - (método tradicional).....	182
Tabla 154. Valor imputado con la media ponderada en la relación proceso-recurso $p_{34}r_{33}$, aplicando la pauta uno - (método tradicional).....	182
Tabla 155. Valor histórico de la relación proceso-recurso $p_{35}r_{12}$, registro único - (método tradicional)...	183
Tabla 156. Valor imputado con la media ponderada en la relación proceso-recurso $p_{35}r_{12}$, aplicando la pauta uno - (método tradicional).....	183
Tabla 157. Valor asignado según las prestaciones del nodo para los criterios de $p_{35}r_{24}$, aplicando la pauta dos - (método tradicional).	183
Tabla 158. Valoraciones asignadas a los criterios para calcular la prioridad nodal, con valores imputados y asignado en algunas relaciones de proceso-recurso - (método tradicional).	184
Tabla 159. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia nodal que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo, con valores imputados/asignados - (método tradicional).	186
Tabla 160. Prioridades nodales de los procesos para acceder a cada recurso p_{11} a p_{25} - (método tradicional).	190
Tabla 161. Prioridades nodales de los procesos para acceder a cada recurso p_{31} a p_{37} - (método tradicional).	191
Tabla 162. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, de p_{11} al p_{25} - (método tradicional).	191
Tabla 163. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, de p_{31} al p_{37} - (método tradicional).	191
Tabla 164. Prioridades globales finales de los procesos para acceder a cada recurso, de p_{11} al p_{25} - (método tradicional).....	192
Tabla 165. Prioridades globales finales de los procesos para acceder a cada recurso, de p_{31} al p_{37} - (método tradicional).....	193
Tabla 166. Prioridades globales finales para asignar los recursos, constituye la Función de Asignación para Sistemas Distribuidos (FASD), del método tradicional (MT).....	194
Tabla 167. Prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso, constituye la Función de Asignación para Sistemas Distribuidos Ordenada (FASDO) - MT.....	194
Tabla 168. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (segunda iteración) - MT.....	195
Tabla 169. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (tercera iteración) - MT.	195
Tabla 170. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso	

(cuarta iteración) - MT.	196
Tabla 171. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (quinta iteración) - MT.	196
Tabla 172. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (sexta iteración) - MT.	196
Tabla 173. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (séptima iteración) - MT.	197
Tabla 174. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (octava iteración) - MT.	197
Tabla 175. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (novena iteración) - MT.	197
Tabla 176. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décima iteración) - MT.	197
Tabla 177. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo primera iteración) - MT.	198
Tabla 178. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo segunda iteración) - MT.	198
Tabla 179. Asignaciones obtenidas en (La Red Martínez, 2017), del escenario E1, del escenario E2 y del método tradicional (MT).	198
Tabla 180. Comparativa de asignaciones obtenidas en cada iteración de (La Red Martínez, 2017) con respecto al escenario E1.	200
Tabla 181. Comparativa de asignaciones obtenidas en cada iteración de (La Red Martínez, 2017) con respecto al escenario E2.	201
Tabla 182. Comparativa de asignaciones obtenidas en cada iteración de (La Red Martínez, 2017) con respecto al método tradicional.	201
Tabla 183. Comparativa de las posiciones globales obtenidas en las asignaciones de (La Red Martínez, 2017) con respecto al escenario E1.	202
Tabla 184. Comparativa de las posiciones globales obtenidas en las asignaciones de (La Red Martínez, 2017) con respecto al escenario E2.	202
Tabla 185. Comparativa de las posiciones globales obtenidas en las asignaciones de (La Red Martínez, 2017) con respecto al método tradicional.	202
Tabla 186. Cantidad de nodos imputados/asignados con las pautas establecidas para la carga computacional del escenario E1.	203
Tabla 187. Cantidad de nodos imputados/asignados con las pautas establecidas para la carga computacional del escenario E2.	203
Tabla 188. Cantidad de relaciones recurso-proceso imputados/asignados con las pautas establecidas para calcular la prioridad o preferencia de los procesos del escenario E1.	204
Tabla 189. Cantidad de relaciones recurso-proceso imputados/asignados con las pautas establecidas para calcular la prioridad o preferencia de los procesos del escenario E2.	204
Tabla 190. Cantidad de relaciones recurso-proceso imputados/asignados con las pautas establecidas para calcular la prioridad o preferencia de los procesos para el método tradicional.	205

Lista de figuras

Figura 1. Demostración de cómo funciona el K-Means.....	29
Figura 2. Demostración de cómo funciona el KNN	31
Figura 3. Modelo de decisión.	41
Figura 4. Pautas establecidas para la imputación o asignación de la carga computacional - E1.....	45
Figura 5. Pautas establecidas para la prioridad o preferencia de procesos – E1.....	52
Figura 6. Pautas establecidas para la imputación/asignación de la carga computacional – E2.....	117
Figura 7. Pautas establecidas para la imputación/asignación de la prioridad o preferencia de los procesos – E2.....	119
Figura 8. Estados de los nodos.....	156
Figura 9. Modelo global de decisión para acceso a recursos compartidos, considerando la imputación de datos	164
Figura 10. Modelo de decisión propuesto.....	165
Figura 11. Diagrama de flujo para aplicar la imputación de valores faltantes en la gestión de recursos compartidos en sistemas distribuidos.....	166
Figura 12. Comparación del Modelo propuesto y los modelos tradicionales	168



IMPUTACIÓN DE DATOS EN LA GESTIÓN DE RECURSOS COMPARTIDOS EN SISTEMAS DISTRIBUIDOS

Diego David Duré Attis

Orientador: Prof. Mgter. Federico Agostini

Co-orientador: Prof. Dr. David Luis la Red Martínez

Resumen

Los sistemas distribuidos se componen de múltiples nodos que intercambian información entre sí, para mantener la red conectada, cada nodo es único e independiente y puede tener una capacidad de procesamiento diferente a los demás. En esta propuesta se considera que hay un nodo central que se encarga de recibir y mantener actualizada la información de control de todos ellos, y es el encargado de la asignación de recursos en la modalidad de exclusión mutua, asegurando la disponibilidad de los mismos y respetando las prioridades de los procesos.

En un determinado ciclo de recolección de información de gestión, necesaria para ejecutar y asegurar lo mencionado anteriormente, el nodo central puede recibir de alguno de los nodos información de control incompleta o inexistente, por ejemplo, de los criterios de evaluación de carga nodal, de las prioridades o preferencias de los procesos, etc. Estos datos faltantes, constituyen un obstáculo importante en la gestión de recursos. Las técnicas de imputación de datos permiten estimarlos utilizando diferentes algoritmos, mediante los cuales, se puede imputar una característica importante para una instancia en particular.

Esta investigación propone el uso de imputación de valores faltantes sobre la información de control en el contexto de los Sistemas Distribuidos. Se incorpora una capa de imputación/asignación a un modelo de decisión, que permite completar los valores faltantes con valores estimados, necesarios para establecer un correcto orden de asignación de recursos a los procesos.

Palabras clave: *Sistemas Operativos, Operadores de Agregación, Modelo de Decisión, Imputación de Datos, Amputación de Datos, K-Means, K-NN, Medias Ponderadas.*



IMPUTACIÓN DE DATOS EN LA GESTIÓN DE RECURSOS COMPARTIDOS EN SISTEMAS DISTRIBUIDOS

Diego David Duré Attis

Orientador: Prof. Mgter. Federico Agostini

Co-orientador: Prof. Dr. David Luis La Red Martínez

Abstract

Distributed systems are composed of multiple nodes that exchange information with each other, to keep the network connected, each node is unique and independent and has a different processing capacity than the others. In this proposal it is considered that there is a central node that is responsible for receiving and keeping the control information of all of them updated and is responsible for the allocation of resources in the modality of mutual exclusion, ensuring their availability and respecting the priorities of the processes.

In a certain cycle of management information collection, necessary to execute and ensure the above, the central node may receive incomplete or non-existent control information from any of the nodes, for example, from the criteria for evaluating nodal load, from the priorities or preferences of the processes, etc. These missing data constitute a major obstacle in resource management. Data imputation techniques allow to estimate them using different algorithms, by means of which, an important characteristic can be imputed for a particular instance.

This research proposes the use of imputation of missing values on control information in the context of Distributed Systems. An imputation/assignment layer is incorporated into a decision model, which allows to complete the missing values with estimated values, necessary to establish a correct order of assignment of resources to the processes.

Keywords: *Operating Systems, Aggregation Operators, Decision Model, Data Imputation, Data Amputation, K-Means, K-NN, Weighted Means.*



IMPUTACIÓN DE DATOS EN LA GESTIÓN DE RECURSOS COMPARTIDOS EN SISTEMAS DISTRIBUIDOS

Diego David Duré Attis

Orientador: Prof. Mgter. Federico Agostini

Co-orientador: Prof. Dr. David Luis La Red Martínez

Ñemombyky

Tembiapo ojapóva computadora omosarambi marandu red rupive, upe computadora ipeteĩ ha ikatu ojapo tembiapo ñambuéva. Ko mba`e jehechaukápe ojeikuaa oĩha peteĩ nodo imbaretevéva oguerúva ha oguerékóva ipype marandu ojejesarekóva hese, upe guive oñembosarambi mba`e ha omopeteĩ, ikatuhaguãicha ojeguereko oñeikotevẽ jave, ojehechakuaava`erã mba`e ra`ẽpa oñeikoteveve.

Pe oñembyaty jave marandu nodo imbaretevéva ikatu oguahẽ chupe ambue nodo-gui marandu noimbáiva térã ndaipóriva. Pe marandu noimbáiva ha'e apañuãi ojejapohaguã pe jegueru ha jegueraha marandu opa tenda oguerékóva computadora-pe.

Ko tembiapo ñehesa`ỹijo ohechauka mba`éicha ojeipuru pe imputación ndoguerékóiva marandu pe sistema oñemboja`óvape. Oñemoinge oe imputación/asignación ikatuhaguãicha oñemoimba umi valor ndojeguerékóiva pe valor oñeimo`ãvare, ikatuhaguãicha osẽ hekopete pe tembiapo.

Ñe`ẽ imbaretevéva: Sistemas Operativos, Operadores de Agregación, Modelo de Decisión, Imputación de Datos, Amputación de Datos, K-Means, K-NN, Medias Ponderadas.

Capítulo I - Introducción

1.1. Introducción

Los sistemas distribuidos están compuestos por un grupo de ordenadores cuyos componentes (hardware y software) están conectados en red, y les permite comunicarse y coordinar acciones para lograr un objetivo. En los sistemas informáticos distribuidos donde existen múltiples procesos que cooperan para el logro de una determinada función, es indispensable disponer de modelos de decisión que permitan a los grupos de procesos acceder a los recursos compartidos que sólo se pueden acceder en la modalidad de exclusión mutua. Estos procesos se encuentran distribuidos en múltiples nodos, en los cuales se define una interfaz entre las aplicaciones y el sistema operativo, que a través de un runtime (software en tiempo de ejecución complementario al sistema operativo presente en cada uno de los nodos) incluido en esa interfaz, gestiona los procesos y recursos compartidos. Los runtime interactúan entre sí para intercambiar información y uno de ellos actúa como coordinador global para las asignaciones de recursos a procesos.

En cierto momento alguno de los nodos del sistema distribuido puede ser considerado inaccesible, no pudiendo brindar información al runtime central (coordinador), que a su vez, podría postergar de manera indefinida solicitudes de recursos por parte de los procesos a siguientes rondas de asignación. Esto ocurre específicamente por la falta de información de control sobre los requerimientos de los procesos que pertenecen a ese nodo cuya información de control no está disponible.

Algunos modelos de investigaciones sobre comunicación en sistemas distribuidos se presentan en (Tanenbaum et al., 1996; Tanenbaum, 2009), en (Tanenbaum & Van Steen, 2008) y en (La Red Martínez, 2004) se describen los principales algoritmos de comunicación en sistemas

distribuidos (algoritmos clásicos de las ciencias de la computación).

Investigaciones relacionadas a sincronización en sistemas distribuidos se presenta en (Agrawal & El Abbadi, 1991) donde se propone una solución eficiente y tolerante a fallas para el problema de la exclusión mutua distribuida, desde la óptica de las ciencias de la computación, en (Ricart & Agrawala, 1981), (Cao & Singhal, 2001) y en (Lodha & Kshemkalyani, 2000) donde se presentan unos algoritmos para gestionar la exclusión mutua en redes de computadoras, conforme a las ciencias de la computación; en (Stallings, 2005) se detallan los principales algoritmos de las ciencias de la computación para la gestión distribuida de procesos, los estados globales distribuidos y la exclusión mutua distribuida.

Los métodos considerados innovadores se expresan en (La Red Martínez, 2017), en (La Red Martínez et al., 2017; La Red Martínez, et al., 2018) donde se desarrollan operadores de agregación para asignaciones de recursos en sistemas distribuidos, en (Agostini, 2019), (Agostini & La Red Martínez, 2019) y (Agostini et al., 2019) se presentan modelos de decisión innovadores para la gestión de recursos en sistemas distribuidos.

La presente investigación pretende generar modelos de decisión considerando la imputación de datos para los escenarios posibles cuando la información de variables que indican el estado de carga llega incompleta y/o sea inexistente; ambos casos en el momento que la información sea solicitada por el nodo central y se presenten dichos escenarios.

Investigaciones relacionadas a imputación de datos se encuentran en (Jönsson & Wohlin, 2004) donde se presenta una evaluación del método k-NN utilizando los datos de Likert en un contexto de ingeniería de software, (Mehala & Vivekanandan, 2008) analizan el uso del algoritmo K-Means evaluando la eficiencia del algoritmo de imputación K-Means como un método de imputación para tratar los datos perdidos, en (Primorac et al., 2020) muestran una metodología de

evaluación del desempeño de métodos de imputación mediante una métrica tradicional complementada con un nuevo indicador, basado en el promedio normalizado de la raíz cuadrada del error cuadrático medio (RMSE: Root Mean Squared Error). Otros principales métodos de imputación de datos se describen en (Song et al., 2005), (Vallejos et al., 2011), (Husson et al., 2019), (Choudhury & Pal, 2019), entre otros.

Por tal motivo se pretende generar un método de imputación/asignación para estimar las variables de estados de los nodos considerados inaccesibles, que a través de un modelo de decisión, permita sustituir los datos faltantes por los valores de sustitución o valores imputados. Los valores faltantes en la información de estado de carga enviada por los nodos del sistema distribuido deben ser reemplazados por valores estimados aplicando un método de imputación, para que el modelo de decisión pueda establecer un correcto orden de asignación de recursos. Las técnicas de imputación de datos permiten estimar valores faltantes en el conjunto de datos utilizando diferentes algoritmos, mediante los cuales se puede imputar una característica importante para una instancia en particular.

La propuesta es presentar un método innovador para la gestión de recursos compartidos en sistemas distribuidos, basado en (La Red Martínez, 2017) y en (Agostini et al., 2019), donde se desarrollan operadores de agregación para asignar recursos en sistemas distribuidos, pero que no consideran la posibilidad de información de control faltante. Este trabajo consiste en resolver esta problemática, agregando una capa de imputación/asignación de información de control, a un modelo de decisión existente para gestión de recursos y procesos en sistemas distribuidos, considerando como punto de inicio las premisas y las estructuras de datos mencionadas en esas publicaciones.

1.2. El problema

Planteamiento del problema

En los sistemas de procesamiento distribuido es frecuentemente necesario coordinar la asignación de recursos compartidos que deben ser asignados a los procesos en la modalidad de exclusión mutua; en tales casos se debe decidir el orden en que los recursos compartidos serán asignados a los procesos que requieren de los mismos.

El escenario general propuesto en este trabajo consiste en que en cada nodo se define una interfaz entre las aplicaciones y el sistema operativo, que a través de un runtime (software en tiempo de ejecución complementario al sistema operativo presente en cada uno de los nodos) incluido en esa interfaz, gestiona los procesos y recursos compartidos y define el escenario específico correspondiente a cada situación. Los runtime interactúan entre sí para intercambiar información, además, existe un runtime central, encargado de recopilar la información de todos los nodos, aplicar el proceso de agregación y obtener la lista de asignaciones de recursos a procesos. Este escenario se basa en la propuesta realizada (La Red Martínez, 2017), (La Red Martínez et al., 2017; La Red Martínez, et al., 2018), (Agostini & La Red Martínez, 2019) y (Agostini et al., 2019).

El problema se presenta cuando en cierto momento uno de los nodos del sistema distribuido es considerado inaccesible, no pudiendo brindar información al runtime central que a su vez iría postergando de manera indefinida solicitudes de procesos y recursos a siguientes rondas. Esto ocurre específicamente por la falta de información sobre los requerimientos de los procesos que pertenecen a ese nodo. Se tendrá en cuenta:

- **Escenario 1:** consiste en un nodo central que aloja al runtime encargado de recopilar la información de control respecto del estado de los nodos y de los requerimientos de acceso de

procesos a recursos, proporcionada por todos los nodos que están compartiendo información. En diversas situaciones, esta información puede llegar al nodo central de manera incompleta; es decir, alguno de los criterios de evaluación de carga (CPU, Memoria Virtual, Prioridad de proceso, etc.) puede no estar disponible.

- **Escenario 2:** consiste en un nodo central que aloja al runtime encargado de recopilar la información de control respecto del estado de los nodos y de los requerimientos de acceso de procesos a recursos, proporcionada por todos los nodos que están compartiendo información. Se considera la situación particular en la que no se recibe de algún nodo la totalidad de información requerida de los criterios de evaluación de carga (CPU, Memoria Virtual, Prioridad de proceso, etc.), ya sea, por caída del enlace o por sobrecarga del nodo que debe suministrar información de control al runtime central.

Enunciado preciso del eje de la situación problemática

Como se ha mencionado anteriormente en la descripción de la situación problemática, el runtime gestiona los procesos y recursos compartidos y define el escenario específico correspondiente a cada situación. Los runtime interactúan entre sí para intercambiar información, y mediante un nodo central se recopila la información de todos los nodos, se aplica el proceso de agregación y se obtiene la lista de asignaciones de recursos a procesos.

Durante el funcionamiento de los sistemas de procesamiento distribuido, en cada ciclo se obtiene una imagen global del sistema (macro imagen), para poder verificar los estados de los nodos del SD (Sistema Distribuido). Una vez verificado e identificado la existencia de nodos accesibles y no accesibles, se atienden las solicitudes de los procesos y recursos correspondientes a los nodos accesibles. Aquellas asignaciones que no han podido ser resueltas, se agrupan con las nuevas solicitudes y se vuelve a tener una macro imagen (entendiéndose por macro imagen a la

información referida al estado de los nodos respecto de la carga computacional actual y la demanda de asignación de recursos a procesos) para verificar si el nodo que estaba inaccesible ahora ya no lo está.

El nodo central puede no recibir o recibir información incompleta de algún nodo. Por tal motivo se pretende mediante un método de imputación/asignación estimar las variables de estados de los nodos considerados inaccesibles por los motivos mencionados, y crear un modelo de decisión agregándole la imputación de datos, que permita sustituir los datos faltantes por los valores de sustitución o valores imputados.

En caso de que la información suministrada por los nodos esté completa, el modelo de decisión se ejecuta si necesidad de la imputación de datos, según la propuesta de (Agostini, 2019).

Fuentes de interés

Las principales fuentes de interés son las siguientes:

- a) Los métodos considerados tradicionales para la asignación de recursos en sistemas distribuidos respetando la exclusión mutua en el acceso a los recursos compartidos, según se expresa en (Ricart & Agrawala, 1981), (Cao & Singhal, 2001), (Lodha & Kshemkalyani, 2000), y otros.
- b) Los métodos considerados innovadores expresados en (La Red Martínez, 2017),
- c) (La Red Martínez et al., 2017; La Red Martínez, et al., 2018), (Agostini & La Red Martínez, 2019), (Agostini et al., 2019) y (Agostini, 2019).
- d) Los principales métodos de imputación de datos descriptos en (Jönsson & Wohlin, 2004), (Song et al., 2005), (Mehala & Vivekanandan, 2008), (Vallejos et al., 2011), (Husson et al., 2019), (Choudhury & Pal, 2019), (Primorac et al., 2020) y otros.
- e) Los modelos de decisión utilizando operadores de agregación adecuados mostrados en

(Chiclana et al., 2001), (Peláez et al., 2004, 2007), (Dong et al., 2016) y otros.

- f) Las técnicas de amputación que generan conjuntos de datos con valores faltantes a partir de conjuntos de datos completos expresados en (Twala, 2009), (Schouten et al., 2018), (Schouten & Vink, 2018) y (Primorac et al., 2020).

Formulación del problema

¿Cuáles son los nuevos modelos de decisión que habrá que desarrollar incorporando imputación de datos para la toma de decisiones en grupos de procesos, que trasciendan el enfoque tradicional de las ciencias de la computación?

Sistematización del problema

¿Cuáles son los nuevos modelos de decisión que habrá que desarrollar incorporando imputación de datos para la toma de decisiones en grupos de procesos, que trasciendan el enfoque tradicional de las ciencias de la computación, en las siguientes circunstancias?

- a) Que los procesos accedan a recursos compartidos en la modalidad de exclusión mutua, incorporando mecanismos de imputación de datos para los casos en que la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos sea *incompleta*.
- b) Que los procesos accedan a recursos compartidos en la modalidad de exclusión mutua, incorporando mecanismos de imputación de datos para los casos en que la información sobre las variables que indican el estado de carga de alguno de los nodos sea *inexistente*.

En caso de que la información suministrada por los nodos esté completa, el modelo de decisión se ejecuta sin necesidad de la imputación de datos, según la propuesta realizada en (Agostini, 2019).

1.3. Objetivos

Objetivo General

Generar modelos de decisión considerando imputación de datos para la gestión de recursos y procesos en sistemas distribuidos.

Objetivos específicos

- Generar modelos de decisión y sus correspondientes operadores de agregación para la gestión de grupos de procesos que pueden compartir recursos, estudiando la aplicación de imputación de datos para estimar las variables que indican el estado de carga de los nodos considerados inaccesibles, para los siguientes tipos de situaciones:

- ** Implementar imputación de datos cuando la información de las variables que indican el estado de carga de alguno o algunos de los nodos sea *incompleta*, para que el modelo de decisión pueda establecer un orden de asignación de recursos.

- ** Aplicar imputación de datos cuando la información sobre las variables que indican el estado de carga de alguno de los nodos sea *inexistente*, para que el modelo de decisión pueda establecer un orden de asignación de recursos.

- Ejecutar el modelo de decisión sin necesidad de la imputación de datos cuando la información suministrada por los nodos esté completa.

- Comparar los resultados al aplicar el modelo de decisión cuando los datos están completos respecto del modelo de decisión que resuelva la falta de datos con imputación.

1.4. Justificación

En los sistemas distribuidos es necesario coordinar la asignación de recursos compartidos que deben ser asignados a los procesos en la modalidad de exclusión mutua; en tales casos se debe decidir el orden en que los recursos compartidos serán asignados a los procesos que requieren de

los mismos.

Los valores faltantes en la información de estado de carga enviada por los nodos del SD serán completados por valores estimados aplicando un método de imputación o asignación, para que el modelo de decisión pueda establecer un correcto orden de asignación de recursos. Las técnicas de imputación de datos permiten estimar valores faltantes en el conjunto de datos utilizando diferentes algoritmos, mediante los cuales se puede imputar una característica importante para una instancia en particular. Para simular los posibles datos faltantes se aplicarán algunos de los mecanismos de amputación más conocidos.

Con el trabajo de investigación se podrá agregar una capa de imputación/asignación a un modelo de decisión que no consideraba información de control faltante, para analizar el escenario que se presente en el momento que el nodo central reciba la información de los nodos del sistema distribuido, para luego definir cuál de los métodos de imputación se utilizará y ejecutará cuando la información de las variables que indican el estado de carga sea incompleta y/o inexistente, y finalmente el modelo de decisión permita establecer un correcto orden de asignación de recursos.

1.5. Hipótesis y objeto

Hipótesis

Si se generan modelos de decisión y sus correspondientes operadores de agregación considerando la imputación de datos se podrá estimar la información de variables que indican el estado de carga de los nodos cuando sea incompleta o inexistente.

Objeto

El objeto de la investigación es la asignación de recursos en sistemas distribuidos, especialmente recursos compartidos en la modalidad de exclusión mutua distribuida.

1.6. Antecedentes

Para poder realizar la investigación se han buscado y analizado materiales de gran importancia referentes a:

- a) Comunicación en sistemas distribuidos.
- b) Sincronización en sistemas distribuidos.
- c) Uso compartido de recursos y exclusión mutua en sistemas distribuidos.
- d) Toma de decisiones en grupo y modelos de decisión.
- e) Amputación de datos para generar valores faltantes en múltiples variables para cualquier conjunto de datos completos.
- f) Imputación y principales métodos.
- g) Imputación de datos mediante métodos de imputación en el contexto de modelos de decisión.
- h) Uso de operadores de agregación para representar y evaluar la carga de trabajo de los nodos de procesamiento en un sistema distribuido.

Aunque el número de antecedentes proporcionado en la bibliografía es considerable, no se han encontrado antecedentes que aborden simultáneamente la toma de decisiones relacionada con la gestión de recursos compartidos en sistemas distribuidos bajo diferentes condiciones, utilizando imputación de datos. Por esta razón se considera que el trabajo de investigación propuesto reúne las condiciones de originalidad necesarias, siendo además relevante académicamente ya que:

- a) Se espera producir conocimiento científico específico no disponible actualmente.
- b) El desarrollo de modelos de decisión con imputación de datos para toma de decisiones en grupos de procesos en el contexto de sistemas complejos y esquemas de autorregulación, constituye un aporte significativo a las ciencias de la computación, en especial en

relación con la gestión del acceso a recursos compartidos desde procesos distribuidos, mejorando la gestión de los mismos por parte de los sistemas operativos.

c) Se podría utilizar el modelo de decisión y sus operadores de agregación para la asignación de recursos a procesos, contribuyendo así a mejorar el desempeño de los sistemas distribuidos considerando información global del sistema, aplicando imputación de datos para la información faltante, características que no se contemplan en los métodos tradicionales.

1.7. Encuadre o Marco Teórico

En este apartado se describen los principales sustentos teóricos tenidos en cuenta para la realización del trabajo.

Comunicación en sistemas distribuidos.

En (Birman & Joseph, 1987) se presenta un protocolo de comunicación confiable en presencia de fallas, desde la óptica de las ciencias de la computación; en (Birman et al., 1991) se presenta un protocolo de comunicación multicast para grupos de procesos en la modalidad atómica, desde la óptica de las ciencias de la computación; en (Birman, 2005) se estudian tecnologías, servicios web y aplicaciones de sistemas distribuidos confiables; en (Cicirelli & Nigro, 2016) se presenta un novedoso algoritmo de exclusión mutua distribuida; en (Kamei & Kakugawa, 2018) se presenta un algoritmo distribuido basado en el intercambio de mensajes asincrónico para el problema global del acceso a regiones críticas; en (Joseph & Briman, 1989) se presentan protocolos confiables de comunicaciones en la modalidad de broadcast, desde la óptica de las ciencias de la computación; en (Kaashoek, 1992) se estudia en profundidad la comunicación entre grupos de procesos analizándose protocolos como el FLIP: Fast Local Internet Protocol y el BP: Broadcast Protocol, analizando la comunicación entre grupos confiable y eficiente, la programación paralela, la programación tolerante a fallos utilizando broadcasting y planteando una

arquitectura de tres niveles, el inferior para el protocolo FLIP, el medio para la comunicación en grupos de procesos y el superior para las aplicaciones. Información adicional puede encontrarse en (Van Renesse et al., 2003); en (La Red Martínez, 2004) se describen los principales algoritmos de comunicación en sistemas distribuidos (algoritmos clásicos de las ciencias de la computación); en (Macedonia et al., 1995) se describe una arquitectura de red para entornos virtuales de gran escala para soportar la comunicación entre grupos de procesos distribuidos; en (Silberschatz et al., 2006) se presentan los principales algoritmos de coordinación distribuida y gestión de la exclusión mutua (algoritmos clásicos de las ciencias de la computación); en (Tanenbaum et al., 1996; Tanenbaum, 2009) y en (Tanenbaum & Van Steen, 2008) se describen los principales algoritmos de comunicación en sistemas distribuidos (algoritmos clásicos de las ciencias de la computación).

Sincronización en sistemas distribuidos.

En (Agrawal & El Abbadi, 1991) se presenta una solución eficiente y tolerante a fallas para el problema de la exclusión mutua distribuida, desde la óptica de las ciencias de la computación; en (Alagarsamy, 2003) se analiza los principales algoritmos de exclusión mutua; en (Bagrodia, 1989) se analiza el diseño y la evaluación de rendimiento de algoritmos distribuidos para la sincronización de procesos, se presenta una solución sencilla para el problema de coordinación de comité, que abarca los problemas de sincronización y exclusión asociados con la implementación de encuentro de múltiples vías; en (Joshi & Holzmann, 2007) analizan la verificabilidad en un sistema de archivos; en (La Red Martínez, 2004) se describen los principales algoritmos de sincronización en sistemas distribuidos (algoritmos clásicos de las ciencias de la computación); en (Lamport, 1978) se estudia el tiempo, los relojes y el ordenamiento de eventos en sistemas distribuidos; se analiza el ordenamiento parcial, los relojes lógicos, el ordenamiento total de eventos, el comportamiento anormal, los relojes físicos, etc.; en (Ricart & Agrawala, 1981), (Cao

& Singhal, 2001) y en (Lodha & Kshemkalyani, 2000) se presentan unos algoritmos para gestionar la exclusión mutua en redes de computadoras, conforme a las ciencias de la computación; en (Stallings, 2005) se detallan los principales algoritmos de las ciencias de la computación para la gestión distribuida de procesos, los estados globales distribuidos y la exclusión mutua distribuida; en (Tanenbaum et al., 1996; Tanenbaum, 2009) se describen los principales algoritmos de sincronización en sistemas distribuidos (algoritmos clásicos de las ciencias de la computación); en (La Red Martínez, 2017) se desarrollan operadores de agregación para asignaciones de recursos en sistemas distribuidos.

Sistemas de soporte a la decisión.

En (Chen, 2001) se estudia la aplicación de métodos lingüísticos de toma de decisiones para tratar el problema de evaluación de calidad de servicio; en (La Red et al., 2011) presenta un modelo de decisión en grupo con la utilización de operadores de agregación de la familia OWA (Ordered Weighted Averaging: Promedio de Pesos Ordenados); en (Fodor & Roubens, 1994) se estudia el modelado de preferencias difuso para el soporte de decisiones multicriterio; en (Fullér, 1996) se presenta la utilización de operadores de agregación de la familia OWA (Ordered Weighted Averaging: Promedio de Pesos Ordenados) para la toma de decisiones; en (García-Melón et al., 2006) se presenta el ANP (Analytic Network Process: Proceso Analítico en Red) y su aplicación a un caso concreto de toma de decisiones. El ANP se basa en el MCDA (Multiple Criteria Decision Analysis: Análisis de Decisión Multi-Criterio); en (Greco et al., 2002) se presentan metodologías para resolver problemas en presencia de múltiples atributos y criterios. En (Marakas, 2002) se estudian los sistemas de soporte de decisión; en (Peláez & Doña, 2003; Peláez et al., 2004, 2007, 2009) se analizan operadores de agregación (de mayoría) en grupo que buscan la representación de la mayoría; en (Saaty, 1980) se presenta el AHP (Analytic Hierarchy Process:

Proceso Analítico Jerárquico) para la toma de decisiones; en (La Red Martínez & Acosta, 2015b) se presentan las principales propiedades matemáticas y las medidas de comportamiento relacionadas con los operadores de agregación; en (La Red Martínez & Pinto, 2015) se presenta una revisión acerca de los operadores de agregación, especialmente los de la familia OWA; en (Agostini et al., 2018) y (Agostini & La Red Martínez, 2019) se presentan operadores de agregación y modelos de decisión para la gestión de recursos en sistemas distribuidos según distintos escenarios.

Toma de decisiones en grupo.

En (Chao et al., 2016) se estudia la forma de obtener un vector de prioridades colectivo a partir de diferentes formatos de expresión de las preferencias por parte de los decisores. El modelo puede reducir la complejidad de la toma de decisiones y evitar la pérdida de información cuando se transforman los diferentes formatos en un formato único de expresión de las preferencias; en (Chiclana et al., 2000), (Chiclana et al., 2001) y (Chiclana et al., 2004) se presentan varios operadores de agregación utilizables para la toma de decisiones en grupos; en (Dong et al., 2016) se define un problema complejo y dinámico de toma de decisiones en grupo con múltiples atributos y se propone un método de resolución que utiliza un proceso de consenso para grupos de atributos, de alternativas y de preferencias, presentándose un modelo de decisión para problemas del mundo real; en (La Red Martinez & Acosta, 2014) se presenta una nueva perspectiva para los modelos de decisión para la sincronización de procesos en sistemas distribuidos; en (La Red Martínez & Acosta, 2015a) se realiza una revisión del modelado de preferencias para los modelos de decisión; en (La Red et al., 2011) y (La Red Martínez, 2004) se presenta el operador WKC-OWA para agregar información en problemas de decisión democrática; en (Lu & Ruan, 2007), (Martínez et al., 2007) y (Martinez et al., 2006) se estudia la toma de decisiones en grupo con multiobjetivos,

el tratamiento de información heterogénea en procesos ingenieriles de evaluación y en sistemas de ingeniería; en (Peláez et al., 2003, 2004, 2007) se analizan operadores de agregación de mayoría y sus posibles aplicaciones a la toma de decisiones en grupo; en (Yager, 1988, 1993) se presentan y analizan los operadores OWA (Ordered Weighted Averaging: Promedio de Pesos Ordenados) aplicados a la toma de decisiones multicriterio; en (Yager & Kacprzyk, 1997; Yager & Pasi, 2002) se presentan los operadores OWA y sus aplicaciones en la toma de decisiones multiagente.

Imputación y Principales Métodos.

En (Wilks, 1932) propone el reemplazo de los datos faltantes por la media de los datos presentes de la variable, aplicado cuando existen pocos datos faltantes, ya que tienden a distorsionar la distribución de las variables.

En (Scheffer, 2002) se expone que la falta de datos es común e incierta en muchos campos de investigación.

(De Leeuw, 2001) describe el problema de la falta de datos en las encuestas y da sugerencias sobre cómo lidiar con ello.

(Cartwright et al., 2003) comparan la imputación de la media de la muestra y el k-NN.

En (Gu & Gao, 2005; Nelwamondo & Marwala, 2007) se exponen que la imputación es una clase de procedimientos que tiene por objeto completar los valores que faltan con los estimados.

En (Huisman, 2014) se propone varios métodos de tratamiento para superar los problemas creados por los datos faltantes, investigando el uso de algunos procedimientos de imputación sencillos para manejar los datos de la red.

(Chang & Ge, 2011) comparan diez métodos de imputación para tratar el problema del valor perdido de los datos de las micro matrices.

En (Nilsson, 2016) se explora un enfoque de imputación en datos incompletos sobre el desempleo del mercado laboral en varios países.

(Berchtold & Surís, 2017) utilizan datos de un estudio real para comparar diferentes estrategias que combinan la imputación múltiple y el método de las ecuaciones encadenadas.

(Casleton et al., 2018) proponen dos métodos para la imputación de datos de múltiples fuentes, que aprovechan la posible correlación entre los datos de diferentes sensores. Estos métodos se aplican a los datos recogidos de una variedad de sensores que vigilan una instalación experimental.

(Lu & Mei, 2018) proponen un método de imputación de datos faltantes basado en un auto-organizador de máquina de aprendizaje extremo.

(Kalton & Kasprzyk, 1982) analizan y evalúan los métodos de imputación más utilizados, como las técnicas de ajuste ponderado cuando hay pérdida de un registro completo, y las técnicas de imputación para los casos de registros con pérdidas de sólo algunas variables.

(Laird & D.B., 1987; Little & Rubin, 2002) exponen que la imputación estadística es uno de los métodos más utilizados, que consiste en reemplazar los datos faltantes para una característica dada por la media, moda o mediana de todos los valores conocidos de ese atributo en la clase a la que pertenece la instancia con el atributo faltante.

(Downey & King, 1998) evalúan dos métodos para imputar los datos de Likert, que se utilizan a menudo en las encuestas.

(Raaijmakers, 1999) presenta un método relacionado a la imputación media, para imputar los valores perdidos de Likert en gran escala.

(Myrtveit et al., 2001) comparan cuatro métodos para tratar los datos faltantes: supresión por lista, imputación media, máxima verosimilitud de la información completa e imputación de

patrones de respuesta similares.

(Strike et al., 2001) describen una simulación de eliminación por lista, imputación media e imputación hot deck. En estos tres estudios, el contexto es la estimación del coste del software.

(Musil et al., 2002) definen que los métodos de imputación consisten en sustituir los valores que faltan por los estimados sobre la base de alguna información disponible en el conjunto de datos.

(Gediga & Düntsch, 2003) presentan un método de imputación basado en el análisis de datos de reglas no numéricas. Su método no hace suposiciones sobre la distribución de los datos, y trabaja con coherencia entre los casos más que con la distancia.

(Song et al., 2005) evalúan la diferencia entre MCAR (datos perdidos completamente al azar) y MAR (datos perdidos al azar) utilizando k-NN y la imputación con la media.

(Preda et al., 2005) expresan que hay muchas opciones que varían desde métodos de imputación como la media o la moda, hasta algunos métodos más robustos basados en relaciones entre atributos.

En (Vallejos et al., 2011) se presenta un caso de estudio sobre la comparación de bases reales y bases de datos imputadas aplicando técnicas de minerías de datos.

(Choudhury & Pal, 2019) proponen un método de imputación que utiliza una red neuronal autoorganizada, haciendo un uso innovador de los datos de entrenamiento sin valores faltantes para entrenar al autocodificador para que esté mejor equipado para predecir los valores faltantes.

(Husson et al., 2019) proponen un método de imputación único para los datos de niveles múltiples, que puede utilizarse para completar datos cuantitativos, categóricos o mixtos.

(Primorac et al., 2020) determinan un modelo de decisión para la validación de Métodos de Imputación mediante la utilización de algoritmos de minería de datos.

(Madhu & Rajinikanth, 2012) proponen una nueva medida de indexación al algoritmo de imputación, para los valores de datos faltantes de los atributos, para calcular la medida de similitud entre dos elementos típicos cualquiera del conjunto de datos.

(Farhangfar et al., 2007) desarrollan un marco unificado que sirva de apoyo a una serie de métodos de imputación para valores perdidos en las bases de datos.

(Aljuaide & Sasi, 2016) presentan una comparación de técnicas de imputación para los valores que faltan en los conjuntos de datos.

Imputación k-Means

(Dubes & Jain, 1988; Kaufman & Rousseeuw, 1990) expresan que la imputación k-Means es una extensión de la k-Means básica, que calcula los valores no disponibles; es decir, que es similar al algoritmo de la k-Means, pero maneja los datos faltantes en un conjunto de datos.

(Mehala & Vivekanandan, 2008) analizan el uso del algoritmo k-Means evaluando su eficiencia como un método de imputación para tratar los datos perdidos, comparando su rendimiento con el rendimiento obtenido por la media, la mediana, la moda y el C4.5 (entiéndase por C4.5 como método de aprendizaje basados en relaciones entre atributos).

(Ngoc Chau & Phung, 2018) proponen un algoritmo de auto entrenamiento utilizando k-Means y modelos de árboles aleatorios, llamado k-Means Tri Forest. k-Means ayuda a formar un método de imputación de datos incompletos basado en clúster.

Imputación k-NN

(Cover & Hart, 1967) indican que en la imputación k-NN los valores perdidos de una instancia se imputan considerando un número determinado de instancias que son más similares a la instancia de interés. La similitud de dos instancias se determina utilizando una función de distancia.

(Sande, 1983) expresa que el método k-Nearest Neighbor (k-NN) es un método común del método de imputación hot deck, en el que se seleccionan k donantes de los vecinos (es decir, los casos completos) de tal manera que minimicen alguna medida de similitud.

(Todeschini, 1990) propone un nuevo método basado en el algoritmo k-NN para la estimación de los valores perdidos en un conjunto de datos multivariantes.

(Huisman, 2000) presenta una comparación de imputación incluyendo k-NN con k igual a 1, en el contexto de la investigación del ADN.

(Troyanskaya et al., 2001) informan en el contexto de la investigación del ADN sobre una comparación de tres métodos de imputación: uno basado en la descomposición del valor singular (SVD), una variante de k-NN y un promedio de fila.

(Batista & Monard, 2015) comparan k-NN con los algoritmos de aprendizaje de máquina C4.5 y C2 (métodos internos utilizados para tratar datos faltantes).

Para el método k-NN un parámetro importante es su valor de k. (Clarke et al., 1974) sugieren, en el contexto de la estimación de la densidad de probabilidad dentro de la clasificación de patrones el uso de k igual a $\sqrt{N}$, donde N corresponde al número de vecinos. Por otra parte, (Cartwright et al., 2003) sugieren un bajo k, típicamente 1 o 2, pero señalan que si k es igual a 1 es sensible a los valores atípicos y por consiguiente usan k igual a 2. (Chen & Shao, 2000; Huisman, 2000; Myrtveit et al., 2001; Strike et al., 2001) utilizan un valor de k igual 1. (Batista & Monard, 2015) reportan k igual a 10 para grandes conjuntos de datos.

(Troyanskaya et al., 2001) argumentan que el método es bastante insensible a la elección de k. Como k aumenta, la distancia media a los donantes se hace mayor, lo que implica que los valores de sustitución podrían ser menos precisos. Eventualmente, a medida que k se aproxima a N, el método converge a la imputación media ordinaria donde también contribuyen los casos más

distantes.

(Cartwright et al., 2003) comparan la imputación de la media de la muestra y el k-NN.

(Jönsson & Wohlin, 2004) presentan una evaluación del método k-NN utilizando los datos de Likert en un contexto de ingeniería de software. Simulan el método con diferentes valores de k y para diferentes porcentajes de datos que faltan.

En (García-Laencina et al., 2012) a partir de la técnica clásica basada en los vecinos más cercanos, presentan dos técnicas de vecindad que mejoran los resultados obtenidos en problemas de diagnóstico médica bajo dos enfoques diferentes: la imputación de datos orientada a mejorar la tasa de reconocimiento y la imputación orientada a mejorar la calidad de los datos estimados.

Amputación de Datos

(Twala, 2009) realiza una comparación empírica de las técnicas de manejo de datos faltantes utilizando árboles de decisión.

En (Schouten et al., 2018) se presenta el método de amputación multivariado, que genera amputaciones fiables y permite una regulación adecuada de los problemas de datos faltantes.

(Schouten & Vink, 2018) simulan datos completos y generan valores faltantes según varios tipos de mecanismos MCAR, MAR y MNAR.

(Santos et al., 2019) estudian los diferentes enfoques de generación de datos sintéticos faltantes encontrados en la literatura y analizan sus detalles prácticos, elaborando sus fortalezas y debilidades.

En (Primorac et al., 2020) se utiliza el método de amputación multivariado y posterior imputación, que les permite parametrizar de manera sencilla los procedimientos de amputación y los métodos de imputación facilitando la gestión de los respectivos archivos originales, amputados e imputados.

Medias Ponderadas

En (Abril et al., 2014) presentan una variación del enlace de registros (hacen referencia a un mismo individuo, entre diferentes bases de datos) basada en una media ponderada para calcular distancias entre registros.

En (Gómez et al., 2014) muestran una aplicación de las condiciones de estabilidad estricta al problema de pérdida de datos en un proceso de agregación de información, utilizando las familias de la media ponderada y de la media ponderada cuasi-aritmética.

En (Roy et al., 2017) proponen una combinación de filtro vectorial mediano adaptativo (VMF) y filtro de media ponderada para eliminar el ruido impulsivo de alta densidad de las imágenes en color (un píxel no ruidoso se sustituye por la media ponderada de los píxeles buenos de la ventana de procesamiento).

En (Forc et al., 2018) proponen la adaptación automática del operador de pooling aprendiendo pesos de medias ponderadas ordenadas en Redes Neuronales Convolucionales.

En (Baez et al., 2019) abordan mediante una revisión teórica de trabajos de investigación, la evolución de la toma de decisiones empresariales a través de la media ponderada ordenada.

En (Yoo et al., 2020) se propone un enfoque de conjunto de medias ponderadas sustitutivas para reducir la incertidumbre a escala regional en el cálculo de la evapotranspiración (pérdida de humedad de una superficie por evaporación) potencial.

En (Yugander et al., 2020) mejoran las imágenes ruidosas de resonancia magnética (RM) cerebrales, utilizando el filtrado medio ponderado adaptativo (AWMF) y el filtrado homomórfico.

En (Si et al., 2021) proponen un nuevo método llamado Máquina de Aprendizaje Extremo, basada en la máxima discrepancia de media ponderada para tratar con los problemas de adaptación de dominios no supervisados.

1.8. Relevancia

Académica

Se trata de una investigación que conlleva en sí una fuerte innovación y que llama constantemente a la reflexión sobre las formas de encarar la solución a problemas de las ciencias de la computación desde una óptica no convencional, proponiéndose la utilización de distintas disciplinas en las que ya se han realizado trabajos y publicaciones, y que ahora se tiene previsto integrar en la búsqueda de nuevas soluciones a los problemas planteados de gestión de recursos compartidos en grupos de procesos operando en sistemas distribuidos.

Social

El propósito del presente plan de investigación es el de desarrollar modelos de decisión que permitan la toma de decisiones incorporando imputación de datos para los casos en que exista información faltante, aplicables a grupos de procesos computacionales; en tal sentido podría considerarse que la relevancia social está dada, ya que la informática y las comunicaciones son hoy parte de la vida diaria de millones de personas alrededor del mundo, considerándose que la informática se ha vuelto ubicua y perversiva, ubicua en el sentido de estar presente de manera muchas veces inadvertida en gran cantidad de actos de la vida cotidiana, y perversiva al permitir el acceso a los sistemas informáticos casi desde cualquier lugar y en cualquier momento, haciendo uso de redes de comunicaciones muchas veces inalámbricas.

Es un hecho además que actualmente la mayoría de los sistemas informáticos utilizados operan de manera distribuida sobre redes de comunicaciones, debiendo realizar permanentemente operaciones de coordinación y toma de decisiones en grupos de procesos.

Los sistemas antes muy brevemente descriptos son masivamente utilizados por la población en general, de manera directa o indirecta, razón por la cual se considera que cualquier

forma de mejorar el desempeño de dichos sistemas redundaría en beneficio de la sociedad, en el contexto de la llamada Sociedad de la Información y el Conocimiento, con aplicaciones relacionadas, por ejemplo, con el e-government (gobierno electrónico), el e-commerce (comercio electrónico), el e-learning (aprendizaje electrónico), la e-democracy (democracia electrónica) y la m-cognocracy (gobierno electrónico móvil en la sociedad del conocimiento).

Relevancia Científica

El desarrollo de modelos de decisión para toma de decisiones en grupos de procesos en el contexto de sistemas complejos y esquemas de autorregulación, constituye un aporte significativo de las ciencias cognitivas a las ciencias de la computación, en especial en relación con la gestión del acceso a recursos compartidos desde procesos distribuidos, mejorando la gestión de los mismos por parte de los sistemas operativos.

1.9. Metodología

Dimensión de la Estrategia General

a) Tipo de Estrategia y Fundamentación

Considerando el objeto, problema y objetivos, se decidió utilizar el tipo de investigación cuantitativa, con medición de variables, formulación de hipótesis, y validación de hipótesis.

b) Universo

El universo fue formado por recursos y procesos alojados en los nodos que compone el sistema distribuido, donde se originan requerimientos de los accesos compartidos en la modalidad de exclusión mutua.

c) Unidad de Análisis

La unidad de análisis estuvo integrada por la relación proceso-recurso requerido, donde los procesos y recursos están en diferentes nodos del sistema distribuido.

d) Selección de Casos

Se trabajó con los escenarios indicados en los objetivos específicos y detallados en el problema.

e) Énfasis en el Proceso Lineal

Conforme a las características de la investigación propuesta, se ha optado por la lógica cuantitativa, basada en la búsqueda y verificación de relaciones causa-efecto, que se puedan generalizar; es decir, hay un proceso lineal de relación teoría-empiría, con diferentes espacios y momentos para la definición de las hipótesis, la obtención de los datos, su análisis y su posterior interpretación.

Dimensión de las Técnicas de Obtención y Análisis de Información Empírica

Técnicas de Obtención y Análisis

La información obtenida y analizada, surgió de los procesos de simulación de los modelos de decisión y operadores de agregación que fueron desarrollados. A continuación, se describen las estructuras de datos con las que se trabajó y el tratamiento y análisis de datos efectuados.

Estructuras de datos

El sistema de matrices de datos utilizado contempló las siguientes premisas y estructuras de datos. Operan grupos de procesos distribuidos en nodos de procesos que acceden a recursos críticos compartidos en la modalidad de exclusión mutua distribuida, que, ante la demanda de recursos por parte de los procesos, se debe decidir cuáles serán las prioridades para asignar los recursos a los procesos que los requieren (sólo intervienen como alternativas de asignación a los procesos aquellos recursos disponibles, es decir, no asignados ya a determinados procesos):

- El permiso de acceso a los recursos compartidos propios de un nodo no dependerá sólo de si los nodos los están utilizando o no, sino del valor de agregación de las opiniones (prioridades) de los distintos nodos respecto de otorgar el acceso a los recursos compartidos (alternativas).
- Las opiniones (prioridades) de los distintos nodos respecto de otorgar el acceso a los

recursos compartidos (alternativas) dependerá de la consideración del valor de variables que representen el estado de cada uno de los distintos nodos. Cada nodo deberá expresar sus prioridades para la asignación de los distintos recursos compartidos respecto de los requerimientos de recursos de cada proceso de cada grupo.

La denominación para las variables que fueron utilizadas es la siguiente:

Nodos que alojan procesos: $1, \dots, n$.

Procesos alojados en cada uno de los n nodos: $1, \dots, p$.

Grupos de procesos distribuidos: $1, \dots, g$.

Tamaño de cada uno de los g grupos de procesos: $1, \dots, t$.

Recursos críticos compartidos en la modalidad de exclusión mutua distribuida disponibles en cada uno de los n nodos: $1, \dots, r$.

- Estados posibles de cada uno de los p procesos:

Los posibles estados de cada uno de los p procesos son: a) Grupo al que pertenece el proceso (0 significa proceso independiente); b) Requiere sincronización con los demás procesos del grupo al que pertenece (los procesos del grupo deben estar activos en sus respectivos procesadores en un mismo lapso de tiempo), (0 significa que no requiere sincronización, 1 significa que sí la requiere); c) En espera de un recurso compartido con el grupo de procesos al que pertenece; d) En espera de un recurso no compartido con el grupo de procesos al que pertenece; e) En ejecución con permiso de acceso a un recurso compartido con el grupo de procesos al que pertenece; f) En ejecución sin permiso de acceso a un recurso compartido con el grupo de procesos al que pertenece; y g) Inactivo.

- Estado posible de cada uno de los n nodos:

Los posibles estados de cada uno de los n nodos son: a) Número de procesos; b) Prioridades

de los procesos; c) Uso de CPU; d) Uso de memoria principal; y e) Uso de memoria virtual.

- Estado posible de cada uno de los r recursos:

Los posibles estados de cada uno de los r recursos críticos compartidos en la modalidad de exclusión mutua distribuida existentes en el nodo: a) Asignado a un proceso local; b) Asignado a un proceso remoto; c) Disponible.

- Predisposición (prioridad nodal) para otorgar el acceso a cada uno de los r recursos críticos compartidos (alternativas) en la modalidad de exclusión mutua distribuida (resultará de la consideración de las variables representativas del estado del nodo, para cada recurso crítico compartido existente). Se obtendrá una tupla por cada uno de los n nodos, cada tupla contendrá r valores (prioridades nodales) para compartir los recursos críticos.

Estado global del sistema:

a) Número de grupos de procesos y tamaño (número t de procesos) de cada uno de los g grupos.

b) Porcentajes de consenso requeridos para otorgar el acceso a cada uno de los r recursos críticos disponibles en cada uno de los n nodos.

c) Predisposición (prioridad global) para otorgar el acceso a cada uno de los r recursos críticos compartidos (alternativas) en la modalidad de exclusión mutua distribuida (resultará de la agregación de las prioridades nodales para cada recurso crítico compartido existente (alternativas). Se obtendrá una tupla con r valores normalizados (prioridades globales) para compartir los recursos críticos.

d) Decisión de acceso a los r recursos críticos en función de contrastar las prioridades globales normalizadas para compartir los mismos con los porcentajes de consenso requeridos para otorgar los respectivos accesos.

e) El estado global del sistema deberá actualizarse reiteradamente mientras haya alguno o algunos de los p procesos que requieran acceder a alguno o algunos de los r recursos compartidos.

El sistema se autorregula reiteradamente en función del estado local de los n nodos y del estado global del sistema, produciéndose una actualización de los estados locales de los nodos como consecuencia de la evolución de sus respectivos procesos y de las decisiones de acceso a los recursos críticos producidas teniendo en cuenta el estado global del sistema: el sistema distribuido en el que se ejecutan grupos de procesos que acceden a recursos críticos, se observa a sí mismo y produce decisiones de accesos a recursos que modifican el estado del sistema y lo reajustan reiterativamente.

Evaluación de los Métodos de Imputación

Para definir los métodos de imputación del modelo de decisión para la gestión de recursos y procesos en SD, se evaluaron varios de ellos para luego determinar cuáles son los más apropiados para cada uno de los escenarios planteados, teniendo como referencia investigaciones de métodos de imputación en (Nilsson, 2016), (Berchtold & Surís, 2017), (Casleton et al., 2018), (Husson et al., 2019), (Primorac et al., 2020), entre otros. Entre los métodos analizados, se ha considerado una comparación entre el método K-Means y el K-NN.

Definición del Método de imputación K-Means

Es un algoritmo de agrupación que forma parte del aprendizaje automático no supervisado. Se utiliza para crear cluster a partir de muchos puntos de datos no etiquetados, donde los datos que tienen una propiedad o un patrón similar se agrupan en el mismo cluster. La " K " en K representa el número de cluster que forma, que es que es asignado arbitrariamente.

¿Cómo funciona K-Means?

El método de imputación K-Means funciona de la siguiente manera:

- Paso 1, después de recoger y cargar los datos al modelo, primero se decide el valor de " K ", que no es más que el número de cluster. El valor de " K " se decide generalmente en función de la comprensión del dominio y de los datos.
- Paso 2, a continuación, se selecciona aleatoriamente el número K de centroides (centro de los puntos de datos, inicialmente se elige al azar por lo que en realidad no está en un centro) que representa el centro de un clúster particular. Por ejemplo, si el valor de K es 2 entonces se tiene dos cluster y cada cluster tendrá un centroide.
- Paso 3, luego se calcula la distancia de cada punto de datos a cada centroide creado, se asigna los puntos de datos a un determinado clúster en función de la distancia más cercana que tenga a un centroide. Por ejemplo, si la distancia de un punto de datos " P " al centroide " A " es de 5 y al centroide " B " es de 7, el punto de datos " P " se asignará al centroide " A ".
- Paso 4, después de asignar cada punto de datos a su clúster más cercano, se calcula el centroide real de cada clúster.
- Paso 5, se repiten los pasos 3 y 4 hasta que esté totalmente optimizado.

Se puede visualizar en la Figura 1 un pequeño ejemplo de cómo trabaja el algoritmo de clasificación no supervisada K-Means, suponiendo que se tiene datos sin etiquetas y se quiere separarlos en clústeres, siguiendo las secuencias de los pasos 1 al 10, con sus respectivas descripciones en cada paso.

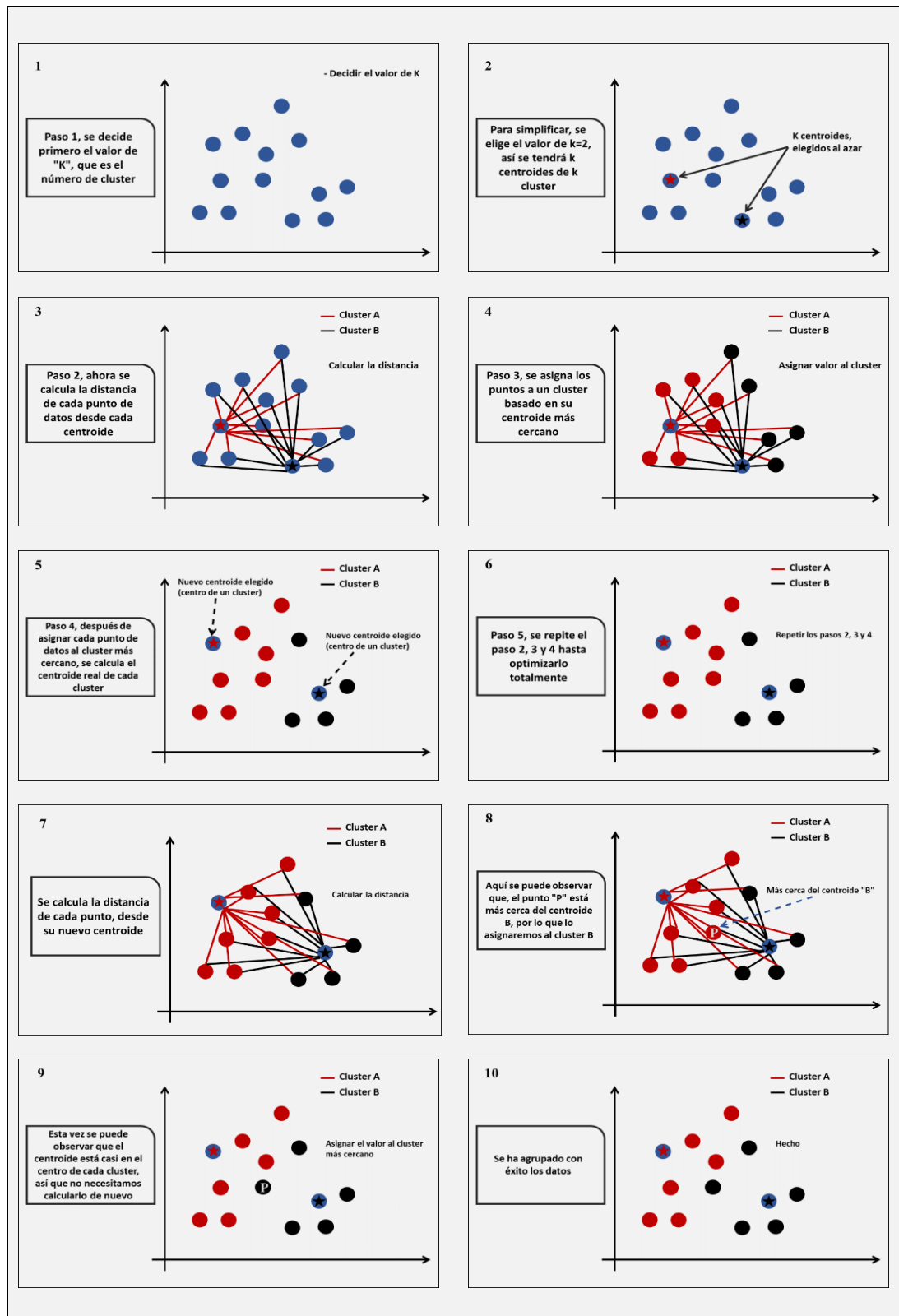


Figura 1. Demostración de cómo funciona el K-Means

Fuente: Elaboración propia.

Definición del Método de imputación KNN

Es un algoritmo de clasificación de aprendizaje supervisado. Se utiliza para clasificar un punto de datos a una determinada clase, basándose en su propiedad y patrones, por ejemplo, si se tienen dos clases, perro y gato y se pasa una imagen de un gato entonces se clasificará como una clase de gato.

"K" en KNN es un parámetro que se refiere al número de vecinos más cercanos a incluir en la mayoría del proceso de selección.

¿Cómo funciona KNN?

Para entender cómo funciona KNN se explica a continuación el proceso paso a paso:

- Paso 1, se elige el punto de datos que se quiere clasificar, llámese "*M*" por un momento.
- Paso 2, al igual que K-Means, se decide el valor de *K*, pero aquí no es el valor del clúster, sino el número de vecinos más cercanos al punto que se quiere clasificar, que en este caso es "*M*". El valor de *K* se decide arbitrariamente dependiendo de los datos.
- Paso 3, decidido el valor de "*K*" (que son los *K* puntos de datos más cercanos de la clase disponible a los puntos de datos que se está clasificando en "*M*") se crea un círculo imaginario que cubra esos valores *K* alrededor del punto "*M*", se cuentan los puntos de datos de la clase que son máximos dentro del círculo. El punto de datos se asignará a la clase que tenga el máximo punto de datos dentro de ese círculo.

Se puede visualizar en la Figura 2 un ejemplo de demostración del algoritmo de clasificación supervisada KNN, siguiendo las secuencias de los pasos 1 al 6 con sus respectivas descripciones en cada paso.

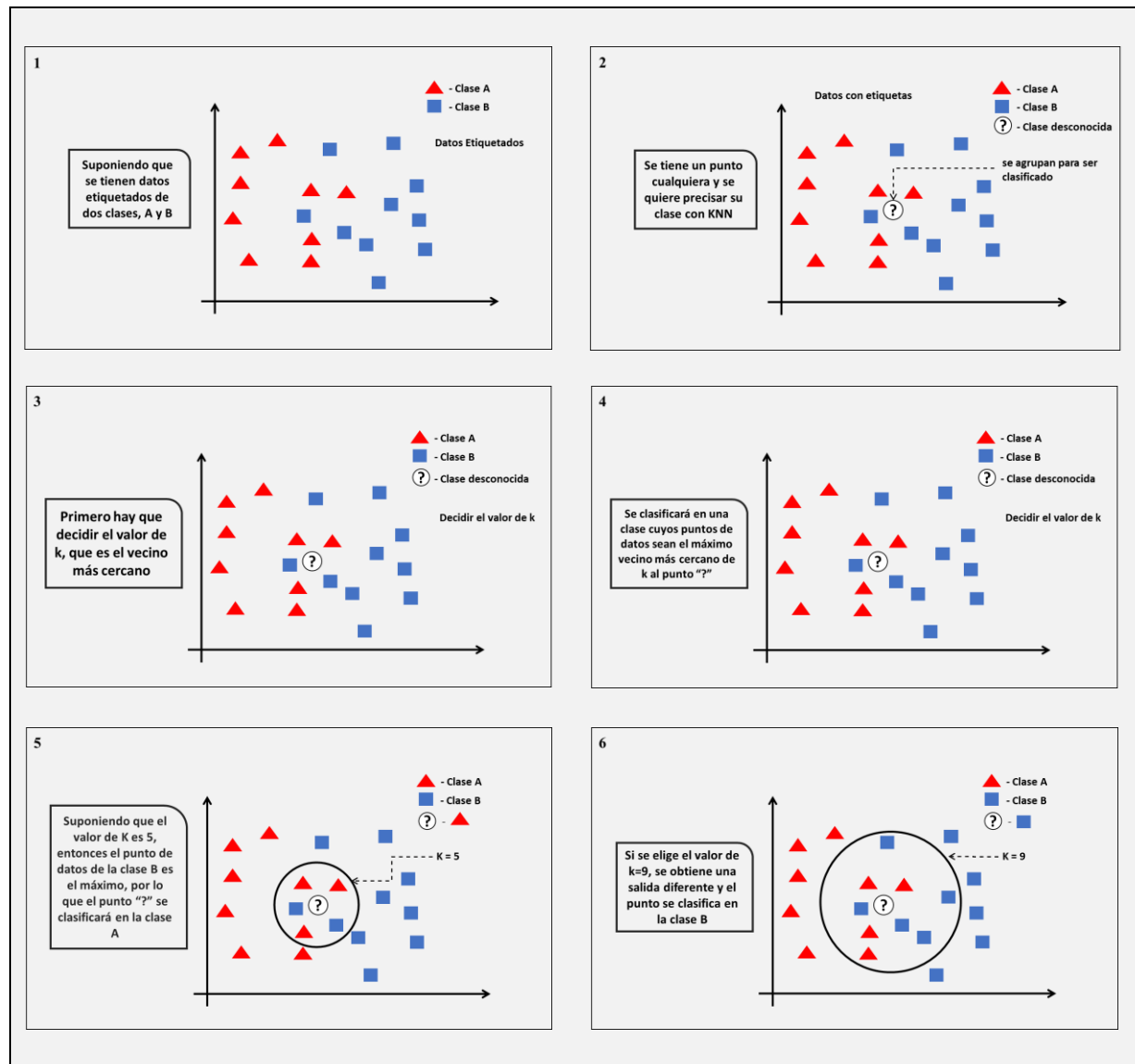


Figura 2. Demostración de cómo funciona el KNN

Fuente: Elaboración propia.

Comparaciones entre los Métodos de Imputación K-Means y KNN

La diferencia principal entre K-Means y KNN es que, K-Means es un algoritmo de agrupación de aprendizaje no supervisado, mientras que KNN es un algoritmo de clasificación de aprendizaje supervisado. K-Means crea clases a partir de datos no etiquetados, mientras que KNN clasifica los datos en clases disponibles a partir de datos etiquetados.

Tratamiento y análisis de datos

Se efectuaron simulaciones con los modelos de decisión propuestos y los modelos de decisión computacionales clásicos, a los efectos de analizar el comportamiento de estos para las mismas condiciones de carga de trabajo y de consumo de recursos, dando lugar a un esquema iterativo de modificación de los modelos propuestos para intentar lograr un desempeño de los mismos al menos equivalente al de los modelos computacionales clásicos.

Para crear los escenarios propuestos se aplicaron técnicas de amputación, que permitieron generar conjuntos de datos con valores faltantes a partir de conjuntos de datos completos, referente a la información de control respecto del estado de carga de los nodos y de los requerimientos de acceso de procesos a recursos, proporcionada por todos los nodos del SD.

1.10. Estructura de la Tesis

Para la elaboración de la presente investigación se procedió a presentar el problema del cual se inició el trabajo, como también los antecedentes principales y los conceptos teóricos. Además, se mencionó la metodología a utilizar y posteriormente se indican los capítulos restantes con los cuales se ha estructurado esta tesis.

Capítulo II – Capa de imputación/asignación para datos faltantes. En este Capítulo se han efectuado los siguientes pasos:

- Estudiado el escenario general E1.
- Formalizado el operador de agregación para el escenario E1.
- Validado empíricamente el operador de agregación para el escenario E1.
- Definido y validado teóricamente el modelo global de decisión del escenario E1.
- Definido detalladamente las etapas, niveles o capas del modelo de decisión E1 y validado empíricamente.

- Evaluado y comparado los resultados del modelo de decisión del escenario E1.

Capítulo III - Capa de imputación/asignación para datos inexistente. En este capítulo se han efectuado los siguientes pasos:

- Estudiado el escenario general E2.
- Formalizado el operador de agregación para el escenario E2.
- Validado empíricamente el operador de agregación para el escenario E2.
- Definido y validado teóricamente el modelo global de decisión del escenario E2.
- Definido detalladamente las etapas, niveles o capas del modelo de decisión E2 y validado empíricamente.
- Evaluado y comparado los resultados del modelo de decisión del escenario E1.

Capítulo IV - Modelo de decisión: Se ha desarrollado y explicado el modelo de decisión propuesto para poder aplicar cada uno de los operadores de agregación para cada escenario.

Capítulo V - Conclusiones y futuras líneas de investigación: se han comentado las principales conclusiones y se han indicado las posibles líneas futuras de investigación.

1.11. Discusiones y Comentarios

Se elaboró un despliegue de los antecedentes de investigaciones y el marco teórico donde se especifican aspectos relevantes y actualizados del estado del arte en las diferentes áreas con relación a la presente tesis (comunicación y sincronización en sistemas distribuidos, sistemas de soporte de decisión, toma de decisiones en grupo, imputación y principales métodos, técnicas de amputación). Se definió la hipótesis y los objetivos de la investigación donde se pueden apreciar las situaciones en las que los procesos puedan acceder a recursos compartidos en la modalidad de exclusión mutua.

También se describieron las estructuras de datos, los elementos que intervienen dentro del

sistema, los nodos, recursos y procesos, con prioridades y cargas de trabajo y que en diferentes situaciones permiten que el sistema se autorregule de manera reiterada en función del estado local de los n nodos y del estado global del sistema.

Para finalizar se describieron los capítulos conformados en la estructura de la tesis. A parte de lo expuesto anteriormente se prosigue con dos capítulos referentes a los distintos escenarios propuestos, un capítulo sobre el modelo de decisión propuesto, otro capítulo sobre comparaciones de los resultados cuando los datos están completos respecto del modelo de decisión que resuelva la falta de datos con imputación, finalizándose luego con las conclusiones y líneas futuras de investigación.

En este trabajo se ha generado un método de imputación/asignación para estimar las variables de estados de los nodos considerados inaccesibles, que a través de un modelo de decisión, permita sustituir los datos faltantes por los valores de sustitución o valores imputados. Los valores faltantes en la información de estado de carga enviada por los nodos del sistema distribuido deben ser reemplazados por valores estimados aplicando un método de imputación, para que el modelo de decisión pueda establecer un correcto orden de asignación de recursos. Las técnicas de imputación de datos permiten estimar valores faltantes en el conjunto de datos utilizando diferentes algoritmos, mediante los cuales se puede imputar una característica importante para una instancia en particular.

La propuesta es presentar un método innovador para la gestión de recursos compartidos en sistemas distribuidos, basado en (La Red Martínez, 2017) y en (Agostini et al., 2019), donde se desarrollan operadores de agregación para asignar recursos en sistemas distribuidos, pero que no consideran la posibilidad de información de control faltante. Este trabajo consiste en resolver esta problemática, agregando una capa de imputación/asignación de información de control, a un

modelo de decisión existente para gestión de recursos y procesos en sistemas distribuidos, considerando como punto de inicio las premisas y las estructuras de datos mencionadas en esas publicaciones.

Lo mencionado con anterioridad se inicia en el siguiente capítulo.

Capítulo II – Capa de imputación/asignación para el modelo de decisión

2.1. Trabajos previos

En (La Red Martínez, 2017) y en (La Red Martínez et al., 2017; La Red Martínez, et al., 2018) se presentan nuevos operadores de agregación para modelos de decisión aplicados a la gestión de recursos compartidos en sistemas distribuidos.

En (Agostini & La Red Martínez, 2019), (Agostini et al., 2018) se presentan operadores de agregación y modelos de decisión para la gestión de recursos en sistemas distribuidos según distintos escenarios.

En (Primorac et al., 2020) muestran una metodología de evaluación del desempeño de métodos de imputación mediante una métrica tradicional complementada con un nuevo indicador, basado en el promedio normalizado de la raíz cuadrada del error cuadrático medio (RMSE: Root Mean Squared Error)

2.2. Propuesta de Solución

La imputación de datos será el método a utilizar en el modelo decisión propuesto para resolver el problema de datos faltantes. En base a revisiones bibliográficas se ha optado por utilizar para el escenario E1 el método K-Means, por ser uno de los métodos más utilizados en investigaciones donde los resultados obtenidos son comparados con otros métodos, se puede mencionar:

En (Mehala & Vivekanandan, 2008) se evalúan la eficiencia del algoritmo de imputación K-Means como un método de imputación para tratar los valores faltantes, comparando su rendimiento con el rendimiento obtenido por la Media, la Mediana, la Moda y los C4.5. Para los conjuntos de datos Bupa, Breast Cancer y Pima la imputación K-Means proporciona un buen

resultado en la mayoría de los casos.

En (Patil et al., 2010) se proponen un enfoque novedoso para la estimación del método de datos faltantes utilizando la técnica K-Means con distancia ponderada. El análisis de los resultados muestra que dicho enfoque resulta con un mejor rendimiento.

(Rahman & Davis, 2013) exploran el uso de una técnica de aprendizaje por máquina como método de imputación de valores erróneos para datos cardiovasculares incompletos, y en la mayoría de los casos en K-Means encontraron mejor rendimiento.

En (Primorac et al., 2020) se imputan un conjunto de datos con MV (Missing Values, Valores Perdidos) mediante los métodos de imputación por Medias, K-NN, k-Means y Hot-Deck. Los resultados muestran que el error para el método de imputación K-Means es el más bajo considerando la totalidad de conjuntos de datos.

Método de Imputación K-Means aplicado

Es un algoritmo de clasificación no supervisada que agrupa a las observaciones de forma tal que todas las que se encuentren en el mismo grupo sean lo más semejantes entre sí, la semejanza se calcula con una métrica de similitud, calculando la media del grupo seleccionado por la semejanza (imputación con la media), eso evita el sesgo en la varianza y en la desviación estándar (estimación en base a los valores parecidos y no al total). En K-Means la cantidad de valores parecidos surge de la clusterización realizada (si se tiene que imputar el valor de una variable de un registro de un grupo de 100 registros que componen el clúster, se tomarán los valores disponibles de esa variable a imputar dentro del grupo de los 100).

Funcionamiento del algoritmo K-Means

Según lo expuesto en (Mehala & Vivekanandan, 2008) K-Means define los centroides de k , uno para cada grupo. Estos centroides deben ser colocados de una manera a mejorar, porque una ubicación diferente causa un resultado diferente. Por lo tanto, la mejor opción es colocarlos lo más lejos posible uno del otro. El siguiente paso es tomar cada punto perteneciente a un grupo de datos y asociarlo al centroide más cercano. Cuando no hay ningún punto pendiente, el primer paso se completa y se hace un agrupamiento. En este punto hay que recalcular los nuevos centroides k como baricentros de los grupos resultantes del paso anterior. Después de estos nuevos centroides k , hay que hacer un nuevo agrupamiento entre los mismos puntos del grupo de datos y el nuevo centroide más cercano. Se ha generado un bucle. Como resultado de este bucle, puede que los centroides k cambien su ubicación paso a paso hasta que no se hagan más cambios. Calcula la media de todos los puntos del grupo. Consideremos que el número k de grupos es el valor x_{ij} de la clase m -ésima, C_{km} falta en el grupo k -ésima entonces será reemplazado por:

$$\hat{x}_{kij} = \sum_{i: x_{kij} \in C_{km}} \frac{x_{kij}}{n_{km}}$$

donde, n_{km} representa el número de valores no faltantes en la j -ésima parte de la clase k -ésima.

La limitación de este método es que sustituirá dentro de cada clúster o agrupamiento todos los valores que faltan por el mismo valor, aunque la variabilidad definida está asociada a los datos reales que faltan. En consecuencia, se modifica la distribución estadística de los datos y ello repercutirá en la calidad de los mismos. Este es un enfoque simple y el más utilizado en la literatura. Una vez agrupados todos los registros, el K-Means calcula la media del grupo seleccionado por la semejanza, para obtener el valor imputado para una variable de un registro.

Aplicación del método K-Means con el lenguaje de programación R

Sin pérdida de generalidad, para la imputación de datos en esta propuesta se utiliza el lenguaje de programación R (R Core Team, 2020). El método de imputación seleccionado es el K-Means, éste imputa el valor faltante utilizando el lenguaje de programación mencionado, primeramente, calcula la cantidad óptima de clúster, y seguidamente utiliza la información de los clústers obtenidos para imputar los valores faltantes.

Amputación de Datos

En esta propuesta, se realiza la amputación con el objetivo de generar datos faltantes en un archivo histórico de información de control, para luego imputar los datos de control cuyo faltante se ha generado, para poder comparar los resultados con datos completos (estos datos son los utilizados en (La Red Martínez, 2017), específicamente en el ejemplo donde se presentan las valoraciones de los criterios establecidos para la carga computacional actual de los nodos y las prioridades o preferencias de los procesos) con los resultados logrados con la imputación propuesta.

La amputación es el proceso por el cual se generan conjuntos de datos con valores faltantes a partir de conjuntos de datos completos (Schouten et al., 2018).

En (Primorac et al., 2020) se generaron conjuntos de datos con valores faltantes a partir de un conjunto de datos completo, para luego imputarlos mediante diferentes métodos de imputación. El entorno de trabajo desarrollado para realizar los experimentos de amputación y posterior imputación resultó apropiado y permite la incorporación a futuro de otros mecanismos de amputación y otros métodos de imputación.

Existen diferentes mecanismos para generar la pérdida de datos, (Rubin, 1976) clasificó

los mecanismos en tres tipos diferentes:

- Falta completamente aleatoria (**MCAR**: missing completely at random): representa una situación para la cual la ausencia es independientemente de las variables. Es una ausencia de datos debida exclusivamente al azar.
- Falta aleatoria (**MAR**: missing at random): se refiere cuando la ausencia de datos está ligada a las variables independientes del estudio, pero no a la variable dependiente.
- Falta no aleatoria (**NMAR**: not missing at random): estos casos se dan cuando la pérdida de datos se debe a la variable dependiente, y posiblemente a alguna variable independiente.

Aplicación del mecanismo de amputación con el lenguaje de programación R

Sin pérdida de generalidad, para la amputación de datos se utiliza el lenguaje de programación R (The R Project for Statistical Computing, 2020). Para obtener los valores perdidos se adopta el método de amputación multivariado descrito en (Schouten et al., 2018) implementado en la función “ampute” del paquete “mice” (Multivariate Imputation by Chained Equations, 2020) del software R (The R Project for Statistical Computing, 2020), parametrizado en un 10%.

2.3. Descripción de la capa de imputación/asignación (E1)

La propuesta mencionada (La Red Martínez, 2017) no considera la información de control faltante, por lo que en este trabajo se propone incorporar una capa de imputación/asignación a un modelo de decisión para imputar los valores de las variables de categorías de carga computacional y los valores de las variables de criterios de carga computacional, como se observa en la Figura 3.

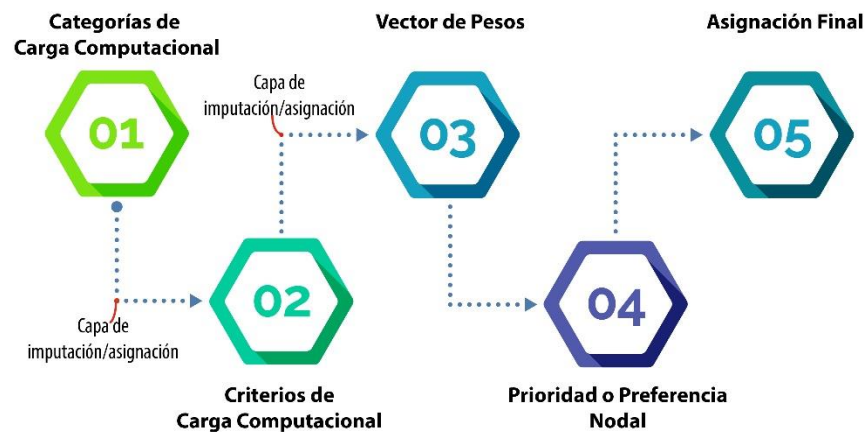


Figura 3. Modelo de decisión.

Fuente: Elaboración propia

Para imputar los valores de las variables de categorías de carga computacional y los valores de las variables de criterios de carga computacional la capa de imputación/asignación cuenta con las siguientes etapas:

2.3.1. Identificar el propósito del problema.

2.3.2. Identificar las alternativas.

2.3.3. Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa.

A continuación, se describirá cada una de las etapas mencionadas.

2.3.1. Identificar el propósito del problema

En esta investigación se pretende resolver las solicitudes de los diferentes recursos por parte de los procesos ubicados en los nodos del sistema distribuido. La información de control compartida por los distintos nodos es recibida por el runtime alojado en el nodo central, y a través de un operador de agregación determina y soluciona la siguiente premisa: que los procesos accedan a recursos compartidos en la modalidad de exclusión mutua, incorporando mecanismos de

imputación de datos para los casos en que la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos sea incompleta.

En un ciclo de recolección de información el nodo central puede recibir de alguno de los nodos información incompleta; por ejemplo, alguno de los criterios utilizados para calcular la carga computacional o las prioridades o preferencia de los procesos. Para solucionar la problemática planteada se desea mejorar un modelo de decisión y su correspondiente operador de agregación incorporando una capa de imputación/asignación, para que el mecanismo pueda completar los valores faltantes de la información de control, y finalmente el modelo de decisión pueda establecer un orden de asignación de recursos, respetando la modalidad de exclusión mutua.

En base a revisiones bibliográficas y por ser uno de los métodos más fiables y de menor consumo de recursos computacionales, y considerando que es uno de los métodos más utilizados, en razón de sus resultados, en investigaciones donde los resultados obtenidos son comparados con otros métodos, como en las propuestas de (Mehala & Vivekanandan, 2008) y (Primorac et al., 2020), se opta por utilizar el método de imputación K-Means.

2.3.2. *Identificar las alternativas*

Las alternativas ayudaran a resolver la situación problemática que se describió en el punto anterior. Entre las aplicaciones y el sistema operativo de cada nodo se define una interfaz, donde por medio de un runtime incluido en esa interfaz se gestiona los procesos y recursos compartidos, para luego definir el escenario correspondiente.

En un determinado ciclo de recolección de información proporcionada por todos los nodos, el nodo central puede recibir información incompleta de uno o algunos de los criterios utilizados para el cálculo de la carga computacional de cada nodo.

Además, para cada nodo, se consideran distintos criterios para calcular la prioridad o preferencia que tiene cada uno, en las asignaciones de recursos a procesos. Al igual que el caso anterior, algunos de los valores de estos criterios podrían no estar disponibles en un ciclo de recolección de información de control.

En caso de que la información de control suministrada por los nodos esté completa, el modelo de decisión se ejecuta sin necesidad de la imputación/asignación de datos, según la propuesta de (La Red Martínez, 2017) y (Agostini, 2019).

2.3.3. *Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa*

Indicadores de las prestaciones actuales de cada nodo

Para obtener un indicador de las prestaciones actuales de cada nodo se adoptarán las j prestaciones en los n nodos. El runtime del nodo que arranca estará informando al runtime del nodo central sus características, por ejemplo, la velocidad de procesamiento, capacidad de memoria, velocidad de transmisión de datos, velocidad de entrada de salida, entre otros. Los valores que se asumirán para los indicadores de prestaciones de los n nodos se obtienen de la Tabla 1.

Tabla 1. Valores de las características referentes a las prestaciones de cada nodo.

Nodos	Prestaciones				
1	$pres_{11}$	...	$pres_{1j}$	...	$pres_{1s}$
...	...	...	...	...	...
i	$pres_{i1}$	...	$pres_{ij}$	...	$pres_{is}$
...	...	...	...	...	...
n	$pres_{n1}$	...	$pres_{nj}$	...	$pres_{ns}$

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Establecimiento de las prestaciones para todos los nodos.

$prestaciones = \{pres_{ij}\}$ con $i = 1, \dots, n$ (n° de nodos en el sistema distribuido) y $j = 1, \dots, s$ (prestaciones para cada nodo).

Cálculo de la carga computacional actual de los nodos

Para obtener un indicador de la carga computacional actual de cada nodo se adoptarán los e criterios en los n nodos. Los valores que se asumirán para los indicadores de carga computacional de los n nodos y el cálculo de carga promedio para cada nodo se obtienen de la Tabla 2.

Tabla 2. Valoración de criterios para calcular la carga computacional en cada nodo.

Nodos	Criterios			
1	C_{11}	C_{12}	...	C_{1e}
....	...	...	...	...
i	C_{i1}	C_{i2}	...	C_{ie}
....	...	...	...	...
n	C_{n1}	C_{n2}	...	C_{ne}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Eventualmente todos los nodos podrían utilizar el mismo conjunto de criterios.

Cálculo de la carga computacional de cada nodo:

$$carga_i = (\text{valor}(C_{i1}) + \dots + \text{valor}(C_{ie})) / e \text{ con } i = 1, \dots, n$$

Puede presentarse la situación de que en cierto ciclo de recolección el nodo central reciba información incompleta de uno o algunos de los criterios de la Tabla 2, utilizados para el cálculo de la carga computacional, entonces se aplica la imputación de datos para dichos criterios, para imputar se tendrá en cuenta ciertas pautas, sabiendo que son las variables de estados consideradas para los nodos.

Imputación o asignación de valores faltantes de los criterios de carga computacional

En función a la información de control histórica (disponible en el nodo central y que corresponde a la información de carga computacional y preferencias de los procesos en ciclos

anteriores) disponible para los métodos de imputación/asignación, como se indica en la Figura 4, se describirán distintas pautas.

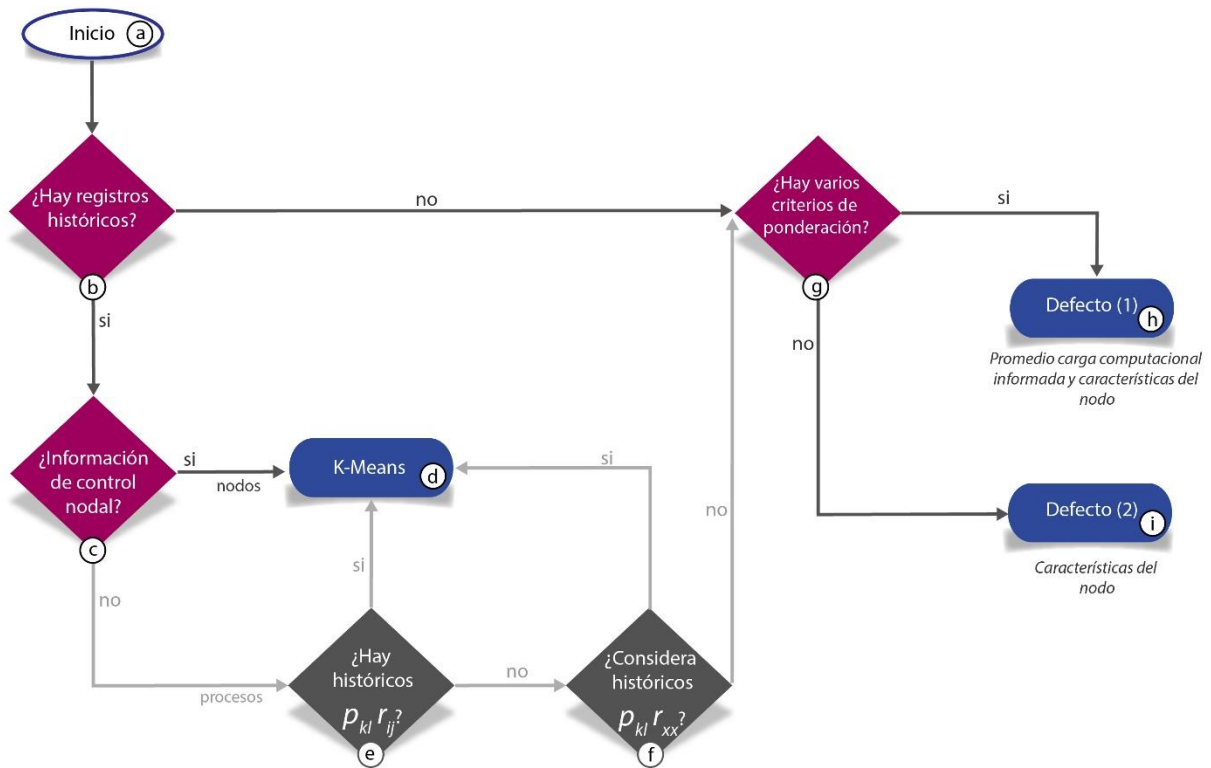


Figura 4. Pautas establecidas para la imputación o asignación de la carga computacional - E1

Fuente: Elaboración propia.

Notas: $p_{kl} r_{ij}$ representa la relación específica entre un proceso y un recurso, mientras que $p_{kl} r_{xx}$ representa la relación de un proceso con respecto a todos los posibles recursos que este solicita.

- Asignación por defecto (1) se utilizará cuando se tenga información sobre otros criterios para obtener el promedio de la carga computacional informada, para luego sumar el valor en función a las características del nodo, y dividir en dos.

- Asignación por defecto (2) se utiliza para este ejemplo por no tener información sobre otros criterios, asignando un valor en función a la característica del nodo.

La carga nodal se refiere a las variables que indican el estado de carga de los nodos, analizando la Figura 4 se establecen las siguientes pautas para resolver el problema del valor faltante:

- Primera pauta; para el criterio kj del nodo k , **si se cuenta** con información histórica en (b), y la información de control es de los nodos en (c), se puede realizar la imputación con K-

Means en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información del nodo k , se obtiene entonces, el valor faltante del criterio kj .

- Segunda pauta; para el criterio kj del nodo k , **si no se cuenta** con información histórica en **(b)**, se realiza la asignación de un valor por **defecto (1)** en **(h)** (se considera cuando se tiene más de un criterio de ponderación), en función de la carga computacional informada y las características del nodo k .
 - Se calcula el promedio de la carga computacional existente basadas en la información disponible de los valores de criterios de carga computacional informados por el nodo central para el criterio kj del nodo k . El promedio carga computacional existente es igual a la suma de los valores disponibles de criterios de carga computacional dividido la cantidad de valores disponibles.
 - Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio kj se asume el valor de acuerdo a sus prestaciones (un valor distinto para los nodos de Alta, Media y Bajas prestaciones).
 - Se promedia los valores del primer y segundo paso. El promedio del primer y segundo paso es igual al valor del promedio de la carga computacional existente sumado al valor de las prestaciones del nodo, todo ello, dividido por la cantidad de criterios disponibles.
- Tercera pauta; para el criterio kj del nodo k , **si no se cuenta** con información histórica en

(b), se puede realizar la asignación de un valor por **defecto (2)** en (i) (se considera cuando se tiene solo un criterio de ponderación), en función de las características del nodo k .

- Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio kj se asume el valor de acuerdo a sus prestaciones.

Al no tener información histórica se está en presencia de un nodo recientemente integrado al sistema, con lo cual se asume que la carga computacional inicial de dicho nodo representa un valor relativamente bajo de consumo de recursos para un nodo de altas prestaciones, un valor intermedio para un nodo de prestaciones medias, y un valor alto para un nodo de prestaciones bajas.

Mediante la imputación o asignación de valores faltantes se obtiene los valores faltantes de los criterios en cuestión, permitiendo de esta manera obtener los indicadores de carga computacional actual de cada nodo como lo indica inicialmente.

El establecimiento de las categorías de carga computacional y de los vectores de pesos asociados a ellos

En esta propuesta, serán utilizados los mismos criterios que en (La Red Martínez, 2017) y (Agostini, 2019).

Para establecer las categorías de carga computacional actual de cada nodo se pueden adoptar distintos criterios; en esta propuesta las categorías serán: Alta (si la carga es mayor al 70%), Media (si la carga está entre el 40% y el 70% inclusive) y Baja (si la carga es menor al 40%), como se verá en el ejemplo.

Establecimiento del n° de categorías para determinar la carga de los nodos:

$$\text{card}(\{\text{categorías}\}) = a$$

Establecimiento de las categorías que se aplicarán (podrán diferir de un nodo a otro):

$\text{categorías} = \{\text{cat}_{ij}\}$ con $i = 1, \dots, n$ (n° de nodos en el sistema distribuido) y $j = 1, \dots, a$ (n° máximo de categorías para cada nodo), lo que se puede expresar mediante la Tabla 3.

Tabla 3. Categorías para medir la carga computacional en cada nodo.

Nodos	Categorías			
1	cat_{11}	cat_{12}	...	cat_{1a}
...	...	...	...	...
i	cat_{i1}	cat_{i2}	...	cat_{ia}
...	...	...	...	...
n	cat_{n1}	cat_{n2}	...	cat_{na}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Eventualmente todos los nodos podrían utilizar el mismo conjunto de categorías.

Los vectores de pesos asociados a las categorías de carga computacional se utilizan para ponderar el impacto en consumo de recursos que significaría asignar cada recurso solicitado al proceso que lo solicita y en función de ello priorizar nodalmente esa asignación. Para establecer los vectores de pesos asociados a las categorías de carga computacional actual de cada nodo se pueden adoptar distintos criterios; en esta propuesta los criterios serán: N° de procesos en el nodo, % de uso de CPU, % de uso de memoria, % de uso de memoria virtual, prioridad del proceso (prioridad del proceso en el nodo donde se ejecuta), sobrecarga de memoria (memoria adicional que requerirá disponer el recurso solicitado, si el dato está disponible), sobrecarga de procesador (uso adicional de procesador que requerirá disponer el recurso solicitado, si el dato está disponible) y sobrecarga de entrada / salida (entrada / salida adicional que requerirá disponer el recurso

solicitado, si el dato está disponible), como se verá en el ejemplo.

En esta propuesta se hará una clasificación de los criterios, que serán agrupados de acuerdo a:

- 1) Criterios relacionados con las variables de estados principales del nodo.
- 2) Criterios referentes a las prioridades de proceso.
- 3) Criterios que hacen referencia a las sobrecargas (valores que indican cuánto le costaría al nodo hacer la asignación de los recursos a los procesos).

Esta clasificación podría variar para cada nodo y para cada circunstancia.

Establecimiento del n° de criterios para determinar la prioridad o preferencia que se otorgará en cada nodo según su carga a cada pedido de un recurso compartido hecho por cada proceso:

$$card(\{critpref\}) = e$$

Establecimiento de los criterios que se aplicarán (iguales para todos los nodos):

criterios para preferencias = $\{cp_{isj}\}$ con $i = 1, \dots, a$ (n° de categorías de carga computacional), $s = 1, \dots, t$ (n° de clasificación de criterios) y $j = 1, \dots, e$ (n° máximo de criterios)

Una vez determinadas las categorías para indicar la carga de los nodos y los criterios que se aplicarán para evaluar la prioridad a otorgar a cada requerimiento de recursos de cada proceso, se podrán establecer los valores correspondientes a los criterios constituyendo así los vectores de pesos para las distintas categorías de carga.

Establecimiento de los vectores de pesos que se aplicarán (iguales para todos los nodos):

$pesos = \{w_{isj}\}$ con $i = 1, \dots, a$ (n° de categorías de carga computacional), $s = 1, \dots, t$ (n° de clasificación de criterios) y $j = 1, \dots, e$ (n° máximo de criterios), lo que se puede expresar mediante la Tabla 4.

Tabla 4. Pesos asignados a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo.

Categoría	Pesos																
	Clasificación 1					...	Clasificación s					...	Clasificación t				
1	w_{111}	...	w_{11j}	...	w_{11e}	...	w_{1s1}	...	w_{1sj}	...	w_{1se}	...	w_{1t1}	...	w_{1tj}	...	w_{1te}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
i	w_{i11}	...	w_{i1j}	...	w_{i1e}	...	w_{is1}	...	w_{isj}	...	w_{ise}	...	w_{it1}	...	w_{itj}	...	w_{ite}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
a	w_{a11}	...	w_{a1j}	...	w_{a1e}	...	w_{as1}	...	w_{asj}	...	w_{ase}	...	w_{at1}	...	w_{atj}	...	w_{ate}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

La asignación de pesos a los distintos criterios será función de estudios estadísticos previamente realizados acerca del sistema distribuido; habrá entonces una función de asignación de pesos a los criterios para constituir los vectores de pesos de cada categoría de carga:

$w_{isj} = norm(función(cp_{isj}))$ con $i = 1, \dots, a$ (n° de categorías de carga computacional), $s = 1, \dots, t$ (n° de clasificación de criterios) y $j = 1, \dots, e$ (n° máximo de criterios); *norm* indica que los valores deben estar normalizados (en el intervalo de 0 a 1 inclusive) y con la restricción de que la sumatoria de los elementos de un vector de pesos debe dar 1:

$\sum \{w_{isj}\} = 1$ con $s = 1, \dots, t$ para cada i constante y con $j = 1, \dots, e$ para cada s constante, esta clasificación varía para cada nodo y para cada circunstancia, cada criterio pertenece a una sola clasificación.

Esto significa que la sumatoria de los pesos asignados a los distintos criterios será 1 para cada una de las categorías, o lo que es lo mismo, que la suma de elementos del vector de pesos de

cada categoría es 1.

Cálculo de las prioridades o preferencias de los procesos teniendo en cuenta el estado del nodo

Las prioridades o preferencias se calculan en cada nodo para cada solicitud de recursos originada en cada proceso, se establecen los criterios que se aplicarán (iguales para todos los nodos y que pueden pertenecer a una sola clasificación):

valoraciones ($r_{ij} p_{kl}$) = $\{c_{smu}\}$ con $i = 1, \dots, n$ (nodo donde reside el recurso), $j = 1, \dots, r$ (recurso en el nodo i), $k = 1, \dots, n$ (nodo donde reside el proceso), $l = 1, \dots, p$ (proceso en el nodo k), $s = 1, \dots, t$ (nº de clasificación de los criterios), $m = 1, \dots, e$ (valoración de criterios), $o = 1, \dots, u$ (nº de valoración ($r_{ij} p_{kl}$), que será único y permitirá identificar la solicitud recurso-proceso) y los que se pueden expresar mediante la Tabla 5.

Tabla 5. Valoraciones asignadas a los criterios según la clasificación, para calcular las prioridades o preferencias de los procesos.

Recursos Procesos	Criterios																
	Clasificación 1					...	Clasificación s					...	Clasificación t				
$r_{11}p_{11}$	C_{111}	...	C_{1m1}	...	C_{1e1}	...	C_{s11}	...	C_{sm1}	...	C_{se1}	...	C_{t11}	...	C_{tm1}	...	C_{te1}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{ij}p_{kl}$	C_{11o}	...	C_{1mo}	...	C_{1eo}	...	C_{s1o}	...	C_{smo}	...	C_{seo}	...	C_{t1o}	...	C_{tmo}	...	C_{teo}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{np}p_{rp}$	C_{11u}	...	C_{1mu}	...	C_{1eu}	...	C_{s1u}	...	C_{smu}	...	C_{seu}	...	C_{t1u}	...	C_{tmu}	...	C_{teu}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Puede presentarse que en un determinado ciclo de recolección de información de control el nodo central reciba de alguno o algunos nodos información incompleta sobre los criterios de evaluación de carga, que son los valores que permiten calcular las prioridades o preferencias de los procesos, entonces se aplica el método de imputación o asignación para completar los valores

faltantes.

Imputación de las valoraciones faltantes de los criterios de prioridad o preferencia de los procesos - E1

Para los métodos de imputación o asignación se describirán distintas pautas, como se indica en la Figura 5.

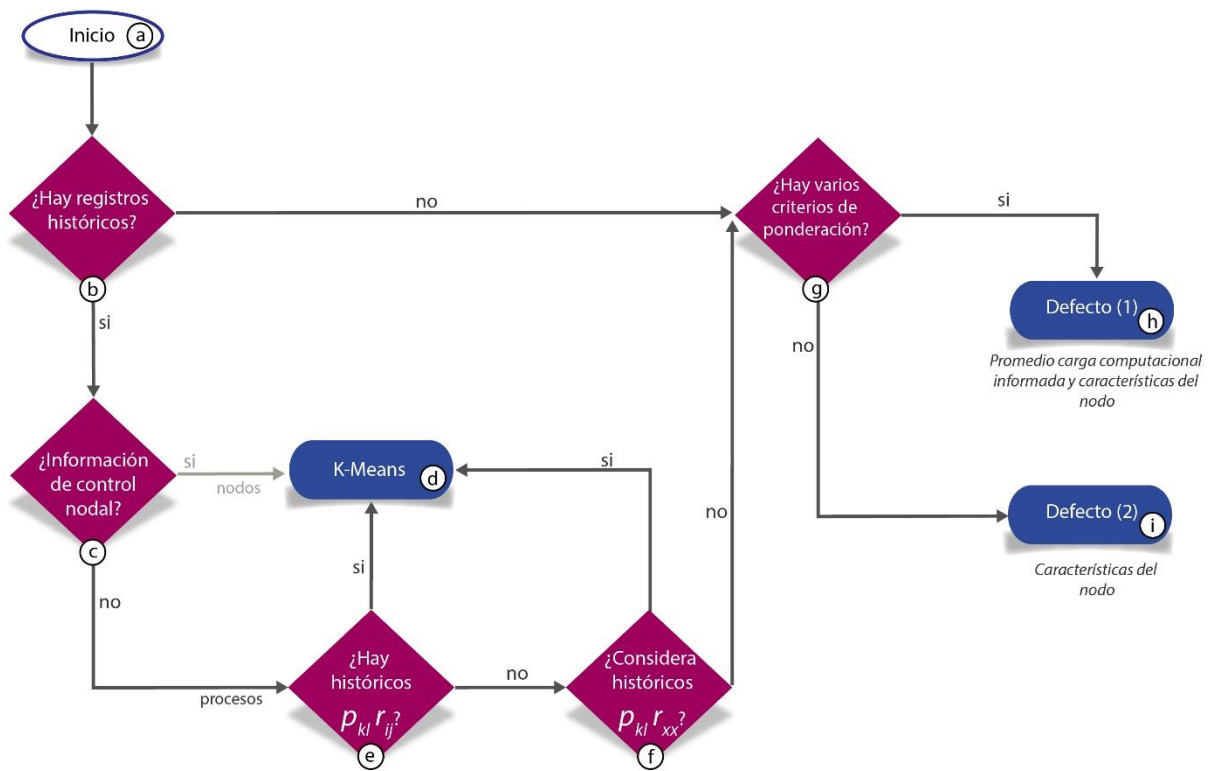


Figura 5. Pautas establecidas para la prioridad o preferencia de procesos – E1

Notas: $p_{kl} r_{ij}$ representa la relación específica entre un proceso y un recurso, mientras que $p_{kl} r_{xx}$ representa la relación de un proceso con respecto a todos los posibles recursos que este solicita.

- Asignación por defecto (1) se utilizará cuando se tenga información sobre otros criterios para obtener el promedio de la carga computacional informada, para luego sumar el valor en función a las características del nodo, y dividir en dos.

- Asignación por defecto (2) se utiliza para este ejemplo por no tener información sobre otros criterios, asignando un valor en función a la característica del nodo.

Para la imputación o asignación de valores faltantes se tendrá en cuenta las **clasificaciones** y sus respectivas pautas, que serán descriptas a continuación.

La **clasificación 1** contempla las variables medidas por el sistema operativo que representan la carga del nodo, por ejemplo, las variables de estados principales del nodo, que son las medidas que toma el sistema operativo cada vez que le llega un requerimiento, que se indican en la Tabla 6.

Tabla 6. Clasificación 1 - Variables de estados principales del nodo.

Recursos Procesos	Criterios																
	Clasificación 1					...	Clasificación s					...	Clasificación t				
$r_{11}p_{11}$	c_{111}	...	c_{1m1}	...	c_{1e1}	...	c_{s11}	...	c_{sm1}	...	c_{se1}	...	c_{t11}	...	c_{tm1}	...	c_{te1}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{ij}p_{kl}$	c_{11o}	...	c_{1mo}	...	c_{1eo}	...	c_{s1o}	...	c_{smo}	...	c_{seo}	...	c_{t1o}	...	c_{tmo}	...	c_{teo}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{np}p_{rp}$	c_{11u}	...	c_{1mu}	...	c_{1eu}	...	c_{s1u}	...	c_{smu}	...	c_{seu}	...	c_{t1u}	...	c_{tmu}	...	c_{teu}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Notas: - Las celdas sombreadas no corresponden a la clasificación 1.

- Los criterios de la Tabla 6 son los mismos que se han utilizado y han sido establecidos en la Tabla 5.

Clasificación 1

Las pautas que se tendrán en cuenta para imputar información de control faltante referente a la clasificación 1 se obtienen analizando la Figura 5.

- Primera pauta; para el criterio c_{m1} de la relación recurso-proceso, **si se cuenta** con información histórica en **(b)**, la información de control en **c** corresponde a los procesos, y hay históricos de $r_{ij}p_{kl}$ en **(e)**, se aplica la imputación con K-Means en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación recurso-proceso $r_{ij}p_{kl}$, se obtiene entonces, el valor faltante del criterio c_{m1} .
- Segunda pauta; para el criterio c_{m1} de la relación recurso-proceso, **si se cuenta** con

información histórica en **(b)**, la información de control en **(c)** corresponde a los procesos, pero **no se cuenta** con información histórica específica de esa relación $r_{ij} p_{kl}$ en **(e)**, se considera **información histórica del proceso con relación a otros recursos** $r_{xx} p_{kl}$ en **(f)**, entonces se aplica la imputación con K-Means en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación proceso p_{kl} con otros recursos r_{xx} , se obtiene entonces, el valor faltante del criterio c_{m1} .

- Tercera pauta; para el criterio c_{m1} de la relación recurso-proceso, **si no se cuenta** con información histórica en **(b)**, pero **no se cuenta** con información histórica específica de esa relación $r_{ij} p_{kl}$ en **(e)**, tampoco se cuenta con **información histórica del proceso con relación a otros recursos** $r_{xx} p_{kl}$ en **(f)**, se puede realizar la asignación de un valor por **defecto (1)** en **(h)**, en función de la carga computacional informada y las características del nodo:
 - Se calcula el promedio de los valores de los criterios c_{m1} (correspondientes a la clasificación 1), disponibles para la relación recurso-proceso $r_{ij} p_{kl}$. El promedio de los criterios existentes es igual a la suma de los valores disponibles de criterios c_{m1} , dividido por la cantidad de criterios disponibles.
 - Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio c_{m1} se asume el valor correspondiente.
 - Se promedia los valores del primer y segundo paso, esto es igual al valor del

promedio de los criterios disponibles sumado al valor de las prestaciones del nodo, todo ello, dividido por la cantidad de criterios disponibles.

- Cuarta pauta; para el criterio cm_1 de la relación recurso-proceso, **si no se cuenta** con información histórica en **(b)**, se puede realizar la asignación de un valor por **defecto (2)** en **(i)**, en función de las características del nodo k .

Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio cm_1 se asume el valor de acuerdo a sus prestaciones.

La **Clasificación s** considera los valores del sistema operativo asignados por defecto o asignado externamente, por ejemplo, la prioridad del proceso, que se indica en la Tabla 7.

Tabla 7. Clasificación s - Valores del sistema operativo asignados por defecto, cuando se genera el proceso.

Recursos Procesos	Criterios																
	Clasificación 1					...	Clasificación s					...	Clasificación t				
$r_{11}p_{11}$	C_{111}	...	C_{1m1}	...	C_{1e1}	...	C_{s11}	...	C_{sm1}	...	C_{se1}	...	C_{t11}	...	C_{tm1}	...	C_{te1}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{ij}p_{kl}$	C_{11o}	...	C_{1mo}	...	C_{1eo}	...	C_{s1o}	...	C_{smo}	...	C_{seo}	...	C_{t1o}	...	C_{tmo}	...	C_{teo}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{np}p_{rp}$	C_{11u}	...	C_{1mu}	...	C_{1eu}	...	C_{s1u}	...	C_{smu}	...	C_{seu}	...	C_{t1u}	...	C_{tmu}	...	C_{teu}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Notas: - Las celdas sombreadas no corresponden a la clasificación s.

- Los criterios de la Tabla 7 son los mismos que se han utilizado y han sido establecidos en la Tabla 5.

Para los métodos de imputación/asignación de la clasificación 2 se describirán distintas pautas, como se indica en la Figura 5.

Clasificación s

Las pautas que se tendrán en cuenta para imputar información de control faltante referente a la clasificación s se obtienen analizando la Figura 5.

- Primera pauta; para el criterio c_{m1} de la relación recurso-proceso, **si se cuenta** con información histórica en **(b)**, la información de control en **c** corresponde a los procesos, y hay históricos de $r_{ij} p_{kl}$ en **(e)**, se aplica la imputación con K-Means en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación recurso-proceso $r_{ij} p_{kl}$, se obtiene entonces, el valor faltante del criterio c_{m1} .
- Segunda pauta; para el criterio c_{m1} de la relación recurso-proceso, **si se cuenta** con información histórica en **(b)**, la información de control en **c** corresponde a los procesos, pero **no se cuenta** con información histórica específica de esa relación $r_{ij} p_{kl}$ en **(e)**, se considera **información histórica del proceso con relación a otros recursos** $r_{xx} p_{kl}$ en **(f)**, entonces se aplica la imputación con K-Means en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación proceso p_{kl} con otros recursos r_{xx} , se obtiene entonces, el valor faltante del criterio c_{m1} .
- Tercera pauta; para el criterio c_{m1} de la relación recurso-proceso, **si no se cuenta** con información histórica en **(b)**, pero **no se cuenta** con información histórica específica de esa relación $r_{ij} p_{kl}$ en **(e)**, tampoco se cuenta con **información histórica del proceso con relación a otros recursos** $r_{xx} p_{kl}$ en **(f)**, se puede realizar la asignación de un valor por **defecto (1)** en **(h)**, en función de la carga computacional informada y las características del

nodo:

- Se calcula el promedio de los valores de los criterios c_{m1} (correspondientes a la clasificación 2), disponibles para la relación recurso-proceso $r_{ij} p_{kl}$. El promedio de los criterios existentes es igual a la suma de los valores disponibles de criterios c_{m1} , dividido por la cantidad de criterios disponibles.
- Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio c_{m1} se asume el valor correspondiente.
- Se promedia los valores del primer y segundo paso, esto es igual al valor del promedio de los criterios disponibles sumado al valor de las prestaciones del nodo, todo ello, dividido por la cantidad de criterios disponibles.
- Cuarta pauta; para el criterio c_{m1} relación recurso-proceso, **si no se cuenta** con información histórica en **(b)**, se puede realizar la asignación de un valor por **defecto (2)** en **(i)**, en función de las características del nodo k .

Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio ij se asume el valor de acuerdo a sus prestaciones.

Clasificación t

La **clasificación t** considera los valores que indican cuánto le costaría al nodo hacer la asignación, el impacto esperado de atender el requerimiento (se tiene en cuenta la relación de

recurso-proceso), como se indica en la Tabla 8.

Tabla 8. Clasificación t – Valores que indican cuánto costaría al nodo hacer la asignación, el impacto esperado de atender el requerimiento (se tiene en cuenta la relación de recurso-proceso).

Recursos Procesos	Criterios																
	Clasificación 1					...	Clasificación s					...	Clasificación t				
$r_{11}p_{11}$	C_{111}	...	C_{1m1}	...	C_{1e1}	...	C_{s11}	...	C_{sm1}	...	C_{se1}	...	C_{t11}	...	C_{tm1}	...	C_{te1}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{ij}p_{kl}$	C_{11o}	...	C_{1mo}	...	C_{1eo}	...	C_{s1o}	...	C_{smo}	...	C_{seo}	...	C_{t1o}	...	C_{tmo}	...	C_{teo}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{np}p_{rp}$	C_{11u}	...	C_{1mu}	...	C_{1eu}	...	C_{s1u}	...	C_{smu}	...	C_{seu}	...	C_{t1u}	...	C_{tmu}	...	C_{teu}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Notas: - Las celdas sombreadas no corresponden a la clasificación t .

- Los criterios de la Tabla 8 son los mismos que se han utilizado y han sido establecidos en la Tabla 5.

Las pautas que se tendrán en cuenta para imputar información de control faltante referente a la clasificación t se obtienen analizando la Figura 5.

- Primera pauta; para el criterio c_{m1} de la relación recurso-proceso, **si se cuenta** con información histórica en **(b)**, la información de control en **c** corresponde a los procesos, y hay históricos de $r_{ij} p_{kl}$ en **(e)**, se aplica la imputación con K-Means en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación recurso-proceso $r_{ij} p_{kl}$, se obtiene entonces, el valor faltante del criterio c_{m1} .
- Segunda pauta; para el criterio c_{m1} de la relación recurso-proceso, **si no se cuenta** con información histórica en **(b)**, **tampoco se cuenta** con información histórica específica de esa relación $r_{ij} p_{kl}$ en **(e)**, y al no considerar para esta clasificación la información histórica del proceso con relación a otros recursos $r_{xx} p_{kl}$ en **(f)**, se aplica la asignación de un valor por **defecto (1)** en **(h)**, en función de la carga computacional informada y las características

del nodo:

- Se calcula el promedio de los valores de los criterios c_{m1} (correspondientes a la clasificación 3), disponibles para la relación recurso-proceso $r_{ij} p_{kl}$. El promedio de los criterios existentes es igual a la suma de los valores disponibles de criterios c_{m1} , dividido por la cantidad de criterios disponibles.
- Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio c_{m1} se asume el valor correspondiente.
- Se promedia los valores del primer y segundo paso, esto es igual al valor del promedio de los criterios disponibles sumado al valor de las prestaciones del nodo, todo ello, dividido por la cantidad de criterios disponibles.
- Cuarta pauta; para el criterio ij del nodo k , **si no se cuenta** con información histórica en **(b)**, se puede realizar la asignación de un valor por **defecto (2)** en **i**, en función de las características del nodo k .

Se determina la prestación del nodo en base a las características conocidas. Los valores de las prestaciones mencionadas anteriormente, permiten determinar las características del nodo k , por lo tanto, para el criterio ij se asume el valor de acuerdo a sus prestaciones.

Luego de aplicar la imputación de datos para completar los valores faltantes, se puede tener toda la información necesaria para que el modelo de decisión resuelva los distintos escenarios desarrollados en (La Red Martínez et al., 2017), se prosigue con los cálculos previstos en el mismo.

Las prioridades o preferencias se calculan en cada nodo para cada solicitud de recursos originada en cada proceso; el cálculo considera el vector de peso correspondiente según la carga actual del nodo y el vector de los valores concedidos por el nodo según los criterios de evaluación de la solicitud.

Los vectores de valoraciones que se aplicarán para cada requerimiento de un recurso por parte de un proceso, según los criterios establecidos para la determinación de la prioridad que en cada caso y momento fijará el nodo en el cual se produce el requerimiento, son los siguientes:

valoraciones ($r_{ij} p_{kl}$) = $\{cp_{smu}\}$ con $i = 1, \dots, n$ (nodo donde reside el recurso), $j = 1, \dots, r$ (recurso en el nodo i), $k = 1, \dots, n$ (nodo donde reside el proceso), $l = 1, \dots, p$ (proceso en el nodo k), $s = 1, \dots, t$ (nº de clasificación de criterios), $m = 1, \dots, e$ (criterios de valoración de la prioridad del requerimiento) y $o = 1, \dots, u$ (nº de valoración($r_{ij} p_{kl}$), que será único y permitirá identificar la solicitud recurso-proceso), los que se pueden expresar mediante Tabla 9.

Tabla 9. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo.

Recurso Proceso	Criterios																
	Clasificación 1					...	Clasificación s					...	Clasificación t				
$r_{11} p_{11}$	cp_{111}	...	cp_{1m1}	...	cp_{1e1}	...	cp_{s11}	...	cp_{sm1}	...	cp_{se1}	...	cp_{t11}	...	cp_{tm1}	...	cp_{te1}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{ij} p_{kl}$	cp_{11o}	...	cp_{1mo}	...	cp_{1eo}	...	cp_{s1o}	...	cp_{smo}	...	cp_{seo}	...	cp_{t1o}	...	cp_{tmo}	...	cp_{teo}
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
$r_{np} p_{rp}$	cp_{11u}	...	cp_{1mu}	...	cp_{1eu}	...	cp_{s1u}	...	cp_{smu}	...	cp_{seu}	...	cp_{t1u}	...	cp_{tmu}	...	cp_{teu}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Resumiendo, la prioridad nodal (por ser calculada en el nodo en el que se produce la petición) de un proceso para acceder a un recurso determinado (que puede estar en cualquier nodo) se calcula mediante el producto escalar de los vectores mencionados anteriormente:

$prioridad\ nodal\ (r_{ij}\ p_{kl}) = \sum w_{isj} * cp_{smu}$ con o indicando el vector de pesos según la carga del nodo, manteniendo los demás subíndices los significados explicados anteriormente.

Cálculo de las prioridades o preferencias de los procesos para acceder a los recursos compartidos disponibles

Esta etapa se calcula en el administrador centralizado de recursos compartidos, para luego determinar el orden en que se asignarán los recursos y a qué proceso será asignado cada recurso. Se consideran las prioridades nodales calculadas en la etapa anterior para cada requerimiento de acceso a los recursos por parte de los procesos. A partir de estas prioridades nodales se deben calcular las prioridades globales o finales; es decir, con qué prioridad, o sea en qué orden, los recursos solicitados serán otorgados y a qué procesos se hará dicho otorgamiento. Los requerimientos que no puedan ser atendidos por resultar con bajas prioridades, serán nuevamente considerados en la siguiente iteración del método.

Para el cálculo de las prioridades finales se utiliza la Tabla 10, en la cual se colocan las prioridades o preferencias nodales calculadas en la etapa anterior; en esta tabla cada fila contiene la información de las prioridades nodales de los distintos procesos para acceder a un determinado recurso.

Tabla 10. Prioridades nodales de los procesos para acceder a cada recurso.

Recursos	Prioridades Nodales de los Procesos				
r_{11}	p_{11}	...	cp_{kj}	...	p_{np}
...	...	...	...	...	...
r_{ij}	p_{11}	...	cp_{kj}	...	p_{np}
...	...	...	...	...	...
r_{nr}	p_{11}	...	cp_{kj}	...	p_{np}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Seguidamente corresponde calcular el vector de pesos finales que se utilizará en el proceso

final de agregación para determinar el orden o prioridad de acceso a los recursos.

pesos finales = $\{wf_{kl}\}$ con $k = 1, \dots, n$ (n° de nodos) y $l = 1, \dots, p$ (n° máximo de procesos por nodo), lo que se puede expresar mediante la Tabla 11, donde np es el número de procesos en el sistema y prg_i es la prioridad del grupo de procesos al que pertenece el proceso (explicada anteriormente).

Tabla 11. Pesos asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos.

Procesos	Pasos Finales	
	Si integra un grupo de procesos	Si es independiente
p_{11}	$wf_{11} = (prg_i)lnp$	$wf_{11} = 1/np$
...	...	...
p_{kl}	$wf_{kl} = (prg_i)lnp$	$w_{kl} = 1/np$
...	...	...
p_{np}	$wf_{np} = (prg_i)lnp$	$wf_{np} = 1/np$

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El siguiente paso es normalizar los pesos recientemente obtenidos dividiendo cada uno por la sumatoria de todos ellos, lo cual se indica en la Tabla 12.

Tabla 12. Pesos finales normalizados asignados a los procesos.

Procesos	Pesos Finales Normalizados
p_{11}	$nwf_{11} = wf_{11} / \sum wf_{kl}$
...	...
p_{kl}	$nwf_{kl} = wf_{kl} / \sum wf_{kl}$
...	...
p_{np}	$nwf_{np} = wf_{np} / \sum wf_{kl}$

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Es así como se obtiene un vector de pesos normalizados (en el intervalo de 0 a 1 inclusive) y con la restricción de que la sumatoria de los elementos del vector debe dar 1:

$$\sum \{nwf_{kl}\} = 1 \text{ con } k = 1, \dots, n \text{ (n° de nodos) y } l = 1, \dots, p \text{ (n° máximo de procesos por nodo).}$$

Las prioridades nodales indicadas en la Tabla 10 tomadas fila por fila; es decir, respecto

de cada recurso, se multiplicarán escalarmente por el vector de pesos finales normalizados indicado en la Tabla 12 para obtener las prioridades globales finales de acceso de cada proceso a cada recurso y de allí, el orden o prioridad con que se asignarán los recursos y a qué proceso se asignará cada uno de ellos; esto se indica a continuación.

prioridad final global ($r_{ij} p_{kl}$) = $nwf_{kl} * p_{kl}$ con r_{ij} indicando el recurso j del nodo i , p_{kl} el proceso l del nodo k y el producto la prioridad final global de dicho proceso para acceder al mencionado recurso. El mayor de estos productos hechos para los distintos procesos en relación al mismo recurso indicará cuál de los procesos tendrá acceso al recurso.

La sumatoria de todos estos productos en relación al mismo recurso indicará la prioridad que tendrá dicho recurso para ser asignado, en relación a los demás recursos que también tendrán que ser asignados. Esto constituye lo que se denominará Función de Asignación para Sistemas Distribuidos (**FASD**):

$$FASD(r_{ij}) = \sum nwf_{kl} * p_{kl} = \text{prioridad de asignación del recurso } r_{ij}.$$

Calculando la **FASD** para todos los recursos se obtendrá un vector y, ordenando sus elementos de mayor a menor se obtendrá el orden prioritario de asignación de los recursos. Además, como ya se ha indicado, el mayor de los productos $nwf_{kl} * p_{kl}$ respecto de cada recurso indicará el proceso al cual será asignado el recurso, esto se indica en la Tabla 13.

Tabla 13. Orden o prioridad final de asignación de los recursos a procesos.

Orden de asignación de los recursos	Proceso al que se asignará el recurso
1°: r_{ij} del Máximo($FASD(r_{ij})$)	p_{kl} del Máximo($nwf_{kl} * p_{kl}$) para el r_{ij} seleccionado
2°: r_{ij} del Máximo($FASD(r_{ij})$) para los r_{ij} no asignados	p_{kl} del Máximo($nwf_{kl} * p_{kl}$) para el r_{ij} seleccionado
...	...
último: r_{ij} no asignado	p_{kl} del Máximo($nwf_{kl} * p_{kl}$) para el r_{ij} seleccionado

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

2.4. Ejemplo

En la presente sección se analizará detalladamente un ejemplo con la aplicación del operador de agregación propuesto. El sistema de procesamiento distribuido, las estructuras de datos, los recursos y los procesos que se ejecutan en los diferentes nodos, grupos, cardinalidades, criterios y categorías para evaluar las diferentes cargas y cálculos necesarios, son los mencionados en (La Red Martínez, 2017) y en (Agostini, 2019).

El sistema de procesamiento distribuido tiene tres nodos:

$$nodos = \{1, 2, 3\}$$

Los procesos ejecutados en cada uno de los nodos son: tres procesos en el nodo 1, cinco procesos en el nodo 2 y siete procesos en el nodo 3.

procesos = $\{p_{ij}\}$ con i indicando el nodo y j indicando el proceso, lo que se puede expresar a través de la Tabla 14.

Tabla 14. Procesos ejecutados en cada nodo.

Nodos	Procesos						
1	p_{11}	p_{12}	p_{13}				
2	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}		
3	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Varios procesos son independientes y otros constituyen grupos de procesos cooperativos; en este ejemplo se considerarán cuatro grupos, lo que se puede expresar mediante la Tabla 15.

Tabla 15. Procesos ejecutados en cada grupo.

Grupos	Procesos		
1	p_{11}	p_{25}	p_{37}
2	p_{12}	p_{21}	
3	p_{22}	p_{31}	
4	p_{13}	p_{23}	p_{34}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El n° de procesos en cada grupo indica la cardinalidad del grupo y se representa de la siguiente manera:

$$card = \{card(g_i)\} = \{3, 2, 2, 3\} \text{ con } i \text{ indicando el grupo.}$$

La prioridad de los grupos de procesos se considerará que es la cardinalidad de cada grupo y se representa de la siguiente manera:

$$prg = \{prg_i = card(g_i)\} = \{3, 2, 2, 3\} \text{ con } i \text{ indicando el grupo.}$$

Los recursos compartidos disponibles en los nodos son los siguientes: tres recursos en el nodo 1, cuatro recursos en el nodo 2 y tres recursos en el nodo 3.

$recursos = \{p_{ij}\}$ con i indicando el nodo y j indicando el proceso, lo que se puede expresar mediante la Tabla 16.

Tabla 16. Recursos compartidos disponibles en cada nodo.

Nodos	Recursos			
1	r_{11}	r_{12}	r_{13}	
2	r_{21}	r_{22}	r_{23}	r_{24}
3	r_{31}	r_{32}	r_{33}	

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Las solicitudes de recursos por parte de los procesos se muestran en la Tabla 17.

Tabla 17. Recursos solicitados por los procesos.

Recursos	Procesos											
r_{11}	p_{11}	p_{12}	p_{13}	p_{24}	p_{32}	p_{33}	p_{36}	p_{37}				
r_{12}	p_{11}	p_{12}	p_{13}	p_{21}	p_{23}	p_{24}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{23}	p_{13}	p_{21}	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}				
r_{21}	p_{11}	p_{12}	p_{13}	p_{22}	p_{25}	p_{33}	p_{36}	p_{37}				
r_{22}	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{33}	p_{34}	p_{35}	p_{36}			
r_{23}	p_{11}	p_{21}	p_{24}	p_{32}	p_{33}	p_{34}						
r_{24}	p_{11}	p_{23}	p_{24}	p_{34}	p_{35}	p_{36}						
r_{31}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{31}	p_{34}	p_{35}	p_{36}			
r_{32}	p_{12}	p_{23}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}					
r_{33}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{31}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}	

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Seguidamente se desarrollará específicamente cada una de las etapas de cálculo.

Identificar el propósito del problema

El runtime alojado en el nodo central se encarga de recibir y mantener actualizada la información de control de todos los nodos, para luego realizar la asignación de recursos a procesos en la modalidad de exclusión mutua. Para este ejemplo, en un ciclo de recolección de información de gestión, necesaria para ejecutar y asegurar lo mencionado anteriormente, el nodo central recibe de algunos de los nodos información incompleta de los criterios de evaluación de carga nodal y de las prioridades o preferencias de los procesos para cada nodo.

La falta de información en algunos de estos aspectos, serán resueltos mediante la incorporación de una capa de imputación/asignación, a un modelo de decisión existente para gestión de recursos y procesos en sistemas distribuidos, considerando como punto de inicio las premisas y las estructuras de datos mencionadas en (La Red Martínez, 2017) y (Agostini, 2019).

Identificar las alternativas

El runtime del nodo central recibe información incompleta de algunos de los criterios utilizados, tanto para calcular la carga computacional de cada nodo, y/o para calcular las

prioridades o preferencias de los procesos para cada nodo, se consideran diferentes pautas para completar dichos valores, utilizando métodos de imputación o asignación.

En situaciones donde falta la información de algunos de los criterios para calcular la carga computacional de cada nodo, se considera la imputación si se tiene un histórico de información respecto del nodo donde se encuentra el dato faltante, de no ser así, se hace la asignación en función a la carga computacional informada, y las características del nodo.

Al igual que la sección anterior, se consideran distintos criterios para calcular la prioridad o preferencia de los procesos para cada nodo, algunos de los valores de la información de control de estos criterios pueden no estar disponibles. Se aplica el método de imputación o asignación para obtener los valores que sustituyen a los valores faltantes, dependiendo de la clasificación a la que corresponda el criterio afectado con valor faltante.

Los runtime de cada nodo informan al runtime del nodo central sus características, y así se obtiene un indicador de las prestaciones los mismos.

En caso de que la información suministrada por los nodos esté completa, el modelo de decisión se ejecuta si necesidad de la imputación de datos, según la propuesta de (La Red Martínez, 2017) y (Agostini, 2019).

Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa

Indicadores de las prestaciones de cada nodo

Para obtener un indicador de las prestaciones del nodo, se tienen en cuenta sus características, por ejemplo, la velocidad de procesamiento, capacidad de memoria, velocidad de transmisión de datos, velocidad de entrada de salida, entre otros. El runtime del nodo que arranca

está informando al runtime del nodo central sus características. Los detalles de implementación de las mismas son:

- Nodo de prestaciones Altas se asume un valor de 30%.
- Nodo de prestaciones Medias se asume un valor de 50%.
- Nodo de prestaciones Bajas se asume un valor 70%.

Las prestaciones correspondientes a cada nodo se muestran en la Tabla 18.

Tabla 18. Valores de las características referentes a las prestaciones de cada nodo.

Nodos	Prestaciones
1	Bajas
2	Medias
3	Altas

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Establecimiento de las prestaciones para todos los nodos.

$prestaciones = \{Bajas, Medias, Altas\}$.

Cálculo de la carga computacional actual de los nodos

En un ciclo de recolección de información proporcionada por todos los nodos, el nodo central puede recibir información incompleta de algunos de los criterios utilizados para calcular la carga computacional de cada nodo. Teniendo en cuenta que en (La Red Martínez, 2017) no se consideran valores faltantes en la información de los criterios de carga computacional, entonces, para generar un escenario de valores faltantes se procede a realizar la amputación de datos, que según (Schouten et al., 2018) es el proceso por el cual se generan conjuntos de datos con valores faltantes a partir de conjuntos de datos completos.

Considerando que el ejemplo propuesto tiene solo tres nodos, específicamente sobre los

valores de los criterios para medir la carga computacional, la información faltante de los distintos escenarios se ha generado amputando manualmente las valoraciones, se observan valores faltantes en el % de Uso de CPU del nodo 1, % de Uso de Memoria del nodo 2 y el % de Uso de Oper. de E/S del nodo 3, como se indican en la Tabla 19.

Tabla 19. Valores de los criterios de carga computacional con valor faltante de los criterios en el nodo 1, nodo 2 y nodo 3 – (E1).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	X	90,00	75,00
2	45,00	X	65,00
3	10,00	25,00	X

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Para obtener la carga computacional actual de cada nodo se debe solucionar la problemática de valores faltantes mencionados en la Tabla 19. Para la solución de la misma se aplica el método de imputación/asignación teniendo en cuenta las pautas establecidas.

Imputación o asignación de las valoraciones faltantes de los criterios de carga computacional

Después de haber analizado la información con que se cuenta sobre los criterios de carga computacional de los nodos donde existen valores faltantes en la Tabla 19; y teniendo en cuenta las pautas establecidas (página 43) para la imputación o asignación de valores, queda clasificada como se muestra en la Tabla 20.

Tabla 20. Clasificación de valores faltantes de los criterios de carga computacional de los nodos, de acuerdo a las pautas establecidas – (E1).

Valor Faltante	Pauta 1: Hay información histórica	Pauta 2: valor por defecto (1)	Pauta 3: valor por defecto (2)
Nodo 1 - % Uso CPU	SI	---	---
Nodo 2 - % Uso Memoria	NO	SI	---
Nodo 3 - % Uso de Oper. E/S	SI	---	---

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

- **% Uso CPU – Nodo 1:** como se muestra en la clasificación de la Tabla 20, el nodo 1 cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 21.

Tabla 21. Valores históricos de los criterios de carga computacional del **Nodo 1**. Primeros cinco registros de mil – (E1).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	65,00	37,00	67,00
1	18,00	57,00	40,00
1	38,00	31,00	15,00
1	78,00	32,00	47,00
1	21,00	72,00	42,00

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información del Nodo 1, se obtiene el valor faltante del criterio % Uso de CPU – Nodo 1, como se puede ver en la Tabla 22.

Tabla 22. Valores de los criterios de carga computacional con valor imputado del criterio **% Uso de CPU – Nodo 1**, aplicando la pauta uno - (E1).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	80,00 / 75,69	90,00	75,00
2	45,00	X	65,00
3	10,00	25,00	X

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

En las celdas remarcadas de la Tabla 22, el valor de la izquierda corresponde a los valores obtenidos en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

- **% Uso de Memoria – Nodo 2:** como se muestra en la clasificación de la Tabla 20, para este criterio de carga computacional del nodo 2, no se cuenta con información histórica, esto permite

hacer la asignación de valor por defecto en función de la carga computacional informada y las características del nodo, como lo establece la segunda pauta.

Primer paso: se calcula el promedio de la carga computacional en base a la información disponible de los valores de criterios de carga computacional informados en la Tabla 19 para el criterio % Uso de Memoria – Nodo 2.

Promedio carga computacional existente = valores disponibles de criterios de carga computacional (% Uso de CPU + % Uso de Oper. E/S) / Cantidad valores existente.

Promedio carga computacional existente = (45,00% + 65,00%) / 2

Promedio carga computacional existente = **55,00%**.

Segundo paso: se determina la prestación del nodo en base a las características conocidas.

Los valores de las prestaciones que se establecen en la Tabla 18 permiten determinar que el nodo 2 es de prestaciones medias, por tanto, para el criterio % Uso de Memoria – Nodo 2 se asume un valor de **50,00%**.

Tercer paso: se promedia los valores del primer y segundo paso.

Promedio primer paso y segundo paso = (Valor del promedio de la carga computacional existente + valor de las prestaciones del nodo) / cantidad de criterios disponibles

Promedio primer paso y segundo paso = (55% + 50%) / 2

Promedio primer paso y segundo paso = **52,50%**

Asignando el valor obtenido para el criterio % Uso de Memoria – Nodo 2, en base a la carga computacional informada y las características del nodo, se obtiene la Tabla 23.

Tabla 23. Valores de los criterios de carga computacional con valor asignado del **% Uso de Memoria – Nodo 2**, aplicando la pauta dos – (E1).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	X	90,00	75,00
2	45,00	50,00 / 52,50	65,00
3	10,00	25,00	X

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

En las celdas remarcadas de la Tabla 23, el valor de la izquierda corresponde a los valores obtenidos en (La Red Martínez, 2017), y el de la derecha es el resultado de la asignación en base a la carga computacional informada y las características del nodo.

- **%Uso de Oper. E/S – Nodo 3:** como se muestra en la clasificación de la Tabla 20, el nodo 3 cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo se cuenta con mil registros de valores históricos, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros, ver la Tabla 24.

Tabla 24. Valores históricos del criterio de carga computacional del Nodo 3. Primeros cinco registros de mil – (E1).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
3	88,00	31,00	31,00
3	82,00	21,00	48,00
3	36,00	79,00	80,00
3	14,00	53,00	64,00
3	90,00	62,00	39,00

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información del Nodo 3, se obtiene el valor faltante del criterio **% Uso de Oper. E/S – Nodo 3**, como se puede ver en la Tabla 25.

En las celdas remarcadas de la Tabla 25, el valor de la izquierda corresponde a los valores obtenidos en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

Tabla 25. Valores de los criterios de carga computacional con valor imputado del criterio % Uso de Oper. E/S – Nodo 3, aplicando la pauta uno – (E1).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	X	90,00	75,00
2	45,00	X	65,00
3	10,00	25,00	35,00 / 69,47

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Los valores que se asumirán para los indicadores de carga computacional de los tres nodos y el cálculo de carga promedio para cada nodo, posterior a la imputación/asignación se muestran en la Tabla 26.

Para obtener un indicador de la carga computacional actual de cada nodo se adoptarán los mismos tres criterios en los tres nodos:

$$\text{card}(\{\text{criterios}\}) = 3$$

$\text{criterios} = \{\% \text{ de uso de la CPU, el } \% \text{ de uso de la memoria, } \% \text{ de uso de operaciones de entrada / salida}\}.$

Tabla 26. Valores de los criterios de carga computacional en cada nodo, con valoraciones imputados/asignado – (E1).

Nodos	Valores de los Criterios			Promedio
	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S	
1	75,69	90,00	75,00	80,23
2	45,00	52,50	65,00	54,16
3	10,00	25,00	69,47	34,82

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El establecimiento de las categorías de carga computacional y de los vectores de pesos asociados a ellos

En esta propuesta las categorías serán las mismas para todos los nodos: Alta (si la carga es mayor al 70%), Media (si la carga está entre el 40% y el 70% inclusive) y Baja (si la carga es

menor al 40%).

$$\text{card}(\{\text{categorías}\}) = 3.$$

$$\text{categorías} = \{\text{Alta, Media, Baja}\}.$$

Los valores obtenidos para las categorías de carga en base a los promedios indicados en la Tabla 26, se observa en la Tabla 27.

Tabla 27. Valores de las categorías para medir la carga computacional en cada nodo.

Nodos	Valores de las categorías
1	Alta
2	Media
3	Baja

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Para definir los vectores de pesos asociados a las categorías de carga computacional actual de cada nodo se utilizarán, para todos los nodos y para todas las categorías de carga se hará una clasificación de los siguientes criterios: 1) Clasificación 1 (N° de procesos en el nodo, % de uso de CPU, % de uso de memoria, % de uso de memoria virtual); 2) Clasificación 2 (prioridad del proceso); 3) Clasificación 3 (sobrecarga de memoria, sobrecarga de procesador y sobrecarga de entrada / salida).

Establecimiento del n° de criterios para determinar la prioridad o preferencia que se otorgará en cada nodo según su carga a cada pedido de un recurso compartido hecho por cada proceso:

$$\text{card}(\{\text{critpref}\}) = 8.$$

criterios para preferencias = {N° de procesos en el nodo, % de uso de CPU, % de uso de memoria, % de uso de memoria virtual, prioridad del proceso, sobrecarga de memoria, sobrecarga

de procesador, sobrecarga de entrada / salida}.

A continuación, se deben establecer los valores correspondientes a los criterios constituyendo así los vectores de pesos para las distintas categorías de carga, que serán iguales para todos los nodos, lo cual se indica en la Tabla 28.

Tabla 28. Pesos asignados a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo.

Categorías	Pesos								$\sum w=1$
	Clasificación 1				Clasificación 2	Clasificación 3			
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
Alta	0,050	0,050	0,100	0,500	0,100	0,100	0,050	0,050	1
Media	0,100	0,200	0,300	0,100	0,200	0,050	0,025	0,025	1
Baja	0,100	0,300	0,200	0,200	0,100	0,025	0,025	0,050	1

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

La sumatoria de los pesos asignados a los distintos criterios es 1 para cada una de las categorías, o lo que es lo mismo, que la suma de elementos del vector de pesos de cada categoría es 1.

Cálculo de las prioridades o preferencias de los procesos teniendo en cuenta el estado del nodo - (E1)

Para el mismo ciclo de recolección de información utilizado en la sección anterior (cálculo de carga computacional de cada nodo), tampoco considera valores faltantes en la información de los criterios para calcular las prioridades o preferencias de los procesos para cada nodo; entonces, para generar un escenario de valores faltantes se procede a realizar la amputación de datos, seleccionando para este ejemplo el mecanismo MCAR, parametrizado en un 10%.

Los valores faltantes de los criterios para la relación proceso-recurso se pueden observar en la Tabla 29.

Tabla 29. Valores asignados a los criterios para calcular la prioridad de proceso, amputados con el mecanismo MCAR al 10% - (E1).

Procesos Recursos	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{11}$	0,70	0,50	0,70	0,90	0,80	0,20	0,30	0,40
$p_{11}r_{12}$	0,80	0,70	0,40	0,50	0,30	0,70	0,20	0,40
$p_{11}r_{21}$	0,30	0,40	0,50	0,20	0,90	0,20	0,50	0,70
$p_{11}r_{22}$	0,50	0,50	0,70	0,40	0,80	0,30	0,50	0,60
$p_{11}r_{23}$	X	0,60	0,80	0,80	0,95	0,90	0,70	0,60
$p_{11}r_{24}$	0,30	0,50	0,90	0,20	0,60	0,60	0,70	0,40
$p_{12}r_{11}$	0,40	0,70	0,50	0,90	1,00	0,90	0,80	0,80
$p_{12}r_{12}$	0,20	0,70	0,30	0,70	0,80	0,30	0,80	0,90
$p_{12}r_{21}$	0,70	0,40	0,30	0,70	0,80	0,90	X	0,20
$p_{12}r_{22}$	0,90	0,60	0,70	0,70	0,80	0,20	0,50	0,40
$p_{12}r_{31}$	0,20	0,50	0,70	0,70	0,30	0,20	0,70	0,80
$p_{12}r_{33}$	0,40	0,50	0,70	0,90	0,30	0,40	0,50	0,80
$p_{13}r_{11}$	0,50	0,70	0,70	0,80	0,60	0,50	X	0,80
$p_{13}r_{12}$	0,70	0,80	0,70	0,40	0,90	0,20	0,90	0,70
$p_{13}r_{13}$	0,70	0,60	0,70	0,80	0,90	0,40	0,80	0,70
$p_{13}r_{21}$	0,70	0,40	0,90	0,30	0,50	0,70	0,20	0,30
$p_{13}r_{22}$	0,50	0,90	0,80	0,30	0,50	0,40	0,90	0,30
$p_{13}r_{31}$	0,50	0,70	0,30	0,60	0,80	0,90	0,90	0,50
$p_{13}r_{32}$	0,60	0,90	0,30	0,60	0,40	0,80	0,70	0,80
$p_{13}r_{33}$	0,60	0,20	0,40	0,60	0,90	0,80	0,50	0,80
$p_{21}r_{12}$	0,20	0,10	0,30	0,80	0,70	0,70	0,50	0,80
$p_{21}r_{13}$	0,20	0,40	0,80	0,90	0,70	0,40	0,50	0,20
$p_{21}r_{22}$	0,30	0,50	0,80	0,90	0,50	0,60	0,20	0,40
$p_{21}r_{23}$	0,70	0,50	0,60	0,80	0,50	0,20	0,90	0,40
$p_{21}r_{31}$	0,70	0,50	0,80	0,60	0,90	0,40	0,90	0,90
$p_{21}r_{33}$	0,70	0,40	0,90	0,80	0,40	0,80	0,60	0,60
$p_{22}r_{21}$	0,50	0,40	0,70	0,80	0,60	0,40	0,60	0,90
$p_{22}r_{22}$	X	0,90	0,60	0,80	0,90	0,40	0,20	0,10
$p_{22}r_{31}$	0,50	0,40	0,50	0,20	0,40	0,80	0,20	0,90
$p_{22}r_{33}$	0,80	0,40	0,30	0,20	0,80	0,90	0,50	0,80
$p_{23}r_{12}$	0,50	0,60	0,80	0,30	0,50	0,70	0,50	0,50
$p_{23}r_{24}$	0,60	0,20	0,10	0,70	0,30	0,80	0,90	0,40
$p_{23}r_{31}$	0,30	0,10	0,40	0,80	0,70	0,90	0,40	0,60
$p_{23}r_{32}$	0,40	0,40	0,80	0,20	0,40	0,60	0,60	X
$p_{23}r_{33}$	0,30	0,60	0,80	0,20	0,90	0,60	0,40	0,20
$p_{24}r_{11}$	0,40	0,60	0,70	0,90	0,50	0,70	0,90	0,50
$p_{24}r_{12}$	0,30	0,60	0,70	0,80	X	0,70	0,60	0,90
$p_{24}r_{23}$	0,40	0,90	0,40	0,80	0,80	0,70	0,60	0,30
$p_{24}r_{24}$	0,50	0,80	0,70	0,90	0,30	0,40	0,80	0,70
$p_{25}r_{21}$	0,20	0,80	X	0,90	0,40	0,50	0,60	0,80
$p_{31}r_{13}$	0,60	0,90	0,60	0,90	0,70	0,40	0,90	0,80

Procesos Recursos	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p31r31</i>	0,80	0,30	0,90	0,50	0,70	0,30	0,90	0,70
<i>p31r33</i>	0,40	0,70	0,70	0,50	0,90	0,90	0,70	0,70
<i>p32r11</i>	0,60	0,90	0,80	0,50	0,90	0,70	0,30	0,70
<i>p32r12</i>	0,70	0,40	0,60	0,90	0,80	1,00	0,90	0,70
<i>p32r13</i>	0,80	0,90	1,00	0,70	0,90	0,30	0,50	0,90
<i>p32r23</i>	0,80	0,80	1,00	0,90	0,60	0,80	0,40	0,90
<i>p33r11</i>	0,20	0,70	0,90	0,80	0,60	0,90	1,00	0,30
<i>p33r12</i>	0,90	1,00	0,90	0,70	0,30	0,50	0,70	0,30
<i>p33r13</i>	0,90	0,50	0,70	0,90	0,30	0,40	0,50	X
<i>p33r21</i>	0,40	0,60	0,70	0,80	0,80	0,40	0,80	0,80
<i>p33r22</i>	0,90	0,40	0,70	0,80	0,70	X	0,90	0,70
<i>p33r23</i>	0,40	0,60	0,90	1,00	0,60	0,40	0,70	0,80
<i>p33r32</i>	0,60	0,60	0,70	0,80	0,40	0,50	0,80	0,80
<i>p33r33</i>	0,60	1,00	0,70	0,40	0,60	0,80	0,80	0,90
<i>p34r12</i>	0,80	1,00	0,30	0,50	0,70	0,30	0,80	0,70
<i>p34r13</i>	0,20	1,00	0,80	0,50	0,80	0,60	0,90	0,70
<i>p34r22</i>	0,90	0,80	0,60	0,80	0,90	0,80	0,50	0,60
<i>p34r23</i>	0,40	0,60	0,70	0,80	0,90	0,50	0,90	0,60
<i>p34r24</i>	0,90	0,40	0,70	0,40	0,70	0,50	0,90	0,70
<i>p34r31</i>	0,50	0,60	0,70	0,90	0,70	0,90	0,60	0,70
<i>p34r32</i>	0,80	0,60	0,90	0,50	0,60	0,90	0,30	0,90
<i>p34r33</i>	0,40	0,60	0,80	X	0,90	0,90	0,40	0,30
<i>p35r12</i>	0,20	0,40	0,70	0,40	0,90	0,50	0,70	X
<i>p35r13</i>	0,80	0,90	0,70	0,30	0,90	0,80	0,70	0,40
<i>p35r22</i>	0,30	0,80	0,90	0,40	0,80	0,60	0,30	0,60
<i>p35r24</i>	0,50	X	0,90	0,40	0,60	0,90	0,40	0,70
<i>p35r31</i>	0,90	0,80	0,70	0,40	0,80	0,30	0,40	0,80
<i>p35r32</i>	0,90	0,70	0,80	0,60	0,40	0,90	0,40	0,70
<i>p35r33</i>	0,50	0,70	0,60	0,90	0,80	0,50	0,90	0,70
<i>p36r11</i>	0,80	0,90	0,60	0,50	0,80	0,90	0,90	0,60
<i>p36r12</i>	0,70	0,90	0,40	0,90	0,80	0,90	0,40	0,70
<i>p36r13</i>	0,90	0,90	0,80	0,70	0,80	0,70	0,90	0,70
<i>p36r21</i>	0,80	0,90	0,50	0,30	0,70	0,70	0,90	0,90
<i>p36r22</i>	0,80	0,20	0,10	0,30	0,80	0,70	0,60	0,90
<i>p36r24</i>	0,40	0,70	0,90	0,30	0,80	0,50	0,80	0,90
<i>p36r31</i>	0,50	0,90	0,90	0,70	0,80	0,80	0,80	0,70
<i>p36r32</i>	0,90	0,80	0,60	0,70	0,80	0,50	0,60	0,70
<i>p36r33</i>	0,20	0,70	0,40	0,80	0,80	0,90	0,40	0,50
<i>p37r11</i>	0,90	0,60	0,80	0,60	0,90	0,70	0,90	0,80
<i>p37r12</i>	0,70	0,90	0,80	0,70	0,50	0,80	0,80	0,60
<i>p37r21</i>	0,90	0,70	0,80	0,60	0,60	0,90	0,50	0,40
<i>p37r32</i>	0,80	0,90	0,50	0,30	0,80	0,70	0,40	0,60
<i>p37r33</i>	0,80	0,40	0,60	0,80	0,80	0,60	0,90	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Al no tener información de algunos de los valores de criterios de la Tabla 29, se aplica el método de imputación/asignación teniendo en cuenta la clasificación de los criterios y sus respectivas pautas.

Imputación de las valoraciones faltantes de los criterios de prioridad o preferencia de los procesos - (E1)

Para la imputación o asignación de valores faltantes se tendrá en cuenta las clasificaciones y sus respectivas pautas ya definidas anteriormente.

Clasificación 1 - Variables de estados principales del nodo

Se puede apreciar las valoraciones faltantes en la Tabla 30, que fueron extraídos de la Tabla 29, y que pertenecen a la clasificación 1.

Tabla 30. Valores faltantes para la clasificación 1 – Variables de estados principales del nodo, que son las medidas que toma el sistema operativo cada vez que le llega un requerimiento - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{23}$	x	0,60	0,80	0,80	0,95	0,90	0,70	0,60
$p_{22}r_{22}$	x	0,90	0,60	0,80	0,90	0,40	0,20	0,10
$p_{25}r_{21}$	0,20	0,80	x	0,90	0,40	0,50	0,60	0,80
$p_{34}r_{33}$	0,40	0,60	0,80	x	0,90	0,90	0,40	0,30
$p_{35}r_{24}$	0,50	x	0,90	0,40	0,60	0,90	0,40	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 1.

Después de haber analizado la información con que se cuenta sobre las variables de estados principales de los nodos donde existen valores faltantes para la clasificación 1, y teniendo en cuenta las pautas para la imputación o asignación de valores, la relación de procesos-recurso queda clasificada como se muestra en la Tabla 31.

Tabla 31. Clasificación de valores faltantes para la clasificación 1, de acuerdo a las pautas establecidas – (E1).

Valor Faltante	Pauta 1: Hay históricos $p_{kl} r_{ij}$	Pauta 2: Hay $p_{kl} r_{xx}$	Pauta 3: Valor por defecto (1)	Pauta 3: Valor por defecto (2)
$p_{11}r_{23}$ - N° Proc.	SI	---	---	---
$p_{22}r_{22}$ - N° Proc.	NO	SI	---	---
$p_{25}r_{21}$ - % Mem.	NO	NO	SI	---
$p_{34}r_{33}$ - % MV	NO	SI	---	---
$p_{35}r_{24}$ - % CPU	SI	---	---	---

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Al tener información histórica de la misma relación proceso-recurso, se aplica la imputación teniendo en cuenta la pauta 1, únicamente de no tener dicha información se tendrán en cuenta la información de proceso con otros recursos para así aplicar la pauta 2, y de no tener ambas informaciones recién se puede aplicar la pauta 3.

- $p_{11}r_{23}$: como se muestra en la clasificación de la Tabla 31, para la relación de proceso-recurso del criterio N° **Proc.** se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 32.

Tabla 32. Valores históricos de la relación proceso-recurso $p_{11}r_{23}$. Primeros cinco registros de mil – (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	N° Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{23}$	0,50	0,80	0,30	0,80	0,40	0,30	0,20	0,70
$p_{11}r_{23}$	0,70	0,10	0,70	0,60	1,00	0,20	0,30	0,30
$p_{11}r_{23}$	0,70	0,50	0,70	0,40	0,40	0,40	0,30	0,30
$p_{11}r_{23}$	0,80	0,30	0,30	0,40	0,90	0,90	0,90	0,50
$p_{11}r_{23}$	0,30	1,00	0,60	0,20	0,30	0,90	0,80	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Nº Proc. de p_{11r23}** , como se puede ver en la Tabla 33.

Tabla 33. Valores de la relación proceso-recurso p_{11r23} con valor imputado del criterio Nº Proc., aplicando la pauta uno – (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{11r23}	0,50/0,38	0,60	0,80	0,80	0,95	0,90	0,70	0,60

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

En la celda remarcada de la Tabla 33, para el criterio **Nº Proc.** para p_{11r23} el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

- **p_{22r22}** : como se muestra en la clasificación de la Tabla 31, para la relación de proceso-recurso del criterio de **Nº Proc.** no existe información histórica, pero si se tiene información histórica del proceso con relación a otros recursos, esto permite hacer la imputación con K-Means como lo establece la segunda pauta. Para este ejemplo, la tabla de valores históricos cuenta con tres mil registros, compuesta por las relaciones p_{22r21} , p_{22r31} , y p_{22r33} , se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 34.

Aplicando la imputación de datos con un histórico de valores correspondientes a tres mil ciclos de recolección de información de las relaciones de proceso con otros recursos p_{22r21} , p_{22r31} , y p_{22r33} , se obtiene el valor faltante del criterio **Nº Proc. de p_{22r22}** , como se puede ver en la Tabla 35.

Tabla 34. Valores históricos de la relación proceso con otros recursos de $p_{22}r_{22}$. Primeros cinco registros de tres mil – (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{22}r_{21}$	0,20	0,50	0,50	0,30	1,00	0,20	0,80	0,60
$p_{22}r_{31}$	0,40	0,90	0,10	0,40	0,50	0,70	0,80	0,60
$p_{22}r_{33}$	0,90	1,00	0,60	0,20	1,00	0,70	0,80	0,70
$p_{22}r_{21}$	0,60	0,90	0,50	0,30	0,50	1,00	0,30	0,90
$p_{22}r_{31}$	0,60	0,70	0,40	0,60	0,60	0,60	0,40	0,90

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 35. Valores de la relación proceso-recurso $p_{22}r_{22}$ con valor imputado del criterio Nº **Proc.**, aplicando la pauta dos – (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{22}r_{22}$	0,50/0,47	0,90	0,60	0,80	0,90	0,40	0,20	0,10

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 1.

En la celda remarcada de la Tabla 35, para el criterio Nº **Proc.** para $p_{22}r_{22}$ el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

- $p_{25}r_{21}$: como se muestra en la clasificación de la Tabla 31, para la relación de proceso-recurso del criterio % **Mem.** no se cuenta con información histórica, tampoco se tiene información histórica del proceso con otros recursos, por lo cual debe hacerse la asignación como lo establece la tercera pauta.

Primer paso: se calcula el promedio de los criterios disponibles (correspondientes a la clasificación 1), disponibles en la Tabla 30 para la relación $p_{25}r_{21}$.

Promedio de los criterios disponibles = (Nº Proc. + % CPU + % MV) / Cantidad valores existente.

$$\text{Promedio de los criterios disponibles} = (0,20 + 0,80 + 0,90) / 3$$

$$\text{Promedio de los criterios disponibles} = \mathbf{0,60}$$

Segundo paso: se determina las prestaciones de cada nodo en base a las características conocidas.

Los valores de las prestaciones que se establece en la Tabla 18.

Tabla **18** que el nodo 1 corresponde a prestaciones bajas, al nodo 2 prestaciones medias, y al nodo 3 prestaciones altas.

El caso de la relación proceso-recurso $p_{25}r_{21}$ está situado en el nodo 2, entonces se determina que pertenece a **prestaciones medias igual a 0,50**.

Tercer paso: se promedia los valores del primer y segundo paso.

$$\text{Promedio primer paso y segundo paso} = (\text{Valor del promedio de los criterios disponibles} + \text{valor de la categoría del nodo}) / 2$$

$$\text{Promedio primer paso y segundo paso} = (0,60 + 0,50) / 2$$

$$\text{Promedio primer paso y segundo paso} = \mathbf{0,55}$$

Asignando el valor obtenido para el criterio N° **Proc.** de la relación proceso-recurso $p_{25}r_{21}$, en base a la carga computacional informada y las características del nodo, se obtiene la Tabla 36.

Tabla 36. Valores de la relación proceso-recurso $p_{25}r_{21}$ con valor asignado del criterio **% Mem.**, aplicando la pauta tres - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	N° Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{25}r_{21}$	0,20	0,80	0,80/0,55	0,90	0,40	0,50	0,60	0,80

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 1.

En la celda remarcada de la Tabla 36, para el criterio **% Mem.** para p_{25r21} el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la asignación en función de la carga computacional informada y las características del nodo.

- p_{34r33} : como se muestra en la clasificación de la Tabla 31, para la relación de proceso-recurso del criterio **% MV** no existe información histórica, pero si se tiene información histórica del proceso con relación a otros recursos, esto permite hacer la imputación con K-Means como lo establece la segunda pauta. Para este ejemplo, la tabla de valores históricos cuenta con siete mil registros, compuesta por p_{34r12} , p_{34r13} , p_{34r22} , p_{34r23} , p_{34r24} , p_{34r31} y p_{34r32} , se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 37.

Tabla 37. Valores históricos de la relación proceso con otros recursos de p_{34r33} . Primeros cinco registros de siete mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{34r12}	0,30	0,10	0,40	1,00	1,00	0,50	0,30	0,20
p_{34r13}	0,80	0,70	0,20	1,00	0,60	0,80	0,60	0,30
p_{34r22}	0,40	1,00	0,80	0,20	0,90	0,20	0,30	0,80
p_{34r23}	0,90	0,90	1,00	0,90	0,40	0,30	0,60	0,90
p_{34r24}	0,70	0,30	0,10	1,00	0,90	0,70	0,50	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a siete mil ciclos de recolección de información de las relaciones de proceso con otros recursos p_{34r12} , p_{34r13} , p_{34r22} , p_{34r23} , p_{34r24} , p_{34r31} y p_{34r32} , se obtiene el valor faltante del criterio **% MV** de p_{34r33} , como se puede ver en la Tabla 38.

En la celda remarcada de la Tabla 38, para el criterio **% MV** para p_{34r33} el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el

resultado de la imputación con K-Means.

Tabla 38. Valores de la relación proceso-recurso $p_{22}r_{22}$ con valor imputado del criterio % MV, aplicando la pauta dos - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{34}r_{33}$	0,40	0,60	0,80	0,50/0,39	0,90	0,90	0,40	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 1.

- $p_{35}r_{24}$: como se muestra en la clasificación de la Tabla 31, para la relación de proceso-recurso del criterio % CPU se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 39.

Tabla 39. Valores históricos de la relación proceso-recurso $p_{35}r_{24}$. Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{35}r_{24}$	0,30	0,80	0,50	0,90	1,00	0,40	1,00	0,20
$p_{35}r_{24}$	0,80	0,30	0,20	1,00	0,90	0,30	0,60	0,40
$p_{35}r_{24}$	0,20	0,50	0,30	0,80	0,90	0,40	0,50	0,70
$p_{35}r_{24}$	0,80	0,50	0,70	0,80	0,50	0,70	0,70	0,20
$p_{35}r_{24}$	0,30	0,80	0,20	0,30	0,60	0,40	0,40	0,40

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio % CPU de la relación proceso-recurso $p_{35}r_{24}$, como se puede ver en la Tabla 40.

En la celda remarcada de la Tabla 40, para el criterio % CPU para $p_{35}r_{24}$ el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el

resultado de la imputación con K-Means.

Tabla 40. Valores de la relación proceso-recurso p_{35r24} con valor imputado del criterio % CPU, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{35r24}	0,50	0,80/0,46	0,90	0,40	0,60	0,90	0,40	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 1.

Clasificación 2 - Prioridad Proceso

Se puede observar en la Tabla 41 el valor faltante, extraído de la Tabla 29, y que pertenece a la clasificación 2.

Tabla 41. Valor faltante para la clasificación 2 - Prioridad Proceso, son valores asignados por el sistema operativo por defecto cuando se genera el proceso - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	0,30	0,60	0,70	0,80	X	0,70	0,60	0,90

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 1.

Después de haber realizado un análisis de la información con que se cuenta sobre la variable de prioridad proceso, donde falta un valor para la relación p_{24r12} en la clasificación 2, se tendrán en cuenta las pautas para la imputación o asignación de valores, se determina que:

- Para la relación de proceso-recurso del criterio **Prioridad Proc.** en p_{24r12} se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 42.

Tabla 42. Valores históricos de la relación proceso-recurso p_{24r12} . Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	0,30	0,50	0,10	1,00	0,80	0,30	1,00	0,10
p_{24r12}	0,30	0,30	0,10	0,50	0,50	0,70	0,60	0,90
p_{24r12}	0,20	0,20	0,10	1,00	0,30	1,00	0,30	0,30
p_{24r12}	0,70	1,00	0,30	0,80	0,30	0,60	0,90	0,80
p_{24r12}	0,60	0,90	0,70	0,30	0,80	0,30	0,70	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Prioridad Proc.** de la relación proceso-recurso p_{24r12} , como se puede ver en la Tabla 43.

Tabla 43. Valores de la relación proceso-recurso p_{24r12} con valor imputado del criterio **Prioridad Proc.**, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	0,30	0,60	0,70	0,80	0,90/0,62	0,70	0,60	0,90

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 2.

En la celda remarcada de la Tabla 43 para el criterio **Prioridad Proc.** para p_{24r12} el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

Clasificación 3 – Criterios de Sobrecargas

Se puede apreciar en la Tabla 44 las valoraciones faltantes, que fueron extraídas de la Tabla 29, y que pertenecen a la clasificación 3 (Sobrec. Mem., Sobrec. Proc. y Sobrec. E/S).

Mediante la imputación o asignación de valores faltantes se obtienen Sobrec. Mem.,

Sobrec. Proc., Sobrec. E/S, permitiendo de esta manera obtener los valores pertenecientes a la clasificación 3.

Tabla 44. Valores faltantes para la clasificación 3 – Valores que indican cuánto costaría al nodo hacer la asignación (se tiene en cuenta la relación de proceso-recurso) - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{12}r_{21}$	0,70	0,40	0,30	0,70	0,80	0,90	X	0,20
$p_{13}r_{11}$	0,50	0,70	0,70	0,80	0,60	0,50	X	0,80
$p_{23}r_{32}$	0,40	0,40	0,80	0,20	0,40	0,60	0,60	X
$p_{33}r_{13}$	0,90	0,50	0,70	0,90	0,30	0,40	0,50	X
$p_{33}r_{22}$	0,90	0,40	0,70	0,80	0,70	X	0,90	0,70
$p_{35}r_{12}$	0,70	0,40	0,30	0,70	0,80	0,90	X	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 3.

Después de haber analizado las informaciones con que se cuenta sobre las variables de sobrecarga donde faltan valores para la clasificación 3, y teniendo en cuenta las pautas para la imputación o asignación de valores, la relación de proceso-recurso queda clasificada como se puede ver en la Tabla 45.

Tabla 45. Clasificación de valores amputados para la clasificación 3 de acuerdo a las pautas establecidas - (E1).

Valor Faltante	Pauta 1: hay históricos $p_{kl}r_{ij}$	Pauta 2: valor por defecto (1)	Pauta 3: valor por defecto (2)
$p_{12}r_{21}$ - Sobrec. Proc.	SI	---	---
$p_{13}r_{11}$ - Sobrec. Proc.	SI	---	---
$p_{23}r_{32}$ - Sobrec. E/S	SI	---	---
$p_{33}r_{13}$ - Sobrec. E/S	SI	---	---
$p_{33}r_{22}$ - Sobrec. Mem.	SI	---	---
$p_{35}r_{12}$ - Sobrec. Proc.	No	SI	---

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Al tener información histórica de la misma relación proceso-recurso, se aplica la imputación teniendo en cuenta la pauta 1, únicamente de no tener dicha información se puede aplicar la pauta 2.

- $p_{12}r_{21}$: como se muestra en la clasificación de la Tabla 45, para la relación de proceso-recurso del criterio **Sobrec. Proc.** se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. La tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 46.

Tabla 46. Valores históricos de la relación proceso-recurso $p_{12}r_{21}$. Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{12}r_{21}$	0,70	0,60	0,90	0,60	0,30	0,60	0,30	0,60
$p_{12}r_{21}$	0,40	0,70	0,50	0,60	1,00	0,90	0,40	0,90
$p_{12}r_{21}$	0,50	0,80	1,00	0,70	0,50	0,90	0,60	0,90
$p_{12}r_{21}$	0,20	0,20	0,70	1,00	1,00	0,70	0,70	0,30
$p_{12}r_{21}$	0,50	0,10	0,70	0,50	1,00	0,20	0,20	0,90

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Sobrec. Proc.** de la relación proceso-recurso $p_{12}r_{21}$, como se puede ver en la Tabla 47.

Tabla 47. Valores de la relación proceso-recurso $p_{12}r_{21}$ con valor imputado del criterio **Sobrec. Proc.**, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	Nº Proc.	% MV	Nº Proc.	Sobrec. Mem.	Nº Proc.	Sobrec. E/S
$p_{12}r_{21}$	0,70	0,40	0,30	0,70	0,80	0,90	0,50/0,69	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación 3.

En la celda remarcada de la Tabla 47, para el criterio **Sobrec. Proc.** para $p_{12}r_{21}$ el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

- p_{13r11} : como se muestra en la clasificación de la Tabla 45, para la relación de proceso-recurso del criterio **Sobrec. Proc.** se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 48.

Tabla 48. Valores históricos de la relación proceso-recurso p_{13r11} . Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{13r11}	0,50	0,70	0,60	1,00	0,40	0,80	0,70	0,20
p_{13r11}	0,80	0,70	0,10	0,60	0,40	1,00	0,20	0,20
p_{13r11}	0,60	0,20	0,50	0,30	0,80	0,60	0,80	0,60
p_{13r11}	0,90	0,50	0,90	0,80	0,90	0,30	1,00	0,70
p_{13r11}	0,90	0,20	0,40	0,50	0,30	0,20	0,30	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Sobrec. Proc.** de la relación proceso-recurso p_{13r11} , como se puede ver en la Tabla 49.

Tabla 49. Valores de la relación proceso-recurso p_{13r11} con valor imputado del criterio **Sobrec. Proc.**, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	Nº Proc.	% MV	Nº Proc.	Sobrec. Mem.	Nº Proc.	Sobrec. E/S
p_{13r11}	0,50	0,70	0,70	0,80	0,60	0,50	0,70 / 0,30	0,80

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

En la celda remarcada de la Tabla 49, para el criterio **Sobrec. Proc.** para p_{13r11} el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

- $p_{23}r_{32}$: como se muestra en la clasificación de la Tabla 45, para la relación de proceso-recurso del criterio **Sobrec. E/S** se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 50.

Tabla 50. Valores históricos de la relación proceso-recurso $p_{23}r_{32}$. Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{23}r_{32}$	0,40	0,50	0,40	0,60	0,30	0,90	0,70	0,80
$p_{23}r_{32}$	0,80	0,30	0,70	0,70	0,40	0,20	0,80	0,40
$p_{23}r_{32}$	0,80	0,50	0,90	0,90	0,50	1,00	0,90	0,70
$p_{23}r_{32}$	0,70	0,40	0,60	0,60	0,70	0,40	0,80	0,10
$p_{23}r_{32}$	0,30	0,80	0,90	0,50	0,90	0,40	0,70	0,60

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Sobrec. E/S** de la relación proceso-recurso $p_{23}r_{32}$, como se puede ver en la Tabla 51.

Tabla 51. Valores de la relación proceso-recurso $p_{23}r_{32}$ con valor imputado del criterio **Sobrec. E/S**, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	Nº Proc.	% MV	Nº Proc.	Sobrec. Mem.	Nº Proc.	Sobrec. E/S
$p_{23}r_{32}$	0,40	0,40	0,80	0,20	0,40	0,60	0,60	0,60/0,25

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación.

En la celda remarcada de la Tabla 51, para el criterio **Sobrec. E/S** para $p_{23}r_{32}$ el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la imputación con K-Means.

- **p_{33r13}** : como se muestra en la clasificación de la Tabla 45, para la relación de proceso-recurso del criterio **Sobrec. E/S** se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 52.

Tabla 52. Valores históricos de la relación proceso-recurso **p_{33r13}** . Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r13}	0,40	0,10	0,50	0,50	0,50	0,80	0,60	0,50
p_{33r13}	0,60	0,80	0,70	0,40	0,50	0,60	0,40	0,70
p_{33r13}	0,70	0,70	0,40	0,70	0,70	0,80	0,30	0,70
p_{33r13}	0,20	0,40	0,20	0,20	0,60	0,60	0,70	0,70
p_{33r13}	0,80	0,70	1,00	0,90	0,80	0,40	0,40	0,60

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Sobrec. E/S** de la relación proceso-recurso **p_{33r13}** , como se puede ver en la Tabla 53.

Tabla 53. Valores de la relación proceso-recurso **p_{33r13}** con valor imputado del criterio **Sobrec. E/S**, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r13}	0,90	0,50	0,70	0,90	0,30	0,40	0,50	0,80 / 0,77

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación.

En la celda remarcada de la Tabla 53, para el criterio **Sobrec. E/S** para **p_{33r13}** el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el

resultado de la imputación con K-Means.

- **p_{33r22}** : como se muestra en la clasificación de la Tabla 45, para la relación de proceso-recurso del criterio **Sobrec. Mem.** se cuenta con información histórica, esto permite hacer la imputación con K-Means como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con mil registros, se puede apreciar en una tabla reducida las valoraciones de los cinco primeros registros en la Tabla 54.

Tabla 54. Valores históricos de la relación proceso-recurso **p_{33r22}** . Primeros cinco registros de mil - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r22}	0,50	0,60	1,00	0,70	0,60	0,40	1,00	0,80
p_{33r22}	0,40	0,40	0,20	0,30	0,70	0,50	1,00	0,90
p_{33r22}	0,30	0,10	0,10	0,40	0,80	0,80	0,80	0,60
p_{33r22}	0,20	0,40	1,00	0,80	0,90	1,00	0,50	0,60
p_{33r22}	0,30	0,60	0,40	0,30	0,90	0,20	0,30	0,80

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a mil ciclos de recolección de información de las relaciones de proceso-recurso, se obtiene el valor faltante del criterio **Sobrec. Mem.** de la relación proceso-recurso **p_{33r22}** , como se puede ver en la Tabla 55.

Tabla 55. Valores de la relación proceso-recurso **p_{33r22}** con valor imputado del criterio **Sobrec. Mem.**, aplicando la pauta uno - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r22}	0,90	0,40	0,70	0,80	0,70	0,40 / 0,86	0,90	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación.

En la celda remarcada de la Tabla 55, para el criterio **Sobrec. Mem.** para **p_{33r22}** el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el

resultado de la imputación con K-Means.

- **p_{35r12}** : como se muestra en la clasificación de la Tabla 45, para la relación de proceso-recurso del criterio **Sobrec. E/S** no se cuenta con información histórica, por lo cual debe hacerse la asignación como lo establece la segunda pauta.

Primer paso: se calcula el promedio de los criterios de sobrecarga disponibles (correspondientes a la clasificación 3) basadas en la información disponible en la Tabla 44 para la relación **p_{35r12}** .

Promedio de los criterios de sobrecarga disponibles = Valores de los criterios de sobrecargas existentes (Sobrec. Mem + Sobrec. Proc) / Cantidad de valores disponibles.

Promedio de los criterios de sobrecarga disponibles = $(0,50 + 0,70) / 2$

Promedio de los criterios de sobrecarga disponibles = **0,60**

Segundo paso: se determina las prestaciones de cada nodo en base a las características conocidas.

Los valores de las prestaciones que se establece en la Tabla 18 indican que el nodo 1 corresponde a prestaciones bajas, el nodo 2 a prestaciones medias, y el nodo 3 a prestaciones altas.

El caso de la relación proceso-recurso **p_{35r12}** está situado en el nodo tres, entonces se determina que pertenece a **prestaciones altas igual a 0,70**.

Tercer paso: se promedia los valores del primer y segundo paso.

Promedio primer paso y segundo paso = (Valor del promedio de los criterios de sobrecarga disponibles + valor de las prestaciones del nodo) / 2

Promedio primer paso y segundo paso = $(0,60 + 0,70) / 2$

Promedio primer paso y segundo paso = **0,65**

Asignando el valor obtenido para el criterio **Sobrec. E/S** de la relación proceso-recurso $p_{35}r_{12}$, en base a los valores de los criterios de sobrecarga disponibles y las características del nodo, se obtiene la Tabla 56.

Tabla 56. Valores de la relación proceso-recurso $p_{35}r_{12}$ con valor imputado del criterio **Sobrec. E/S**, aplicando la pauta dos - (E1).

Proceso Recurso	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{35}r_{12}$	0,20	0,40	0,70	0,40	0,90	0,50	0,70	0,90/ 0,65

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: las celdas sombreadas no corresponden a la clasificación.

En la celda remarcada de la Tabla 56, para el criterio **Sobrec. E/S** para $p_{35}r_{12}$ el valor de la izquierda corresponde al valor obtenido en (La Red Martínez, 2017), y el de la derecha es el resultado de la asignación en función de la carga computacional informada y las características del nodo.

Se obtienen los valores de los criterios mediante la imputación o asignación, que serán utilizados para calcular la prioridad o preferencia de procesos para cada nodo, como se observa en la Tabla 57.

Tabla 57. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia de procesos para cada nodo, con valores imputados/asignados - (E1).

Procesos Recursos	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{11}$	0,70	0,50	0,70	0,90	0,80	0,20	0,30	0,40
$p_{11}r_{12}$	0,80	0,70	0,40	0,50	0,30	0,70	0,20	0,40
$p_{11}r_{21}$	0,30	0,40	0,50	0,20	0,90	0,20	0,50	0,70
$p_{11}r_{22}$	0,50	0,50	0,70	0,40	0,80	0,30	0,50	0,60

Procesos Recursos	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{23}$	0,38	0,60	0,80	0,80	0,95	0,90	0,70	0,60
$p_{11}r_{24}$	0,30	0,50	0,90	0,20	0,60	0,60	0,70	0,40
$p_{12}r_{11}$	0,40	0,70	0,50	0,90	1,00	0,90	0,80	0,80
$p_{12}r_{12}$	0,20	0,70	0,30	0,70	0,80	0,30	0,80	0,90
$p_{12}r_{21}$	0,70	0,40	0,30	0,70	0,80	0,90	0,69	0,20
$p_{12}r_{22}$	0,90	0,60	0,70	0,70	0,80	0,20	0,50	0,40
$p_{12}r_{31}$	0,20	0,50	0,70	0,70	0,30	0,20	0,70	0,80
$p_{12}r_{33}$	0,40	0,50	0,70	0,90	0,30	0,40	0,50	0,80
$p_{13}r_{11}$	0,50	0,70	0,70	0,80	0,60	0,50	0,30	0,80
$p_{13}r_{12}$	0,70	0,80	0,70	0,40	0,90	0,20	0,90	0,70
$p_{13}r_{13}$	0,70	0,60	0,70	0,80	0,90	0,40	0,80	0,70
$p_{13}r_{21}$	0,70	0,40	0,90	0,30	0,50	0,70	0,20	0,30
$p_{13}r_{22}$	0,50	0,90	0,80	0,30	0,50	0,40	0,90	0,30
$p_{13}r_{31}$	0,50	0,70	0,30	0,60	0,80	0,90	0,90	0,50
$p_{13}r_{32}$	0,60	0,90	0,30	0,60	0,40	0,80	0,70	0,80
$p_{13}r_{33}$	0,60	0,20	0,40	0,60	0,90	0,80	0,50	0,80
$p_{21}r_{12}$	0,20	0,10	0,30	0,80	0,70	0,70	0,50	0,80
$p_{21}r_{13}$	0,20	0,40	0,80	0,90	0,70	0,40	0,50	0,20
$p_{21}r_{22}$	0,30	0,50	0,80	0,90	0,50	0,60	0,20	0,40
$p_{21}r_{23}$	0,70	0,50	0,60	0,80	0,50	0,20	0,90	0,40
$p_{21}r_{31}$	0,70	0,50	0,80	0,60	0,90	0,40	0,90	0,90
$p_{21}r_{33}$	0,70	0,40	0,90	0,80	0,40	0,80	0,60	0,60
$p_{22}r_{21}$	0,50	0,40	0,70	0,80	0,60	0,40	0,60	0,90
$p_{22}r_{22}$	0,47	0,90	0,60	0,80	0,90	0,40	0,20	0,10
$p_{22}r_{31}$	0,50	0,40	0,50	0,20	0,40	0,80	0,20	0,90
$p_{22}r_{33}$	0,80	0,40	0,30	0,20	0,80	0,90	0,50	0,80
$p_{23}r_{12}$	0,50	0,60	0,80	0,30	0,50	0,70	0,50	0,50
$p_{23}r_{24}$	0,60	0,20	0,10	0,70	0,30	0,80	0,90	0,40
$p_{23}r_{31}$	0,30	0,10	0,40	0,80	0,70	0,90	0,40	0,60
$p_{23}r_{32}$	0,40	0,40	0,80	0,20	0,40	0,60	0,60	0,25
$p_{23}r_{33}$	0,30	0,60	0,80	0,20	0,90	0,60	0,40	0,20
$p_{24}r_{11}$	0,40	0,60	0,70	0,90	0,50	0,70	0,90	0,50
$p_{24}r_{12}$	0,30	0,60	0,70	0,80	0,62	0,70	0,60	0,90
$p_{24}r_{23}$	0,40	0,90	0,40	0,80	0,80	0,70	0,60	0,30
$p_{24}r_{24}$	0,50	0,80	0,70	0,90	0,30	0,40	0,80	0,70
$p_{25}r_{21}$	0,20	0,80	0,55	0,90	0,40	0,50	0,60	0,80
$p_{31}r_{13}$	0,60	0,90	0,60	0,90	0,70	0,40	0,90	0,80
$p_{31}r_{31}$	0,80	0,30	0,90	0,50	0,70	0,30	0,90	0,70
$p_{31}r_{33}$	0,40	0,70	0,70	0,50	0,90	0,90	0,70	0,70
$p_{32}r_{11}$	0,60	0,90	0,80	0,50	0,90	0,70	0,30	0,70
$p_{32}r_{12}$	0,70	0,40	0,60	0,90	0,80	1,00	0,90	0,70
$p_{32}r_{13}$	0,80	0,90	1,00	0,70	0,90	0,30	0,50	0,90
$p_{32}r_{23}$	0,80	0,80	1,00	0,90	0,60	0,80	0,40	0,90
$p_{33}r_{11}$	0,20	0,70	0,90	0,80	0,60	0,90	1,00	0,30

Procesos Recursos	Criterios							
	Clasificación 1				Clasificación 2	Clasificación 3		
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r12}	0,90	1,00	0,90	0,70	0,30	0,50	0,70	0,30
p_{33r13}	0,90	0,50	0,70	0,90	0,30	0,40	0,50	0,77
p_{33r21}	0,40	0,60	0,70	0,80	0,80	0,40	0,80	0,80
p_{33r22}	0,90	0,40	0,70	0,80	0,70	0,86	0,90	0,70
p_{33r23}	0,40	0,60	0,90	1,00	0,60	0,40	0,70	0,80
p_{33r32}	0,60	0,60	0,70	0,80	0,40	0,50	0,80	0,80
p_{33r33}	0,60	1,00	0,70	0,40	0,60	0,80	0,80	0,90
p_{34r12}	0,80	1,00	0,30	0,50	0,70	0,30	0,80	0,70
p_{34r13}	0,20	1,00	0,80	0,50	0,80	0,60	0,90	0,70
p_{34r22}	0,90	0,80	0,60	0,80	0,90	0,80	0,50	0,60
p_{34r23}	0,40	0,60	0,70	0,80	0,90	0,50	0,90	0,60
p_{34r24}	0,90	0,40	0,70	0,40	0,70	0,50	0,90	0,70
p_{34r31}	0,50	0,60	0,70	0,90	0,70	0,90	0,60	0,70
p_{34r32}	0,80	0,60	0,90	0,50	0,60	0,90	0,30	0,90
p_{34r33}	0,40	0,60	0,80	0,39	0,90	0,90	0,40	0,30
p_{35r12}	0,20	0,40	0,70	0,40	0,90	0,50	0,70	0,65
p_{35r13}	0,80	0,90	0,70	0,30	0,90	0,80	0,70	0,40
p_{35r22}	0,30	0,80	0,90	0,40	0,80	0,60	0,30	0,60
p_{35r24}	0,50	0,46	0,90	0,40	0,60	0,90	0,40	0,70
p_{35r31}	0,90	0,80	0,70	0,40	0,80	0,30	0,40	0,80
p_{35r32}	0,90	0,70	0,80	0,60	0,40	0,90	0,40	0,70
p_{35r33}	0,50	0,70	0,60	0,90	0,80	0,50	0,90	0,70
p_{36r11}	0,80	0,90	0,60	0,50	0,80	0,90	0,90	0,60
p_{36r12}	0,70	0,90	0,40	0,90	0,80	0,90	0,40	0,70
p_{36r13}	0,90	0,90	0,80	0,70	0,80	0,70	0,90	0,70
p_{36r21}	0,80	0,90	0,50	0,30	0,70	0,70	0,90	0,90
p_{36r22}	0,80	0,20	0,10	0,30	0,80	0,70	0,60	0,90
p_{36r24}	0,40	0,70	0,90	0,30	0,80	0,50	0,80	0,90
p_{36r31}	0,50	0,90	0,90	0,70	0,80	0,80	0,80	0,70
p_{36r32}	0,90	0,80	0,60	0,70	0,80	0,50	0,60	0,70
p_{36r33}	0,20	0,70	0,40	0,80	0,80	0,90	0,40	0,50
p_{37r11}	0,90	0,60	0,80	0,60	0,90	0,70	0,90	0,80
p_{37r12}	0,70	0,90	0,80	0,70	0,50	0,80	0,80	0,60
p_{37r21}	0,90	0,70	0,80	0,60	0,60	0,90	0,50	0,40
p_{37r32}	0,80	0,90	0,50	0,30	0,80	0,70	0,40	0,60
p_{37r33}	0,80	0,40	0,60	0,80	0,80	0,60	0,90	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Obtenidos los valores por medio de la imputación o asignación se multiplica escalarmente cada vector de valoraciones de cada requerimiento por el vector de pesos correspondiente a la categoría de carga actual del nodo para obtener la prioridad según cada criterio y la prioridad nodal

otorgada a cada requerimiento; esto se puede observar en la Tabla 58.

Tabla 58. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia nodal que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo - (E1).

Procesos Recursos	Criterios								Prioridades Nodales
	Clasificación 1				Clasificación 2	Clasificación 3			
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
$p_{11}r_{11}$	0,7000	0,5000	0,7000	0,9000	0,8000	0,2000	0,3000	0,4000	0,7150
$Pri(p_{11}r_{11})$	0,0350	0,0250	0,0700	0,4500	0,0800	0,0200	0,0150	0,0200	
$p_{11}r_{12}$	0,8000	0,7000	0,4000	0,5000	0,3000	0,7000	0,2000	0,4000	
$Pri(p_{11}r_{12})$	0,0400	0,0350	0,0400	0,2500	0,0300	0,0700	0,0100	0,0200	0,4950
$p_{11}r_{21}$	0,3000	0,4000	0,5000	0,2000	0,9000	0,2000	0,5000	0,7000	0,3550
$Pri(p_{11}r_{21})$	0,0150	0,0200	0,0500	0,1000	0,0900	0,0200	0,0250	0,0350	
$p_{11}r_{22}$	0,5000	0,5000	0,7000	0,4000	0,8000	0,3000	0,5000	0,6000	
$Pri(p_{11}r_{22})$	0,0250	0,0250	0,0700	0,2000	0,0800	0,0300	0,0250	0,0300	0,4850
$p_{11}r_{23}$	0,3800	0,6000	0,8000	0,8000	0,9500	0,9000	0,7000	0,6000	0,7790
$Pri(p_{11}r_{23})$	0,0190	0,0300	0,0800	0,4000	0,0950	0,0900	0,0350	0,0300	
$p_{11}r_{24}$	0,3000	0,5000	0,9000	0,2000	0,6000	0,6000	0,7000	0,4000	
$Pri(p_{11}r_{24})$	0,0150	0,0250	0,0900	0,1000	0,0600	0,0600	0,0350	0,0200	0,4050
$p_{12}r_{11}$	0,4000	0,7000	0,5000	0,9000	1,0000	0,9000	0,8000	0,8000	0,8250
$Pri(p_{12}r_{11})$	0,0200	0,0350	0,0500	0,4500	0,1000	0,0900	0,0400	0,0400	
$p_{12}r_{12}$	0,2000	0,7000	0,3000	0,7000	0,8000	0,3000	0,8000	0,9000	
$Pri(p_{12}r_{12})$	0,0100	0,0350	0,0300	0,3500	0,0800	0,0300	0,0400	0,0450	0,6200
$p_{12}r_{21}$	0,7000	0,4000	0,3000	0,7000	0,8000	0,9000	0,6900	0,2000	0,6495
$Pri(p_{12}r_{21})$	0,0350	0,0200	0,0300	0,3500	0,0800	0,0900	0,0345	0,0100	
$p_{12}r_{22}$	0,9000	0,6000	0,7000	0,7000	0,8000	0,2000	0,5000	0,4000	
$Pri(p_{12}r_{22})$	0,0450	0,0300	0,0700	0,3500	0,0800	0,0200	0,0250	0,0200	0,6400
$p_{12}r_{31}$	0,2000	0,5000	0,7000	0,7000	0,3000	0,2000	0,7000	0,8000	0,5800
$Pri(p_{12}r_{31})$	0,0100	0,0250	0,0700	0,3500	0,0300	0,0200	0,0350	0,0400	
$p_{12}r_{33}$	0,4000	0,5000	0,7000	0,9000	0,3000	0,4000	0,5000	0,8000	
$Pri(p_{12}r_{33})$	0,0200	0,0250	0,0700	0,4500	0,0300	0,0400	0,0250	0,0400	0,7000
$p_{13}r_{11}$	0,5000	0,7000	0,7000	0,8000	0,6000	0,5000	0,3000	0,8000	0,6950
$Pri(p_{13}r_{11})$	0,0250	0,0350	0,0700	0,4000	0,0600	0,0500	0,0150	0,0400	
$p_{13}r_{12}$	0,7000	0,8000	0,7000	0,4000	0,9000	0,2000	0,9000	0,7000	
$Pri(p_{13}r_{12})$	0,0350	0,0400	0,0700	0,2000	0,0900	0,0200	0,0450	0,0350	0,5350
$p_{13}r_{13}$	0,7000	0,6000	0,7000	0,8000	0,9000	0,4000	0,8000	0,7000	0,7400
$Pri(p_{13}r_{13})$	0,0350	0,0300	0,0700	0,4000	0,0900	0,0400	0,0400	0,0350	
$p_{13}r_{21}$	0,7000	0,4000	0,9000	0,3000	0,5000	0,7000	0,2000	0,3000	
$Pri(p_{13}r_{21})$	0,0350	0,0200	0,0900	0,1500	0,0500	0,0700	0,0100	0,0150	0,4400
$p_{13}r_{22}$	0,5000	0,9000	0,8000	0,3000	0,5000	0,4000	0,9000	0,3000	0,4500
$Pri(p_{13}r_{22})$	0,0250	0,0450	0,0800	0,1500	0,0500	0,0400	0,0450	0,0150	
$p_{13}r_{31}$	0,5000	0,7000	0,3000	0,6000	0,8000	0,9000	0,9000	0,5000	
$Pri(p_{13}r_{31})$	0,0250	0,0350	0,0300	0,3000	0,0800	0,0900	0,0450	0,0250	0,6300
$p_{13}r_{32}$	0,6000	0,9000	0,3000	0,6000	0,4000	0,8000	0,7000	0,8000	0,6000
$Pri(p_{13}r_{32})$	0,0300	0,0450	0,0300	0,3000	0,0400	0,0800	0,0350	0,0400	
$p_{13}r_{33}$	0,6000	0,2000	0,4000	0,6000	0,9000	0,8000	0,5000	0,8000	

Procesos Recursos	Criterios								Prioridades Nodales
	Clasificación 1				Clasificación 2	Clasificación 3			
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p₁₃r₃₃)</i>	0,0300	0,0100	0,0400	0,3000	0,0900	0,0800	0,0250	0,0400	0,6150
<i>p₂₁r₁₂</i>	0,2000	0,1000	0,3000	0,8000	0,7000	0,7000	0,5000	0,8000	
<i>Pri(p₂₁r₁₂)</i>	0,0200	0,0200	0,0900	0,0800	0,1400	0,0350	0,0125	0,0200	0,4175
<i>p₂₁r₁₃</i>	0,2000	0,4000	0,8000	0,9000	0,7000	0,4000	0,5000	0,2000	
<i>Pri(p₂₁r₁₃)</i>	0,0200	0,0800	0,2400	0,0900	0,1400	0,0200	0,0125	0,0050	0,6075
<i>p₂₁r₂₂</i>	0,3000	0,5000	0,8000	0,9000	0,5000	0,6000	0,2000	0,4000	
<i>Pri(p₂₁r₂₂)</i>	0,0300	0,1000	0,2400	0,0900	0,1000	0,0300	0,0050	0,0100	0,6050
<i>p₂₁r₂₃</i>	0,7000	0,5000	0,6000	0,8000	0,5000	0,2000	0,9000	0,4000	
<i>Pri(p₂₁r₂₃)</i>	0,0700	0,1000	0,1800	0,0800	0,1000	0,0100	0,0225	0,0100	0,5725
<i>p₂₁r₃₁</i>	0,7000	0,5000	0,8000	0,6000	0,9000	0,4000	0,9000	0,9000	
<i>Pri(p₂₁r₃₁)</i>	0,0700	0,1000	0,2400	0,0600	0,1800	0,0200	0,0225	0,0225	0,7150
<i>p₂₁r₃₃</i>	0,7000	0,4000	0,9000	0,8000	0,4000	0,8000	0,6000	0,6000	
<i>Pri(p₂₁r₃₃)</i>	0,0700	0,0800	0,2700	0,0800	0,0800	0,0400	0,0150	0,0150	0,6500
<i>p₂₂r₂₁</i>	0,5000	0,4000	0,7000	0,8000	0,6000	0,4000	0,6000	0,9000	
<i>Pri(p₂₂r₂₁)</i>	0,0500	0,0800	0,2100	0,0800	0,1200	0,0200	0,0150	0,0225	0,5975
<i>p₂₂r₂₂</i>	0,4700	0,9000	0,6000	0,8000	0,9000	0,4000	0,2000	0,1000	
<i>Pri(p₂₂r₂₂)</i>	0,0470	0,1800	0,1800	0,0800	0,1800	0,0200	0,0050	0,0025	0,6945
<i>p₂₂r₃₁</i>	0,5000	0,4000	0,5000	0,2000	0,4000	0,8000	0,2000	0,9000	
<i>Pri(p₂₂r₃₁)</i>	0,0500	0,0800	0,1500	0,0200	0,0800	0,0400	0,0050	0,0225	0,4475
<i>p₂₂r₃₃</i>	0,8000	0,4000	0,3000	0,2000	0,8000	0,9000	0,5000	0,8000	
<i>Pri(p₂₂r₃₃)</i>	0,0800	0,0800	0,0900	0,0200	0,1600	0,0450	0,0125	0,0200	0,5075
<i>p₂₃r₁₂</i>	0,5000	0,6000	0,8000	0,3000	0,5000	0,7000	0,5000	0,5000	
<i>Pri(p₂₃r₁₂)</i>	0,0500	0,1200	0,2400	0,0300	0,1000	0,0350	0,0125	0,0125	0,6000
<i>p₂₃r₂₄</i>	0,6000	0,2000	0,1000	0,7000	0,3000	0,8000	0,9000	0,4000	
<i>Pri(p₂₃r₂₄)</i>	0,0600	0,0400	0,0300	0,0700	0,0600	0,0400	0,0225	0,0100	0,3325
<i>p₂₃r₃₁</i>	0,3000	0,1000	0,4000	0,8000	0,7000	0,9000	0,4000	0,6000	
<i>Pri(p₂₃r₃₁)</i>	0,0300	0,0200	0,1200	0,0800	0,1400	0,0450	0,0100	0,0150	0,4600
<i>p₂₃r₃₂</i>	0,4000	0,4000	0,8000	0,2000	0,4000	0,6000	0,6000	0,2500	
<i>Pri(p₂₃r₃₂)</i>	0,0400	0,0800	0,2400	0,0200	0,0800	0,0300	0,0150	0,0062	0,5112
<i>p₂₃r₃₃</i>	0,3000	0,6000	0,8000	0,2000	0,9000	0,6000	0,4000	0,2000	
<i>Pri(p₂₃r₃₃)</i>	0,0300	0,1200	0,2400	0,0200	0,1800	0,0300	0,0100	0,0050	0,6350
<i>p₂₄r₁₁</i>	0,4000	0,6000	0,7000	0,9000	0,5000	0,7000	0,9000	0,5000	
<i>Pri(p₂₄r₁₁)</i>	0,0400	0,1200	0,2100	0,0900	0,1000	0,0350	0,0225	0,0125	0,6300
<i>p₂₄r₁₂</i>	0,3000	0,6000	0,7000	0,8000	0,6200	0,7000	0,6000	0,9000	
<i>Pri(p₂₄r₁₂)</i>	0,0300	0,1200	0,2100	0,0800	0,1240	0,0350	0,0150	0,0225	0,6365
<i>p₂₄r₂₃</i>	0,4000	0,9000	0,4000	0,8000	0,8000	0,7000	0,6000	0,3000	
<i>Pri(p₂₄r₂₃)</i>	0,0400	0,1800	0,1200	0,0800	0,1600	0,0350	0,0150	0,0075	0,6375
<i>p₂₄r₂₄</i>	0,5000	0,8000	0,7000	0,9000	0,3000	0,4000	0,8000	0,7000	
<i>Pri(p₂₄r₂₄)</i>	0,0500	0,1600	0,2100	0,0900	0,0600	0,0200	0,0200	0,0175	0,6275
<i>p₂₅r₂₁</i>	0,2000	0,8000	0,5500	0,9000	0,4000	0,5000	0,6000	0,8000	
<i>Pri(p₂₅r₂₁)</i>	0,0200	0,1600	0,1650	0,0900	0,0800	0,0250	0,0150	0,0200	0,5750
<i>p₃₁r₁₃</i>	0,6000	0,9000	0,6000	0,9000	0,7000	0,4000	0,9000	0,8000	
<i>Pri(p₃₁r₁₃)</i>	0,0600	0,2700	0,1200	0,1800	0,0700	0,0100	0,0225	0,0400	0,7725
<i>p₃₁r₃₁</i>	0,8000	0,3000	0,9000	0,5000	0,7000	0,3000	0,9000	0,7000	

Procesos Recursos	Criterios								Prioridades Nodales
	Clasificación 1				Clasificación 2	Clasificación 3			
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p₃₁r₃₁)</i>	0,0800	0,0900	0,1800	0,1000	0,0700	0,0075	0,0225	0,0350	0,5850
<i>p₃₁r₃₃</i>	0,4000	0,7000	0,7000	0,5000	0,9000	0,9000	0,7000	0,7000	
<i>Pri(p₃₁r₃₃)</i>	0,0400	0,2100	0,1400	0,1000	0,0900	0,0225	0,0175	0,0350	0,6550
<i>p₃₂r₁₁</i>	0,6000	0,9000	0,8000	0,5000	0,9000	0,7000	0,3000	0,7000	
<i>Pri(p₃₂r₁₁)</i>	0,0600	0,2700	0,1600	0,1000	0,0900	0,0175	0,0075	0,0350	0,7400
<i>p₃₂r₁₂</i>	0,7000	0,4000	0,6000	0,9000	0,8000	1,0000	0,9000	0,7000	
<i>Pri(p₃₂r₁₂)</i>	0,0700	0,1200	0,1200	0,1800	0,0800	0,0250	0,0225	0,0350	0,6525
<i>p₃₂r₁₃</i>	0,8000	0,9000	1,0000	0,7000	0,9000	0,3000	0,5000	0,9000	
<i>Pri(p₃₂r₁₃)</i>	0,0800	0,2700	0,2000	0,1400	0,0900	0,0075	0,0125	0,0450	0,8450
<i>p₃₂r₂₃</i>	0,8000	0,8000	1,0000	0,9000	0,6000	0,8000	0,4000	0,9000	
<i>Pri(p₃₂r₂₃)</i>	0,0800	0,2400	0,2000	0,1800	0,0600	0,0200	0,0100	0,0450	0,8350
<i>p₃₃r₁₁</i>	0,2000	0,7000	0,9000	0,8000	0,6000	0,9000	1,0000	0,3000	
<i>Pri(p₃₃r₁₁)</i>	0,0200	0,2100	0,1800	0,1600	0,0600	0,0225	0,0250	0,0150	0,6925
<i>p₃₃r₁₂</i>	0,9000	1,0000	0,9000	0,7000	0,3000	0,5000	0,7000	0,3000	
<i>Pri(p₃₃r₁₂)</i>	0,0900	0,3000	0,1800	0,1400	0,0300	0,0125	0,0175	0,0150	0,7850
<i>p₃₃r₁₃</i>	0,9000	0,5000	0,7000	0,9000	0,3000	0,4000	0,5000	0,7700	
<i>Pri(p₃₃r₁₃)</i>	0,0900	0,1500	0,1400	0,1800	0,0300	0,0100	0,0125	0,0385	0,6510
<i>p₃₃r₂₁</i>	0,4000	0,6000	0,7000	0,8000	0,8000	0,4000	0,8000	0,8000	
<i>Pri(p₃₃r₂₁)</i>	0,0400	0,1800	0,1400	0,1600	0,0800	0,0100	0,0200	0,0400	0,6700
<i>p₃₃r₂₂</i>	0,9000	0,4000	0,7000	0,8000	0,8600	0,4000	0,9000	0,7000	
<i>Pri(p₃₃r₂₂)</i>	0,0900	0,1200	0,1400	0,1600	0,0860	0,0100	0,0225	0,0350	0,6635
<i>p₃₃r₂₃</i>	0,4000	0,6000	0,9000	1,0000	0,6000	0,4000	0,7000	0,8000	
<i>Pri(p₃₃r₂₃)</i>	0,0400	0,1800	0,1800	0,2000	0,0600	0,0100	0,0175	0,0400	0,7275
<i>p₃₃r₃₂</i>	0,6000	0,6000	0,7000	0,8000	0,4000	0,5000	0,8000	0,8000	
<i>Pri(p₃₃r₃₂)</i>	0,0600	0,1800	0,1400	0,1600	0,0400	0,0125	0,0200	0,0400	0,6525
<i>p₃₃r₃₃</i>	0,6000	1,0000	0,7000	0,4000	0,6000	0,8000	0,8000	0,9000	
<i>Pri(p₃₃r₃₃)</i>	0,0600	0,3000	0,1400	0,0800	0,0600	0,0200	0,0200	0,0450	0,7250
<i>p₃₄r₁₂</i>	0,8000	1,0000	0,3000	0,5000	0,7000	0,3000	0,8000	0,7000	
<i>Pri(p₃₄r₁₂)</i>	0,0800	0,3000	0,0600	0,1000	0,0700	0,0075	0,0200	0,0350	0,6725
<i>p₃₄r₁₃</i>	0,2000	1,0000	0,8000	0,5000	0,8000	0,6000	0,9000	0,7000	
<i>Pri(p₃₄r₁₃)</i>	0,0200	0,3000	0,1600	0,1000	0,0800	0,0150	0,0225	0,0350	0,7325
<i>p₃₄r₂₂</i>	0,9000	0,8000	0,6000	0,8000	0,9000	0,8000	0,5000	0,6000	
<i>Pri(p₃₄r₂₂)</i>	0,0900	0,2400	0,1200	0,1600	0,0900	0,0200	0,0125	0,0300	0,7625
<i>p₃₄r₂₃</i>	0,4000	0,6000	0,7000	0,8000	0,9000	0,5000	0,9000	0,6000	
<i>Pri(p₃₄r₂₃)</i>	0,0400	0,1800	0,1400	0,1600	0,0900	0,0125	0,0225	0,0300	0,6750
<i>p₃₄r₂₄</i>	0,9000	0,4000	0,7000	0,4000	0,7000	0,5000	0,9000	0,7000	
<i>Pri(p₃₄r₂₄)</i>	0,0900	0,1200	0,1400	0,0800	0,0700	0,0125	0,0225	0,0350	0,5700
<i>p₃₄r₃₁</i>	0,5000	0,6000	0,7000	0,9000	0,7000	0,9000	0,6000	0,7000	
<i>Pri(p₃₄r₃₁)</i>	0,0500	0,1800	0,1400	0,1800	0,0700	0,0225	0,0150	0,0350	0,6925
<i>p₃₄r₃₂</i>	0,8000	0,6000	0,9000	0,5000	0,6000	0,9000	0,3000	0,9000	
<i>Pri(p₃₄r₃₂)</i>	0,0800	0,1800	0,1800	0,1000	0,0600	0,0225	0,0075	0,0450	0,6750
<i>p₃₄r₃₃</i>	0,4000	0,6000	0,8000	0,3900	0,9000	0,9000	0,4000	0,3000	
<i>Pri(p₃₄r₃₃)</i>	0,0400	0,1800	0,1600	0,0780	0,0900	0,0225	0,0100	0,0150	0,5955
<i>p₃₅r₁₂</i>	0,2000	0,4000	0,7000	0,4000	0,9000	0,5000	0,7000	0,6500	

Procesos Recursos	Criterios								Prioridades Nodales
	Clasificación 1				Clasificación 2	Clasificación 3			
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p_{35r12})</i>	0,0200	0,1200	0,1400	0,0800	0,0900	0,0125	0,0175	0,0325	0,5125
<i>p_{35r13}</i>	0,8000	0,9000	0,7000	0,3000	0,9000	0,8000	0,7000	0,4000	
<i>Pri(p_{35r13})</i>	0,0800	0,2700	0,1400	0,0600	0,0900	0,0200	0,0175	0,0200	
<i>p_{35r22}</i>	0,3000	0,8000	0,9000	0,4000	0,8000	0,6000	0,3000	0,6000	0,6975
<i>Pri(p_{35r22})</i>	0,0300	0,2400	0,1800	0,0800	0,0800	0,0150	0,0075	0,0300	
<i>p_{35r24}</i>	0,5000	0,4600	0,9000	0,4000	0,6000	0,9000	0,4000	0,7000	
<i>Pri(p_{35r24})</i>	0,0500	0,1380	0,1800	0,0800	0,0600	0,0225	0,0100	0,0350	0,5755
<i>p_{35r31}</i>	0,9000	0,8000	0,7000	0,4000	0,8000	0,3000	0,4000	0,8000	
<i>Pri(p_{35r31})</i>	0,0900	0,2400	0,1400	0,0800	0,0800	0,0075	0,0100	0,0400	
<i>p_{35r32}</i>	0,9000	0,7000	0,8000	0,6000	0,4000	0,9000	0,4000	0,7000	0,6875
<i>Pri(p_{35r32})</i>	0,0900	0,2100	0,1600	0,1200	0,0400	0,0225	0,0100	0,0350	
<i>p_{35r33}</i>	0,5000	0,7000	0,6000	0,9000	0,8000	0,5000	0,9000	0,7000	
<i>Pri(p_{35r33})</i>	0,0500	0,2100	0,1200	0,1800	0,0800	0,0125	0,0225	0,0350	0,7100
<i>p_{36r11}</i>	0,8000	0,9000	0,6000	0,5000	0,8000	0,9000	0,9000	0,6000	
<i>Pri(p_{36r11})</i>	0,0800	0,2700	0,1200	0,1000	0,0800	0,0225	0,0225	0,0300	
<i>p_{36r12}</i>	0,7000	0,9000	0,4000	0,9000	0,8000	0,9000	0,4000	0,7000	0,7250
<i>Pri(p_{36r12})</i>	0,0700	0,2700	0,0800	0,1800	0,0800	0,0225	0,0100	0,0350	
<i>p_{36r13}</i>	0,9000	0,9000	0,8000	0,7000	0,8000	0,7000	0,9000	0,7000	
<i>Pri(p_{36r13})</i>	0,0900	0,2700	0,1600	0,1400	0,0800	0,0175	0,0225	0,0350	0,8150
<i>p_{36r21}</i>	0,8000	0,9000	0,5000	0,3000	0,7000	0,7000	0,9000	0,9000	
<i>Pri(p_{36r21})</i>	0,0800	0,2700	0,1000	0,0600	0,0700	0,0175	0,0225	0,0450	
<i>p_{36r22}</i>	0,8000	0,2000	0,1000	0,3000	0,8000	0,7000	0,6000	0,9000	0,6650
<i>Pri(p_{36r22})</i>	0,0800	0,0600	0,0200	0,0600	0,0800	0,0175	0,0150	0,0450	
<i>p_{36r24}</i>	0,4000	0,7000	0,9000	0,3000	0,8000	0,5000	0,8000	0,9000	
<i>Pri(p_{36r24})</i>	0,0400	0,2100	0,1800	0,0600	0,0800	0,0125	0,0200	0,0450	0,6475
<i>p_{36r31}</i>	0,5000	0,9000	0,9000	0,7000	0,8000	0,8000	0,8000	0,7000	
<i>Pri(p_{36r31})</i>	0,0500	0,2700	0,1800	0,1400	0,0800	0,0200	0,0200	0,0350	
<i>p_{36r32}</i>	0,9000	0,8000	0,6000	0,7000	0,8000	0,5000	0,6000	0,7000	0,7950
<i>Pri(p_{36r32})</i>	0,0900	0,2400	0,1200	0,1400	0,0800	0,0125	0,0150	0,0350	
<i>p_{36r33}</i>	0,2000	0,7000	0,4000	0,8000	0,8000	0,9000	0,4000	0,5000	
<i>Pri(p_{36r33})</i>	0,0200	0,2100	0,0800	0,1600	0,0800	0,0225	0,0100	0,0250	0,6075
<i>p_{37r11}</i>	0,9000	0,6000	0,8000	0,6000	0,9000	0,7000	0,9000	0,8000	
<i>Pri(p_{37r11})</i>	0,0900	0,1800	0,1600	0,1200	0,0900	0,0175	0,0225	0,0400	
<i>p_{37r12}</i>	0,7000	0,9000	0,8000	0,7000	0,5000	0,8000	0,8000	0,6000	0,7200
<i>Pri(p_{37r12})</i>	0,0700	0,2700	0,1600	0,1400	0,0500	0,0200	0,0200	0,0300	
<i>p_{37r21}</i>	0,9000	0,7000	0,8000	0,6000	0,6000	0,9000	0,5000	0,4000	
<i>Pri(p_{37r21})</i>	0,0900	0,2100	0,1600	0,1200	0,0600	0,0225	0,0125	0,0200	0,7600
<i>p_{37r32}</i>	0,8000	0,9000	0,5000	0,3000	0,8000	0,7000	0,4000	0,6000	
<i>Pri(p_{37r32})</i>	0,0800	0,2700	0,1000	0,0600	0,0800	0,0175	0,0100	0,0300	
<i>p_{37r33}</i>	0,8000	0,4000	0,6000	0,8000	0,8000	0,6000	0,9000	0,3000	0,6475
<i>Pri(p_{37r33})</i>	0,0800	0,1200	0,1200	0,1600	0,0800	0,0150	0,0225	0,0150	

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Cálculo de las prioridades o preferencias de los procesos para acceder a los recursos compartidos disponibles - (E1)

Una vez calculado la prioridad según cada criterio y la prioridad nodal otorgada a cada requerimiento se procede a calcular las prioridades globales o finales, es decir, con qué prioridad, o sea en qué orden, los recursos solicitados serán otorgados, a qué procesos se hará dicho otorgamiento. Para el cálculo de las prioridades finales se utilizan la Tabla 59 y la Tabla 60 (los valores marcados son valores obtenidos con métodos de imputación o asignación, y reemplazan los valores faltantes).

Tabla 59. Prioridades nodales de los procesos para acceder a cada recurso p_{11} al p_{25} - (E1).

Recursos	Prioridades Nodales de los Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{11}	0,7150	0,8250	0,6950	0,0000	0,0000	0,0000	0,6300	0,0000
r_{12}	0,4950	0,6200	0,5350	0,4175	0,0000	0,6000	0,6365	0,0000
r_{13}	0,0000	0,0000	0,7400	0,6075	0,0000	0,0000	0,0000	0,0000
r_{21}	0,3550	0,6495	0,4400	0,0000	0,5975	0,0000	0,0000	0,5750
r_{22}	0,4850	0,6400	0,4500	0,6050	0,6945	0,0000	0,0000	0,0000
r_{23}	0,7790	0,0000	0,0000	0,5725	0,0000	0,0000	0,6375	0,0000
r_{24}	0,4050	0,0000	0,0000	0,0000	0,0000	0,3325	0,6275	0,0000
r_{31}	0,0000	0,5800	0,6300	0,7150	0,4475	0,4600	0,0000	0,0000
r_{32}	0,0000	0,0000	0,6000	0,0000	0,0000	0,5112	0,0000	0,0000
r_{33}	0,0000	0,7000	0,6150	0,6500	0,5075	0,6350	0,0000	0,0000

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Tabla 60. Prioridades nodales de los procesos para acceder a cada recurso p_{31} al p_{37} - (E1).

Recursos	Prioridades Nodales de los Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{11}	0,0000	0,7400	0,6925	0,0000	0,0000	0,7250	0,7200
r_{12}	0,0000	0,6525	0,7850	0,6725	0,5125	0,7475	0,7600
r_{13}	0,7725	0,8450	0,6510	0,7325	0,6975	0,8150	0,0000
r_{21}	0,0000	0,0000	0,6700	0,0000	0,0000	0,6650	0,6950
r_{22}	0,0000	0,0000	0,6635	0,7625	0,6625	0,3775	0,0000
r_{23}	0,0000	0,8350	0,7275	0,6750	0,0000	0,0000	0,0000
r_{24}	0,0000	0,0000	0,0000	0,5700	0,5755	0,6475	0,0000
r_{31}	0,5850	0,0000	0,0000	0,6925	0,6875	0,7950	0,0000
r_{32}	0,0000	0,0000	0,6525	0,6750	0,6875	0,7325	0,6475
r_{33}	0,6550	0,0000	0,7250	0,5955	0,7100	0,6075	0,6125

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Posteriormente se calcula el vector de pesos finales que se utilizará en el proceso final de agregación para determinar el orden o prioridad de acceso a los recursos; esto se muestran en la Tabla 61 y la Tabla 102.

Tabla 61. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, desde p_{11} al p_{25} - (E1).

Pesos	Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
Finales	0,2000	0,1330	0,2000	0,1330	0,1330	0,2000	0,0670	0,2000
Finales Normalizados	0,0970	0,0650	0,0970	0,0650	0,0650	0,0970	0,0320	0,0970

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 62. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, desde p_{31} al p_{37} - (E1).

Pesos	Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
Finales	0,1330	0,0670	0,0670	0,2000	0,0670	0,0670	0,2000
Finales Normalizados	0,0650	0,0320	0,0320	0,0970	0,0320	0,0320	0,0970

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019)

Las prioridades nodales indicadas en la Tabla 59 y la Tabla 60 tomadas fila por fila, es decir, respecto de cada recurso, se multiplicarán escalarmente por el vector de pesos finales normalizados indicados en la Tabla 61 y la Tabla 62, para obtener las prioridades globales finales de acceso de cada proceso a cada recurso y de allí, el orden o prioridad con que se asignarán los recursos y a qué proceso se asignará cada uno de ellos; esto se indican en la Tabla 63 y la Tabla 64.

Tabla 63. Prioridades globales finales de los procesos para acceder a cada recurso, desde p_{11} al p_{25} - (E1).

Recursos	Prioridades Nodales de los Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{11}	0,0694	0,0536	0,0674	0,0000	0,0000	0,0000	0,0202	0,0000
r_{12}	0,0480	0,0403	0,0520	0,0271	0,0000	0,0582	0,0204	0,0000
r_{13}	0,0000	0,0000	0,0720	0,0395	0,0000	0,0000	0,0000	0,0000
r_{21}	0,0344	0,0422	0,0430	0,0000	0,0390	0,0000	0,0000	0,0558
r_{22}	0,0470	0,0416	0,0440	0,0393	0,0451	0,0000	0,0000	0,0000
r_{23}	0,0756	0,0000	0,0000	0,0372	0,0000	0,0000	0,0210	0,0000
r_{24}	0,0393	0,0000	0,0000	0,0000	0,0000	0,0323	0,0201	0,0000
r_{31}	0,0000	0,0377	0,0611	0,0465	0,0291	0,0450	0,0000	0,0000
r_{32}	0,0000	0,0000	0,0582	0,0000	0,0000	0,0496	0,0000	0,0000
r_{33}	0,0000	0,0455	0,0600	0,0423	0,0330	0,0616	0,0000	0,0000

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Tabla 64. Prioridades globales finales de los procesos para acceder a cada recurso, desde p_{31} al p_{37} - (E1).

Recursos	Prioridades Nodales de los Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{11}	0,0000	0,0240	0,0222	0,0000	0,0000	0,0232	0,0700
r_{12}	0,0000	0,0210	0,0251	0,0652	0,0164	0,0240	0,0740
r_{13}	0,0502	0,0270	0,0208	0,0711	0,0223	0,0261	0,0000
r_{21}	0,0000	0,0000	0,0220	0,0000	0,0000	0,0213	0,0674
r_{22}	0,0000	0,0000	0,0212	0,0740	0,0212	0,0121	0,0000
r_{23}	0,0000	0,0270	0,0233	0,0655	0,0000	0,0000	0,0000
r_{24}	0,0000	0,0000	0,0000	0,0553	0,0184	0,0210	0,0000
r_{31}	0,0380	0,0000	0,0000	0,0672	0,0220	0,0260	0,0000
r_{32}	0,0000	0,0000	0,0210	0,0655	0,0220	0,0240	0,0630
r_{33}	0,0426	0,0000	0,0232	0,0578	0,0230	0,0200	0,0600

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

El mayor de estos productos hechos para los distintos procesos en relación al mismo recurso indicará cuál de los procesos tendrá acceso al recurso (en caso de empates se podría optar por dar ganador al proceso identificado con el número menor); esto se muestran en rojo en la Tabla 63 y la Tabla 64.

La sumatoria de todos estos productos en relación al mismo recurso indicará la prioridad que tendrá dicho recurso para ser asignado. Esto constituye la Función de Asignación para Sistemas Distribuidos (**FASD**), que se muestra en la Tabla 65.

La suma de todos estos productos en relación con el mismo recurso indicará la prioridad que deberá asignarse a este recurso, en relación con los demás recursos que también deberán

asignarse.

Tabla 65. Prioridades globales finales para asignar los recursos, constituye la Función de Asignación para Sistemas Distribuidos (FASD) – (E1).

FASD	Prioridad Global Final Para Asignar el Recurso	Proceso
r_{11}	0,3500	r_{11} al p_{37}
r_{12}	0,4717	r_{12} al p_{37}
r_{13}	0,3290	r_{13} al p_{13}
r_{21}	0,3251	r_{21} al p_{37}
r_{22}	0,3455	r_{22} al p_{34}
r_{23}	0,2496	r_{23} al p_{11}
r_{24}	0,1864	r_{24} al p_{34}
r_{31}	0,3726	r_{31} al p_{34}
r_{32}	0,3033	r_{32} al p_{34}
r_{33}	0,4690	r_{33} al p_{23}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019)

El orden final de asignación de los recursos y los procesos destinatarios de los mismos se obtiene ordenando la Tabla 65 (FASD), que se muestra en la Tabla 66.

Tabla 66. Prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso, constituye la Función de Asignación para Sistemas Distribuidos Ordenada (FASDO) – (E1).

FASDO	Prioridad Global Final Para Asignar el Recurso	Proceso
r_{12}	0,4717	r_{12} al p_{37}
r_{33}	0,4690	r_{33} al p_{23}
r_{31}	0,3726	r_{31} al p_{34}
r_{11}	0,3500	r_{11} al p_{37}
r_{22}	0,3455	r_{22} al p_{34}
r_{13}	0,3290	r_{13} al p_{13}
r_{21}	0,3251	r_{21} al p_{37}
r_{32}	0,3033	r_{32} al p_{34}
r_{23}	0,2496	r_{23} al p_{11}
r_{24}	0,1864	r_{24} al p_{34}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019)

El siguiente paso es reiterar el procedimiento como lo establece (La Red Martínez, 2017), retirando de las solicitudes de recursos las asignaciones ya hechas; también debe tenerse en cuenta que los recursos asignados quedarán disponibles cuando los procesos los hayan liberado, pudiendo por lo tanto ser asignados a otros procesos. El resultado de las sucesivas iteraciones finaliza una vez que se han atendido todas las solicitudes de recursos de todos los procesos respetando la exclusión mutua y las prioridades de los procesos, las prioridades nodales y las prioridades finales;

los resultados de las sucesivas iteraciones se muestran en la Tabla 67, Tabla 68, Tabla 69, Tabla 70, Tabla 71, Tabla 72, Tabla 73, Tabla 74, Tabla 75, Tabla 76 y Tabla 77.

Tabla 67. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (segunda iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,4074	r_{33} al p_{13}
0,3977	r_{12} al p_{34}
0,3054	r_{31} al p_{13}
0,2800	r_{11} al p_{11}
0,2715	r_{22} al p_{11}
0,2577	r_{21} al p_{25}
0,2570	r_{13} al p_{34}
0,2378	r_{32} al p_{37}
0,1740	r_{23} al p_{34}
0,1311	r_{24} al p_{11}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 68. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (tercera iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,3474	r_{33} al p_{37}
0,3325	r_{12} al p_{23}
0,2443	r_{31} al p_{21}
0,2245	r_{22} al p_{22}
0,2106	r_{11} al p_{13}
0,2019	r_{21} al p_{12}
0,1859	r_{13} al p_{31}
0,1748	r_{32} al p_{13}
0,1085	r_{23} al p_{21}
0,0918	r_{24} al p_{23}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 69. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (cuarta iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,2874	r_{33} al p_{34}
0,2743	r_{12} al p_{13}
0,1978	r_{31} al p_{23}
0,1794	r_{22} al p_{13}
0,1597	r_{21} al p_{13}
0,1432	r_{11} al p_{12}
0,1357	r_{13} al p_{21}
0,1166	r_{32} al p_{23}
0,0713	r_{23} al p_{32}
0,0595	r_{24} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 70. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (quinta iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,2296	r_{33} al p_{12}
0,2223	r_{12} al p_{11}
0,1528	r_{31} al p_{31}
0,1354	r_{22} al p_{12}
0,1167	r_{21} al p_{22}
0,0962	r_{13} al p_{32}
0,0896	r_{11} al p_{32}
0,0670	r_{32} al p_{36}
0,0443	r_{23} al p_{33}
0,0385	r_{24} al p_{24}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 71. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (sexta iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,1841	r_{33} al p_{31}
0,1743	r_{12} al p_{12}
0,1148	r_{31} al p_{12}
0,0938	r_{22} al p_{21}
0,0777	r_{21} al p_{11}
0,0692	r_{13} al p_{36}
0,0656	r_{11} al p_{36}
0,0430	r_{32} al p_{35}
0,0210	r_{23} al p_{24}
0,0184	r_{24} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 72. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (séptima iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,1415	r_{33} al p_{21}
0,1340	r_{12} al p_{21}
0,0771	r_{31} al p_{22}
0,0545	r_{22} al p_{33}
0,0433	r_{21} al p_{33}
0,0431	r_{13} al p_{35}
0,0424	r_{11} al p_{33}
0,0210	r_{32} al p_{33}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 73. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (octava iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,1069	r_{12} al p_{33}
0,0992	r_{33} al p_{22}
0,0480	r_{31} al p_{36}
0,0333	r_{22} al p_{35}
0,0213	r_{21} al p_{36}
0,0208	r_{13} al p_{33}
0,0202	r_{11} al p_{24}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 74. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (novena iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,0818	r_{12} al p_{36}
0,0662	r_{33} al p_{33}
0,0220	r_{31} al p_{35}
0,0121	r_{22} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 75. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décima iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,0578	r_{12} al p_{32}
0,0430	r_{33} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 76. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo primera iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,0368	r_{12} al p_{24}
0,0200	r_{33} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 77. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo segunda iteración – E1).

Prioridad Global Final Ordenada	Asignación
0,0164	r_{12} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

De esta manera se han atendido todas las solicitudes de recursos de todos los procesos

respetando la exclusión mutua y las prioridades de los procesos, las prioridades nodales y las prioridades finales.

2.5. Evaluación

El método aplicado en el modelo decisión fue muy efectivo, la diferencia que existe entre los valores originales en (La Red Martínez, 2017) y los valores que fueron imputados fueron mínimas.

Una de las limitaciones más importantes al aplicar K-Means en esta propuesta, es que se ha establecido que tiene que estar disponible mínimamente el 70% de información de control en un registro. Una solución a esta problemática y a la situación en la que falten la totalidad de los datos que debería suministrar un nodo en un ciclo de sondeo, es aplicar un nuevo mecanismo que utiliza Medias Ponderadas, en los cuales los valores de los registros más recientes tendrán mayor peso que los más antiguos.

2.6. Discusiones y comentarios

El modelo de decisión propuesto considera la imputación de datos cuando la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos llega incompleta, reemplazando esos valores faltantes por valores estimados, para que finalmente el modelo de decisión establezca un correcto orden de asignación de recursos a los procesos, respetando la exclusión mutua.

Esta propuesta considera la información estimada (resultado de las imputaciones/asignaciones en los diferentes ciclos) como parte de la información histórica para aplicar una nueva imputación/asignación.

Se aplicará el método de imputación K-Means siempre que:

- Se disponga con más de diez registros históricos de información de control, para aplicar el método de imputación K-Means.
- Sólo puede faltar hasta el 30% de los valores de las variables que se recibe en un ciclo de recepción de información de control que envía los nodos al nodo central (información de control de los nodos y para los procesos).

El contenido de este capítulo sirvió de base para las siguientes publicaciones:

- Duré, D., Fornerón, J., Agostini, F., La Red Martínez, D. (2021). Imputación de información de control faltante en la gestión de recursos y procesos. *Conferencia Ibero Americana WWW/Internet, (CIAWI 2021)*. **(PUBLICADO)**
- Duré, D., Agostini, F., La Red Martínez, D. (2021). Imputación de valores faltantes sobre la información de control en el contexto de los sistemas distribuidos. *Congreso Nacional de Ingeniería Informática y Sistemas de la Información, (9º CONAIISI 2021)*. **(PUBLICADO)**.

Capítulo III - Capa de imputación/asignación para datos inexistentes - escenario 2 (E2)

3.1. Trabajos previos

En (La Red Martínez, 2017) y en (La Red Martínez et al., 2017; La Red Martínez, et al., 2018) se presentan nuevos operadores de agregación para modelos de decisión aplicados a la gestión de recursos compartidos en sistemas distribuidos.

En (Agostini & La Red Martínez, 2019) y (Agostini et al., 2018) se presentan operadores de agregación y modelos de decisión para la gestión de recursos en sistemas distribuidos según distintos escenarios.

3.2. Propuesta de Solución

La imputación de datos será el método a utilizar en el modelo decisión propuesto para resolver el problema de datos faltantes. Se opta por utilizar la *Media Ponderada*, utilizando la información de control referida al estado de los nodos de los últimos registros históricos.

Se decide seleccionar el mecanismo mencionado para el escenario E2 en base a revisiones bibliográficas, y lo más importante, que en el momento de realizar los cálculos permitirá dar más peso a la información de control más reciente y menos peso a la información de control más antigua, algunas investigaciones publicadas que han considerado la media ponderada se pueden mencionar en:

En (Abril et al., 2014) presentan una variación del enlace de registros (hacen referencia a un mismo individuo, entre diferentes bases de datos) basada en una media ponderada para calcular distancias entre registros.

En (Gómez et al., 2014) muestran una aplicación de las condiciones de estabilidad estricta

al problema de pérdida de datos en un proceso de agregación de información, utilizando las familias de la media ponderada y de la media ponderada cuasi-aritmética.

En (Roy et al., 2017) proponen una combinación de filtro vectorial mediano adaptativo (VMF) y filtro de media ponderada para eliminar el ruido impulsivo de alta densidad de las imágenes en color (un píxel no ruidoso se sustituye por la media ponderada de los píxeles buenos de la ventana de procesamiento).

En (Forc et al., 2018) proponen la adaptación automática del operador de pooling aprendiendo pesos de medias ponderadas ordenadas en Redes Neuronales Convolucionales.

En (Baez et al., 2019) abordan mediante una revisión teórica de trabajos de investigación, la evolución de la toma de decisiones empresariales a través de la media ponderada ordenada.

En (Yoo et al., 2020) se propone un enfoque de conjunto de medias ponderadas sustitutivas para reducir la incertidumbre a escala regional en el cálculo de la evapotranspiración (pérdida de humedad de una superficie por evaporación) potencial.

En (Yugander et al., 2020) mejoran las imágenes ruidosas de resonancia magnética (RM) cerebrales, utilizando el filtrado medio ponderado adaptativo (AWMF) y el filtrado homomórfico.

En (Si et al., 2021) proponen un nuevo método llamado Máquina de Aprendizaje Extremo, basada en la máxima discrepancia de media ponderada para tratar con los problemas de adaptación de dominios no supervisados.

Media Ponderada

Es un tipo de media donde se otorga ponderaciones diferentes a los valores sobre los que se calcula, siendo una de las más utilizadas por no dar la misma importancia a todos los valores.

Es una medida de tendencia central, que es apropiada cuando en un conjunto de datos cada uno de ellos tiene una importancia relativa (o ponderación) respecto de los demás datos.

En (Freund & Simon, 1994) la media ponderada $\bar{x}_w$ de un conjunto de números, $x_1, x_2, x_3, \dots$, y x_n , cuya importancia relativa se expresa numéricamente por medio de un conjunto de números correspondientes, $w_1, w_2, w_3, \dots$, y w_n , se obtiene mediante la fórmula:

$$\bar{x}_w = \frac{x_1 * w_1 + x_2 * w_2 + x_3 * w_3 + \dots + x_n * w_n}{w_1 + w_2 + w_3 + \dots + w_n} = \frac{\sum x * w}{\sum w}$$

Donde, $\sum x * w$ es la suma de los productos obtenidos de la multiplicación de cada x por el valor relativo correspondiente (ponderación) y $\sum w$ es simplemente la suma de los valores relativos. Nótese que cuando todos los valores relativos son iguales, la fórmula de la media ponderada se reduce a la fórmula de la media ordinaria (aritmética).

La media ponderada no es un método de imputación de datos, pero se va a considerar como tal, cuando no se recibe de algún nodo la totalidad de información de control, dado que se pretende ponderar los valores de las variables de los últimos registros históricos de manera diferente según su orden (antigüedad).

Coeficientes de ponderación para registros históricos

Para generar una variación diferente y gradual (no lineal) en los valores de las ponderaciones, se define una fórmula *ad hoc*, que permite generar una escala progresiva utilizando una exponencial, buscando ponderar con mayor énfasis a los valores de los registros más recientes y menor énfasis a los valores de los registros más antiguos.

$$pond = \frac{1 + (\frac{1}{1 + 3^{n-k}} * 100)}{100}$$

Donde n es la cantidad de registro disponible, y k es el número de orden del registro a calcular. Los coeficientes de ponderación considerados en la propuesta son para cuando se disponga de uno a diez registros históricos, como se indica en la Tabla 78, los pesos van creciendo para los registros más recientes y los pesos de menor valor para los registros más antiguos.

Tabla 78. Coeficientes de ponderación para diez registros históricos.

Orden Registro	Ponderaciones									
	Pond. 10	Pond. 9	Pond. 8	Pond. 7	Pond. 6	Pond. 5	Pond. 4	Pond. 3	Pond. 2	Pond. 1
	Reg.	Reg.	Reg.	Reg.	Reg.	Reg.	Reg.	Reg.	Reg.	Reg.
1	0,01	0,01	0,01	0,01	0,01	0,02	0,05	0,13	0,34	1,00
2	0,01	0,01	0,01	0,01	0,02	0,05	0,12	0,30	0,66	-
3	0,01	0,01	0,01	0,02	0,05	0,12	0,28	0,58	-	-
4	0,01	0,01	0,02	0,05	0,11	0,27	0,55	-	-	-
5	0,01	0,02	0,05	0,11	0,27	0,54	-	-	-	-
6	0,02	0,05	0,11	0,27	0,53	-	-	-	-	-
7	0,05	0,11	0,27	0,53	-	-	-	-	-	-
8	0,11	0,26	0,52	-	-	-	-	-	-	-
9	0,26	0,52	-	-	-	-	-	-	-	-
10	0,51	-	-	-	-	-	-	-	-	-
Σ Pond.	1,00	1,00	1,00	1,00	1,00	1,00	1,00	1,00	1,00	1,00

Nota: Elaboración propia.

3.3. Descripción de la capa de imputación/asignación (E2)

La capa de imputación/asignación consta de las siguientes etapas:

3.3.1. Identificar el propósito del problema.

3.3.2. Identificar las alternativas.

3.3.3. Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa.

A continuación, se describirá cada una de las etapas mencionadas.

3.3.1. *Identificar el propósito del problema*

Se presenta una propuesta de un modelo de decisión considerando la imputación de datos

para la gestión de recursos compartidos en sistemas distribuidos, basado en los trabajos mencionados.

El nodo central, es el encargado de recopilar la información de todos los nodos, aplicar el proceso de agregación y obtener la lista de asignaciones de recursos a procesos. La información de control inexistente de alguno de los nodos será resuelta mediante los métodos comentados a continuación.

Escenario E2: para que los procesos accedan a recursos compartidos en la modalidad de exclusión mutua, se hace necesario incorporar mecanismos de imputación/asignación de datos para los casos en que la información de control sobre las variables que indican el estado de carga de alguno de los nodos sea inexistente.

3.3.2. *Identificar alternativas*

Las premisas, las estructuras de datos mencionados en el **Capítulo II** se utilizan como punto de inicio para definir el escenario E2. Para resolver este problema se agregará una capa de imputación/asignación a modelo de decisión que no consideraba información de control faltante.

En un ciclo de recolección de información proporcionada por todos los nodos, el nodo central puede no recibir de algún nodo la totalidad de información de control de los criterios utilizados para el cálculo de la carga computacional de cada nodo y para calcular la prioridad o preferencia de los procesos.

Como en todos los casos el runtime del nodo que arranca estará informando al runtime del nodo central sus características, favoreciendo a los nodos de mejores prestaciones por sobre los demás.

En caso de que la información de control suministrada por los nodos esté completa, el modelo de decisión se ejecuta sin necesidad de la imputación/asignación de datos.

3.3.3. *Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa*

Indicadores de las prestaciones actuales de cada nodo

Para obtener un indicador de las prestaciones actuales de cada nodo se adoptarán las j prestaciones en los n nodos. Los valores que se asumirán para los indicadores de prestaciones de los n nodos son los utilizados en el **Capítulo II**.

Cálculo de la carga computacional actual de los nodos

Para obtener un indicador de la carga computacional actual de cada nodo se adoptarán los e criterios en los n nodos. Los valores que se asumirán para los indicadores de carga computacional de los n nodos y el cálculo de carga promedio para cada nodo son los utilizados del **Capítulo II**.

Puede presentarse la situación de que en cierto ciclo de recolección el nodo central no reciba la totalidad de información de control de alguno de los nodos, utilizados para el cálculo de la carga computacional, entonces se aplica la imputación o asignación de datos, para imputar se tendrán en cuenta ciertas pautas, sabiendo que son las variables de estados consideradas para los nodos.

Imputación o asignación de valores faltantes de los criterios de carga computacional

En función a la información de control histórica (disponible en el nodo central y que corresponde a la información de carga computacional y preferencias de los procesos en ciclos anteriores) disponible para los métodos de imputación/asignación, como se indica en la Figura 6,

se describirán distintas pautas para las diferentes categorías.

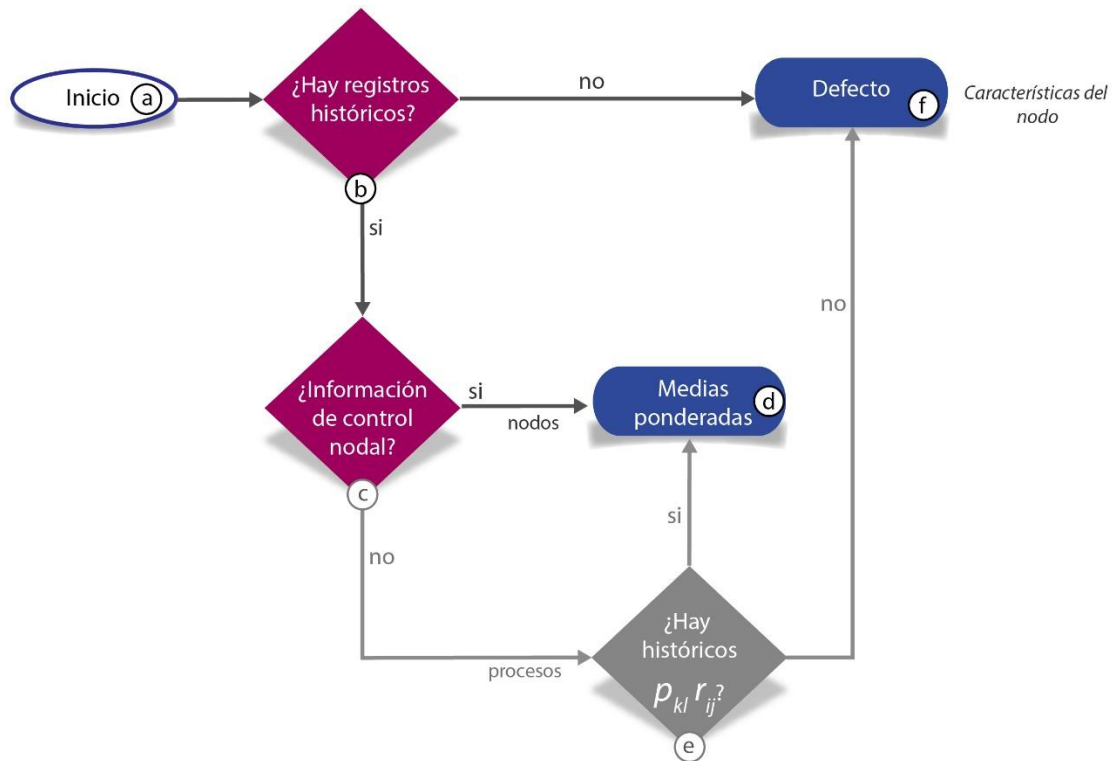


Figura 6. Pautas establecidas para la imputación/asignación de la carga computacional – E2

Fuente: Elaboración propia.

Notas: Asignación por defecto se utiliza para este ejemplo por no tener información sobre otros criterios, asignando un valor en función a la característica del nodo.

La carga nodal se refiere a las variables que indican el estado de carga de los nodos, analizando la Figura 6 se establecen las siguientes pautas para resolver el problema del valor faltante para el escenario E2:

- Primera pauta; para todos los criterios kj del nodo k , **si se cuenta** con información histórica en **(b)**, y la información de control es de los nodos en **(c)**, se puede realizar la imputación con la *Media Ponderada* en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección

de información del nodo k , ponderando con mayor peso a la información más reciente y con menor peso a los más antiguos teniendo en cuenta la Tabla 78, se obtiene entonces, todos los valores faltantes de los criterios kj del nodo k .

- Segunda pauta; para todos los criterios kj del nodo k , **si no se cuenta** con información histórica en **(b)**, se asumirá un valor por defecto en **(f)**, en función a las características del nodo k .
 - Este paso consiste en determinar las prestaciones del nodo k en base a las características conocidas, para cada criterio kj del nodo k se asume el valor correspondiente.

Mediante la imputación o asignación de los valores faltantes en el nodo k se obtienen los valores para los criterios en cuestión, permitiendo de esta manera obtener los indicadores de carga computacional actual de cada nodo como lo indica inicialmente.

El establecimiento de las categorías de carga computacional y de los vectores de pesos asociados a ellos - (E2)

En esta propuesta, se considerarán los mismos criterios, y los mismos vectores de pesos, expuestos en el **Capítulo II**. Teniendo en cuenta que falta la totalidad de información de control, para esta propuesta no se tendrá en cuenta la clasificación de los criterios que se agruparon de acuerdo a los criterios relacionados con las variables de estados principales del nodo, referentes a las prioridades de procesamiento y a los que hacen referencia a las sobrecargas.

Cálculo de las prioridades o preferencias de los procesos teniendo en cuenta el estado del nodo - (E2)

Estas prioridades, al igual que el **Capítulo II**, se calculan en cada nodo para cada solicitud

de recursos originada en cada proceso; el cálculo considera el vector de peso correspondiente según la carga actual del nodo y el vector de los valores concedidos por el nodo según los criterios de evaluación de la solicitud.

Los criterios de evaluación de carga para esta propuesta son las que han sido seleccionadas en el **Capítulo II**, se pueden visualizar en la Tabla 5.

Imputación de las valoraciones faltantes de los criterios de prioridad o preferencia de los procesos – E2

Para los métodos de imputación o asignación se describirán distintas pautas, como se indica en la Figura 7.

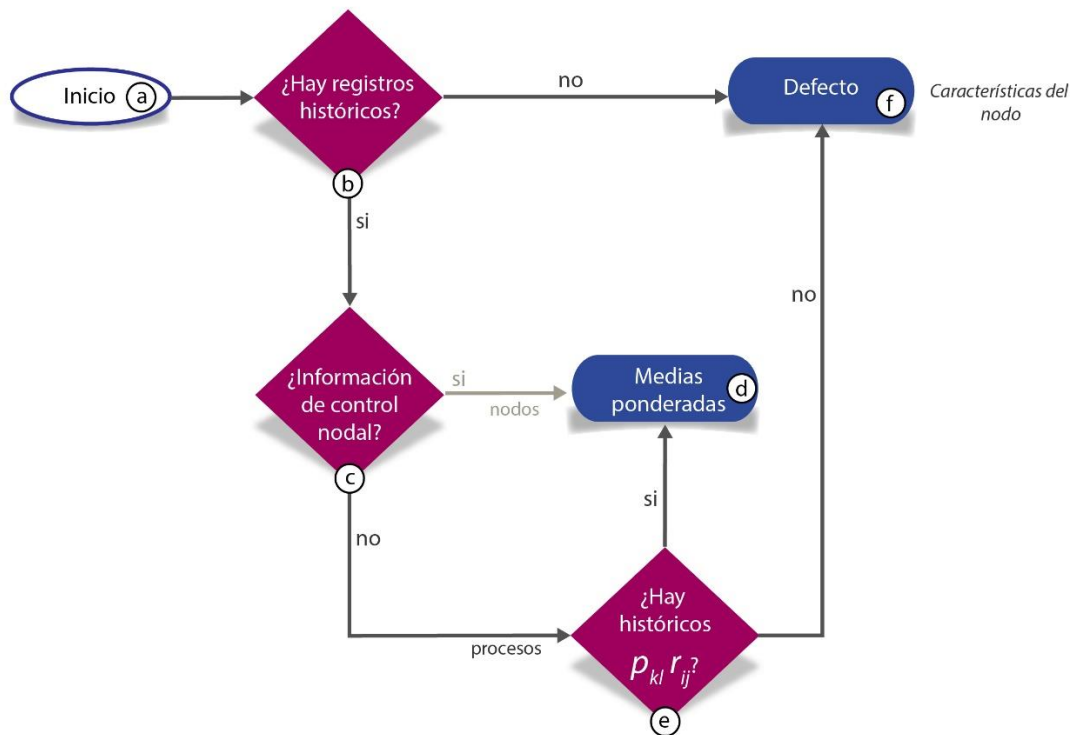


Figura 7. Pautas establecidas para la imputación/asignación de la prioridad o preferencia de los procesos – E2

Fuente: Elaboración propia.

Puede darse la situación de que, en un determinado ciclo de recolección de información de control, el nodo central no reciba la totalidad de información de control sobre los criterios de evaluación de carga de alguno de los nodos, que son los valores que permiten calcular las prioridades o preferencias de los procesos, entonces se aplica el mecanismo de imputación/asignación teniendo en cuenta para las variables dos pautas.

- Primera pauta; para los criterios de la relación recurso-proceso, **si se cuenta** con información histórica en **(b)**, y la información de control corresponde a los procesos en **(c)**, y hay registros históricos de $r_{ij} p_{kl}$ en **(e)**, se aplica la imputación con la *Media Ponderada* en **(d)**. Para ello, se aplica la imputación de datos teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación recurso-proceso $r_{ij} p_{kl}$, ponderando con mayor peso w a la información más reciente y con menor peso w a los más antiguos, se obtiene entonces, los valores faltantes de los criterios.
- Segunda pauta; para los criterios de la relación recurso-proceso, al no tener registros históricos en **(b)**, **no se cuenta** con información histórica específica de $r_{ij} p_{kl}$ en **(e)**, entonces se aplica la asignación de un valor por **defecto** en función a las características del nodo. Este paso consiste en determinar las prestaciones del nodo k en base a las características conocidas, por lo tanto, para cada criterio se asume el valor correspondiente.

Luego de aplicar la imputación o asignación para completar los valores faltantes, se puede tener toda la información necesaria para que el modelo de decisión resuelva los distintos escenarios desarrollados en (La Red Martínez et al., 2017), se prosigue con los cálculos previstos en el mismo.

Los cálculos para las prioridades o preferencias, y los vectores de valoración son los ya expuestos en el punto 2.3.3 del **Capítulo II**, al igual que en la sección anterior, no se consideran

las clasificaciones de los criterios.

El establecimiento de las categorías de carga computacional y de los vectores de pesos asociados a ellos

En esta propuesta serán utilizados los mismos criterios, los mismos vectores de pesos, establecidos el punto 2.3.3 del **Capítulo II**.

Cálculo de las prioridades o preferencias de los procesos para acceder a los recursos compartidos disponibles (se las calcula en el administrador centralizado de recursos compartidos) y determinación del orden en que se asignarán los recursos y a qué proceso será asignado cada recurso

Los cálculos de las prioridades finales, los vectores de pesos finales, los pesos asignados a los procesos, los pesos finales normalizados asignados a los procesos, son los expresados en el punto 2.3.3 del **Capítulo II**, y serán utilizados para realizar los cálculos correspondientes para obtener las obtener **FASD** y posterior **FASDO**.

3.4. Ejemplo

En la presente sección se analizará detalladamente un ejemplo, donde se utilizarán como punto de partida para este ejemplo las premisas y estructura de datos mencionados en el punto 2.4 del **Capítulo II**.

Identificar el propósito del problema

El nodo central se encarga de recibir y mantener actualizada la información de control de todos los nodos, para luego realizar la asignación de recursos en la modalidad de exclusión mutua. El nodo central puede no recibir de alguno de los nodos información de control que permiten

calcular la carga computacional de cada nodo y las prioridades o preferencias de los procesos para cada nodo.

La información de control inexistente de alguno de los nodos, serán resueltos mediante la incorporación de una capa de imputación/asignación, a un modelo de decisión existente para gestión de recursos y procesos en sistemas distribuidos, considerando como punto de inicio la estructura de datos mencionadas en el **Capítulo II**.

Identificar las alternativas

El runtime del nodo central no recibe de algún nodo la totalidad de información de control, utilizados para calcular la carga computacional de cada nodo, y para calcular las prioridades o preferencias de los procesos para cada nodo, se consideran diferentes pautas para completar dichos valores, aplicando mecanismos de imputación o asignación.

Para este ejemplo, para calcular la carga computacional, en uno de los nodos falta la totalidad de información de control, se aplica la imputación de datos si se tiene un histórico de información del nodo en cuestión, de no ser así se hace la asignación de valores por defecto en función a las características del nodo.

Al igual que el caso anterior, se consideran distintos criterios para calcular la prioridad o preferencia de los procesos para cada nodo, la totalidad de los valores de información de control de estos criterios no están disponibles en una o varias relaciones de recurso-proceso. Se aplica el método de imputación o asignación para completar los valores faltantes.

Para todos los casos y al igual que el **Capítulo II**, los runtime de cada nodo informan al runtime del nodo central sus características, y así se obtiene un indicador de las prestaciones de los mismos.

El modelo de decisión se ejecuta si necesidad de la imputación de datos cuando la información suministrada por los nodos esté completa, según la propuesta de (La Red Martínez, 2017) y (Agostini, 2019).

Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa

Indicadores de las prestaciones de cada nodo

Los detalles de implementación para las prestaciones son las utilizadas en el ejemplo (2.4 del **Capítulo II**).

Cálculo de la carga computacional actual de los nodos - (E2)

En un determinado ciclo de recolección de información proporcionada por todos los nodos, el nodo central puede recibir información de control inexistente de alguno de los nodos, que son utilizados para el cálculo de la carga computacional del mismo. Al igual que el **Capítulo II**, para generar un escenario de valores faltantes se procede a realizar la amputación de datos.

Para este ejemplo se considera que la información de control es inexistente en el nodo 1, como se pueden observar en la Tabla 79.

Tabla 79. Valoraciones de los criterios de carga computacional, con valores faltantes en su totalidad en el Nodo 1 - (E2).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	X	X	X
2	45,00	50,00	65,00
3	10,00	25,00	35,00

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Observando que falta la totalidad de valores de criterios en el **Nodo 1**, que a su vez no permite calcular la carga computacional del mismo, entonces se aplica el método de imputación o

asignación para poder completar dichos valores.

Imputación o asignación de las valoraciones faltantes de los criterios de carga computacional (E2)

Teniendo en cuenta las pautas establecidas en **(3.3.3)** para aplicar el método de imputación o asignación y así completar los valores faltantes, se ha verificado la disponibilidad de información de control histórica sobre los criterios de carga computacional perteneciente al nodo 1, se procede a imputar los valores teniendo en cuenta la pauta 1.

- **Nodo 1:** para este caso se cuenta con información histórica, esto permite aplicar la imputación con la *Media Ponderada*, como lo establece la pauta mencionada. Para este ejemplo, los registros de información histórica a utilizar son de los diez últimos ciclos de recolección, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros en la Tabla 80, teniendo en cuenta que la última fila corresponde a la información más reciente.

Tabla 80. Valores históricos de los criterios de carga computacional Nodo 1, últimos cinco registros de diez - (E2).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	65,00	36,00	63,00
1	43,00	38,00	12,00
1	40,00	82,00	12,00
1	52,00	23,00	23,00
1	87,00	24,00	63,00

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con la Media Ponderada, utilizando un histórico de valores correspondientes a los diez últimos ciclos de recolección de información del Nodo 1, se obtienen las valoraciones para el mismo, de los criterios **% Uso de CPU**, **% Uso de Memoria** y **% Uso de Oper. E/S**, como se puede ver en la Tabla 81.

Tabla 81. Valores de los criterios de carga computacional con valores imputados del Nodo 1, aplicando la pauta una - (E2).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	80,00 / 68,47	90,00 / 32,19	75,00 / 43,05
2	45,00	50,00	65,00
3	10,00	25,00	35,00

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

En la celda remarcada de la Tabla 81, los valores de la izquierda corresponden a los valores obtenidos en (La Red Martínez, 2017), y los de la derecha son resultados de la imputación con la *Media Ponderada*.

Si para este ejemplo, no se contara con información histórica, se aplicaría la segunda pauta asumiendo un valor por defecto en función a las características del nodo, determinando las prestaciones del nodo en base a las características conocidas en la Tabla 18.

Tabla **18**, del **Capítulo II**, considerando que al no tener información histórica se está en presencia de un nodo recientemente creado. Se determina que el nodo 1 es de prestaciones baja, el valor para asignar a cada criterio será de 30%, obteniendo las valoraciones como se puede observar en la Tabla 82.

Tabla 82. Valores de los criterios de carga computacional con valores asignado según las prestaciones del Nodo 1, aplicando la pauta 2 - (E2).

Nodos	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S
1	30,00	30,00	30,00
2	45,00	50,00	65,00
3	10,00	25,00	35,00

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Teniendo en cuenta las valoraciones obtenidas en la Tabla 81, posterior a la imputación de

datos, los valores que se asumirán para los indicadores de carga computacional de los tres nodos, y el promedio para cada nodo, puede verse en la Tabla 83.

Tabla 83. Valores de los criterios de carga computacional en cada nodo, y el cálculo del promedio de cada nodo, con valores imputados en el Nodo 1, con la pauta una - (E2).

Nodos	Valores de los Criterios			
	% Uso de CPU	% Uso de Memoria	% Uso de Oper. E/S	Promedio
1	68,47	32,19	43,05	47,90
2	45,00	50,00	65,00	54,16
3	10,00	25,00	35,00	34,82

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Las categorías de carga computacional actual de cada nodo serán:

- Alta (si la carga es mayor al 70,00%).
- Los **Nodo 1** y **Nodo 2** corresponden a categoría Media (la carga está entre el 40,00% y el 70,00%).
- El **Nodo 3** corresponde a categoría Baja (la carga es menor al 40,00%).

Los criterios para determinar la prioridad o preferencia, los valores correspondientes a los criterios y los vectores de pesos para las distintas categorías de carga, son los expuestos en el **Capítulo II**.

Cálculo de las prioridades o preferencias de los procesos teniendo en cuenta el estado del nodo

En el mismo ciclo de recolección de información de control del caso anterior, puede ser inexistente; es decir, faltar totalmente la información de control de los valores en alguna o algunas relaciones de proceso-recurso, utilizados para el cálculo de la prioridad o preferencias de los

procesos para el nodo en cuestión. El escenario propuesto en este ejemplo resulta del escenario planteado en la Tabla 29 del **Capítulo II**, que para este caso se considera que no se tiene la totalidad de información de control de los criterios en un registro de relaciones de proceso-recurso, que se puede observar en la Tabla 84.

Tabla 84. Valores asignados a los criterios para calcular la prioridad nodal, faltando la totalidad de información de control en algunas relaciones de proceso-recurso - (E2).

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{11}$	0,70	0,50	0,70	0,90	0,80	0,20	0,30	0,40
$p_{11}r_{12}$	0,80	0,70	0,40	0,50	0,30	0,70	0,20	0,40
$p_{11}r_{21}$	0,30	0,40	0,50	0,20	0,90	0,20	0,50	0,70
$p_{11}r_{22}$	0,50	0,50	0,70	0,40	0,80	0,30	0,50	0,60
$p_{11}r_{23}$	X	X	X	X	X	X	X	X
$p_{11}r_{24}$	0,30	0,50	0,90	0,20	0,60	0,60	0,70	0,40
$p_{12}r_{11}$	0,40	0,70	0,50	0,90	1,00	0,90	0,80	0,80
$p_{12}r_{12}$	0,20	0,70	0,30	0,70	0,80	0,30	0,80	0,90
$p_{12}r_{21}$	X	X	X	X	X	X	X	X
$p_{12}r_{22}$	0,90	0,60	0,70	0,70	0,80	0,20	0,50	0,40
$p_{12}r_{31}$	0,20	0,50	0,70	0,70	0,30	0,20	0,70	0,80
$p_{12}r_{33}$	0,40	0,50	0,70	0,90	0,30	0,40	0,50	0,80
$p_{13}r_{11}$	X	X	X	X	X	X	X	X
$p_{13}r_{12}$	0,70	0,80	0,70	0,40	0,90	0,20	0,90	0,70
$p_{13}r_{13}$	0,70	0,60	0,70	0,80	0,90	0,40	0,80	0,70
$p_{13}r_{21}$	0,70	0,40	0,90	0,30	0,50	0,70	0,20	0,30
$p_{13}r_{22}$	0,50	0,90	0,80	0,30	0,50	0,40	0,90	0,30
$p_{13}r_{31}$	0,50	0,70	0,30	0,60	0,80	0,90	0,90	0,50
$p_{13}r_{32}$	0,60	0,90	0,30	0,60	0,40	0,80	0,70	0,80
$p_{13}r_{33}$	0,60	0,20	0,40	0,60	0,90	0,80	0,50	0,80
$p_{21}r_{12}$	0,20	0,10	0,30	0,80	0,70	0,70	0,50	0,80
$p_{21}r_{13}$	0,20	0,40	0,80	0,90	0,70	0,40	0,50	0,20
$p_{21}r_{22}$	0,30	0,50	0,80	0,90	0,50	0,60	0,20	0,40
$p_{21}r_{23}$	0,70	0,50	0,60	0,80	0,50	0,20	0,90	0,40
$p_{21}r_{31}$	0,70	0,50	0,80	0,60	0,90	0,40	0,90	0,90
$p_{21}r_{33}$	0,70	0,40	0,90	0,80	0,40	0,80	0,60	0,60
$p_{22}r_{21}$	0,50	0,40	0,70	0,80	0,60	0,40	0,60	0,90
$p_{22}r_{22}$	X	X	X	X	X	X	X	X
$p_{22}r_{31}$	0,50	0,40	0,50	0,20	0,40	0,80	0,20	0,90
$p_{22}r_{33}$	0,80	0,40	0,30	0,20	0,80	0,90	0,50	0,80
$p_{23}r_{12}$	0,50	0,60	0,80	0,30	0,50	0,70	0,50	0,50
$p_{23}r_{24}$	0,60	0,20	0,10	0,70	0,30	0,80	0,90	0,40
$p_{23}r_{31}$	0,30	0,10	0,40	0,80	0,70	0,90	0,40	0,60

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p</i> ₂₃ <i>r</i> ₃₂	X	X	X	X	X	X	X	X
<i>p</i> ₂₃ <i>r</i> ₃₃	0,30	0,60	0,80	0,20	0,90	0,60	0,40	0,20
<i>p</i> ₂₄ <i>r</i> ₁₁	0,40	0,60	0,70	0,90	0,50	0,70	0,90	0,50
<i>p</i> ₂₄ <i>r</i> ₁₂	X	X	X	X	X	X	X	X
<i>p</i> ₂₄ <i>r</i> ₂₃	0,40	0,90	0,40	0,80	0,80	0,70	0,60	0,30
<i>p</i> ₂₄ <i>r</i> ₂₄	0,50	0,80	0,70	0,90	0,30	0,40	0,80	0,70
<i>p</i> ₂₅ <i>r</i> ₂₁	X	X	X	X	X	X	X	X
<i>p</i> ₃₁ <i>r</i> ₁₃	0,60	0,90	0,60	0,90	0,70	0,40	0,90	0,80
<i>p</i> ₃₁ <i>r</i> ₃₁	0,80	0,30	0,90	0,50	0,70	0,30	0,90	0,70
<i>p</i> ₃₁ <i>r</i> ₃₃	0,40	0,70	0,70	0,50	0,90	0,90	0,70	0,70
<i>p</i> ₃₂ <i>r</i> ₁₁	0,60	0,90	0,80	0,50	0,90	0,70	0,30	0,70
<i>p</i> ₃₂ <i>r</i> ₁₂	0,70	0,40	0,60	0,90	0,80	1,00	0,90	0,70
<i>p</i> ₃₂ <i>r</i> ₁₃	0,80	0,90	1,00	0,70	0,90	0,30	0,50	0,90
<i>p</i> ₃₂ <i>r</i> ₂₃	0,80	0,80	1,00	0,90	0,60	0,80	0,40	0,90
<i>p</i> ₃₃ <i>r</i> ₁₁	0,20	0,70	0,90	0,80	0,60	0,90	1,00	0,30
<i>p</i> ₃₃ <i>r</i> ₁₂	0,90	1,00	0,90	0,70	0,30	0,50	0,70	0,30
<i>p</i> ₃₃ <i>r</i> ₁₃	X	X	X	X	X	X	X	X
<i>p</i> ₃₃ <i>r</i> ₂₁	0,40	0,60	0,70	0,80	0,80	0,40	0,80	0,80
<i>p</i> ₃₃ <i>r</i> ₂₂	X	X	X	X	X	X	X	X
<i>p</i> ₃₃ <i>r</i> ₂₃	0,40	0,60	0,90	1,00	0,60	0,40	0,70	0,80
<i>p</i> ₃₃ <i>r</i> ₃₂	0,60	0,60	0,70	0,80	0,40	0,50	0,80	0,80
<i>p</i> ₃₃ <i>r</i> ₃₃	0,60	1,00	0,70	0,40	0,60	0,80	0,80	0,90
<i>p</i> ₃₄ <i>r</i> ₁₂	0,80	1,00	0,30	0,50	0,70	0,30	0,80	0,70
<i>p</i> ₃₄ <i>r</i> ₁₃	0,20	1,00	0,80	0,50	0,80	0,60	0,90	0,70
<i>p</i> ₃₄ <i>r</i> ₂₂	0,90	0,80	0,60	0,80	0,90	0,80	0,50	0,60
<i>p</i> ₃₄ <i>r</i> ₂₃	0,40	0,60	0,70	0,80	0,90	0,50	0,90	0,60
<i>p</i> ₃₄ <i>r</i> ₂₄	0,90	0,40	0,70	0,40	0,70	0,50	0,90	0,70
<i>p</i> ₃₄ <i>r</i> ₃₁	0,50	0,60	0,70	0,90	0,70	0,90	0,60	0,70
<i>p</i> ₃₄ <i>r</i> ₃₂	0,80	0,60	0,90	0,50	0,60	0,90	0,30	0,90
<i>p</i> ₃₄ <i>r</i> ₃₃	X	X	X	X	X	X	X	X
<i>p</i> ₃₅ <i>r</i> ₁₂	X	X	X	X	X	X	X	X
<i>p</i> ₃₅ <i>r</i> ₁₃	0,80	0,90	0,70	0,30	0,90	0,80	0,70	0,40
<i>p</i> ₃₅ <i>r</i> ₂₂	0,30	0,80	0,90	0,40	0,80	0,60	0,30	0,60
<i>p</i> ₃₅ <i>r</i> ₂₄	X	X	X	X	X	X	X	X
<i>p</i> ₃₅ <i>r</i> ₃₁	0,90	0,80	0,70	0,40	0,80	0,30	0,40	0,80
<i>p</i> ₃₅ <i>r</i> ₃₂	0,90	0,70	0,80	0,60	0,40	0,90	0,40	0,70
<i>p</i> ₃₅ <i>r</i> ₃₃	0,50	0,70	0,60	0,90	0,80	0,50	0,90	0,70
<i>p</i> ₃₆ <i>r</i> ₁₁	0,80	0,90	0,60	0,50	0,80	0,90	0,90	0,60
<i>p</i> ₃₆ <i>r</i> ₁₂	0,70	0,90	0,40	0,90	0,80	0,90	0,40	0,70
<i>p</i> ₃₆ <i>r</i> ₁₃	0,90	0,90	0,80	0,70	0,80	0,70	0,90	0,70
<i>p</i> ₃₆ <i>r</i> ₂₁	0,80	0,90	0,50	0,30	0,70	0,70	0,90	0,90
<i>p</i> ₃₆ <i>r</i> ₂₂	0,80	0,20	0,10	0,30	0,80	0,70	0,60	0,90
<i>p</i> ₃₆ <i>r</i> ₂₄	0,40	0,70	0,90	0,30	0,80	0,50	0,80	0,90
<i>p</i> ₃₆ <i>r</i> ₃₁	0,50	0,90	0,90	0,70	0,80	0,80	0,80	0,70
<i>p</i> ₃₆ <i>r</i> ₃₂	0,90	0,80	0,60	0,70	0,80	0,50	0,60	0,70

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{36r33}	0,20	0,70	0,40	0,80	0,80	0,90	0,40	0,50
p_{37r11}	0,90	0,60	0,80	0,60	0,90	0,70	0,90	0,80
p_{37r12}	0,70	0,90	0,80	0,70	0,50	0,80	0,80	0,60
p_{37r21}	0,90	0,70	0,80	0,60	0,60	0,90	0,50	0,40
p_{37r32}	0,80	0,90	0,50	0,30	0,80	0,70	0,40	0,60
p_{37r33}	0,80	0,40	0,60	0,80	0,80	0,60	0,90	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Imputación de las valoraciones de criterios donde se presenta valores faltantes en su totalidad (E2)

Se procede a imputar o asignar dichos valores para completar los datos faltantes, teniendo en cuenta las pautas definidas anteriormente.

Cuando se asigne un valor por defecto se tendrán en cuenta los siguientes detalles de implementación:

- Nodo de prestaciones altas se asume un valor de 0,30.
- Nodo de prestaciones medias se asume un valor de 0,50.
- Nodo de prestaciones bajas se asume un valor 0,70.

Después de haber analizado la información con que se cuenta sobre las variables de estados principales de los nodos donde existen valores faltantes, y teniendo en cuenta las pautas para la imputación o asignación de valores, la relación de proceso-recurso queda clasificada como se muestra en la Tabla 85.

Al tener información histórica de la misma relación proceso-recurso, se aplica la imputación teniendo en cuenta la pauta 1, únicamente de no tener dicha información se puede aplicar la pauta 2.

Tabla 85. Clasificación de valores faltantes de acuerdo a las pautas establecidas, para calcular la prioridad o preferencia de los procesos - E2.

Valor Faltante	Pauta 1: Existe información histórica misma relación proceso-recurso		Pauta 2: No existe información histórica.
	Cant. registros $p_{kl} r_{ij}$	Aplica	
$p_{11}r_{23}$	10	SI	---
$p_{12}r_{21}$	6	SI	---
$p_{13}r_{11}$	9	SI	---
$p_{22}r_{22}$	3	SI	---
$p_{23}r_{32}$	8	SI	---
$p_{24}r_{12}$	10	SI	---
$p_{25}r_{21}$	2	SI	---
$p_{33}r_{13}$	7	SI	---
$p_{33}r_{22}$	5	SI	---
$p_{34}r_{33}$	4	SI	---
$p_{35}r_{12}$	1	SI	---
$p_{35}r_{24}$	0	NO	SI

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

- $p_{11}r_{23}$: como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con diez registros, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 86.

Tabla 86. Valores históricos de la relación proceso-recurso $p_{11}r_{23}$, últimos cinco registros de diez (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{23}$	0,20	0,70	0,50	0,20	1,00	1,00	0,50	0,20
$p_{11}r_{23}$	0,60	0,50	0,20	0,40	0,30	0,70	0,50	0,80
$p_{11}r_{23}$	0,60	0,40	0,30	0,70	0,70	0,50	1,00	0,70
$p_{11}r_{23}$	0,60	0,60	0,60	0,20	0,80	0,50	0,60	0,40

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a diez ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{11}r_{23}$, como se puede ver en la Tabla 87.

Tabla 87. Valores imputados con la media ponderada en la relación proceso-recurso p_{11r23} , aplicando la pauta uno (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{11r23}	0,57	0,48	0,39	0,80	0,58	0,66	0,76	0,43

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- p_{12r21} : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *seis registros*, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 88.

Tabla 88. Valores históricos de la relación proceso-recurso p_{12r21} , últimos cinco registros de seis - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{12r21}	0,60	0,60	0,40	0,50	0,70	0,40	0,40	0,10
p_{12r21}	0,20	0,80	0,10	0,50	1,00	0,80	0,20	0,70
p_{12r21}	0,90	0,40	0,10	0,80	0,80	0,80	0,50	0,50
p_{12r21}	0,80	0,70	0,60	0,70	0,40	1,00	1,00	0,40
p_{12r21}	0,30	0,10	0,40	0,40	0,90	0,90	0,80	0,50

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a seis ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de p_{12r21} , como se puede ver en la Tabla 89.

Tabla 89. Valores imputados con la media ponderada en la relación proceso-recurso p_{12r21} , aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{12r21}	0,50	0,35	0,41	0,53	0,75	0,89	0,78	0,47

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- **p_{13r11}** : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *nueve registros*, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 90.

Tabla 90. Valores históricos de la relación proceso-recurso **p_{13r11}** , últimos cinco registros de nueve - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{13r11}	0,50	0,30	0,20	0,70	0,40	0,70	0,90	0,70
p_{13r11}	0,30	0,80	0,70	0,20	0,60	0,90	0,30	0,60
p_{13r11}	0,70	0,20	0,30	0,80	0,60	0,70	0,70	0,30
p_{13r11}	0,80	0,40	0,20	0,20	0,30	0,50	0,30	0,20
p_{13r11}	0,60	0,60	0,20	0,50	0,50	0,80	0,40	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a seis ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p_{13r11}** , como se puede ver en la Tabla 91.

Tabla 91. Valores imputados con la media ponderada en la relación proceso-recurso **p_{13r11}** , aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{13r11}	0,65	0,50	0,24	0,43	0,46	0,70	0,41	0,25

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- **p_{22r22}** : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *tres registros*, las valoraciones se pueden apreciar en la Tabla 92, siendo el de la última fila el registro más reciente.

Tabla 92. Valores históricos de la relación proceso-recurso $p_{22}r_{22}$, tres únicos registros - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{22}r_{22}$	0,60	0,80	0,80	0,70	0,80	0,20	0,50	0,20
$p_{22}r_{22}$	0,90	0,40	0,70	0,90	0,50	1,00	0,80	0,60
$p_{22}r_{22}$	0,60	1,00	0,80	0,20	0,40	1,00	0,30	0,40

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a tres ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{22}r_{22}$, como se puede ver en la Tabla 93.

Tabla 93. Valores imputados con la media ponderada en la relación proceso-recurso $p_{22}r_{22}$, aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{22}r_{22}$	0,69	0,80	0,77	0,47	0,48	0,90	0,47	0,43

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- $p_{23}r_{32}$: como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *ocho registros*, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 94.

Tabla 94. Valores históricos de la relación proceso-recurso $p_{23}r_{32}$, últimos cinco registros de ocho - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{23}r_{32}$	0,90	0,40	0,50	0,40	0,40	0,20	0,50	0,10
$p_{23}r_{32}$	0,60	0,70	0,70	0,20	0,70	0,90	0,20	0,20
$p_{23}r_{32}$	0,80	0,50	0,80	1,00	1,00	0,20	0,50	0,30
$p_{23}r_{32}$	0,50	0,20	0,60	0,60	0,70	1,00	0,80	0,60
$p_{23}r_{32}$	0,40	0,90	1,00	0,40	0,30	0,30	1,00	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a ocho ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p_{23r32}**, como se puede ver en la Tabla 95.

Tabla 95. Valores imputados con la media ponderada en la relación proceso-recurso **p_{23r32}**, aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{23r32}	0,49	0,64	0,83	0,51	0,51	0,51	0,82	0,32

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- **p_{24r12}**: como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con diez registros, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 96.

Tabla 96. Valores históricos de la relación proceso-recurso **p_{24r12}**, últimos cinco registros de diez - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	0,90	0,60	1,00	1,00	1,00	0,70	0,90	0,90
p_{24r12}	0,30	0,30	0,30	0,80	0,30	0,90	0,60	0,80
p_{24r12}	0,80	0,30	0,30	0,90	0,30	0,40	1,00	0,90
p_{24r12}	0,90	0,40	0,80	1,00	0,30	0,40	0,40	0,80
p_{24r12}	0,60	0,90	0,50	0,30	0,30	0,70	0,80	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Aplicando la imputación de datos con un histórico de valores correspondientes a diez ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p_{24r12}**, como se puede ver en la Tabla 97.

Tabla 97. Valores imputados con la media ponderada en la relación proceso-recurso p_{24r12} , aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	0,69	0,64	0,57	0,61	0,33	0,58	0,70	0,55

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- p_{25r21} : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *dos registros*, las valoraciones se pueden apreciar en la Tabla 98, siendo el de la última fila el registro más reciente.

Tabla 98. Valores históricos de la relación proceso-recurso p_{25r21} , dos únicos registros - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{25r21}	0,60	0,80	0,10	0,80	0,30	0,80	0,30	0,40
p_{25r21}	0,40	1,00	0,70	0,90	0,30	0,50	0,90	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Aplicando la imputación de datos con un histórico de valores correspondientes a dos ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de p_{25r21} , como se puede ver en la Tabla 99.

Tabla 99. Valores imputados con la media ponderada en la relación proceso-recurso p_{25r21} , aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{25r21}	0,47	0,93	0,50	0,87	0,30	0,60	0,70	0,60

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- p_{33r13} : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada

como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *siete registros*, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 100.

Tabla 100. Valores históricos de la relación proceso-recurso *p33r13*, los cinco últimos registros de siete - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p33r13</i>	0,20	0,60	0,80	0,50	0,80	0,90	0,70	0,90
<i>p33r13</i>	0,60	1,00	0,10	0,20	0,80	0,60	0,30	0,60
<i>p33r13</i>	0,80	0,60	0,10	0,40	0,50	0,30	1,00	0,20
<i>p33r13</i>	0,70	0,50	1,00	0,60	1,00	0,90	1,00	0,40
<i>p33r13</i>	0,30	0,60	0,60	1,00	0,80	0,60	0,80	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Aplicando la imputación de datos con un histórico de valores correspondientes a siete ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de *p33r13*, como se puede ver en Tabla 101.

Tabla 101. Valores imputados con la media ponderada en la relación proceso-recurso *p33r13*, aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p33r13</i>	0,48	0,60	0,62	0,76	0,81	0,65	0,84	0,34

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- *p33r22*: como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *cinco registros*, se puede apreciar en la Tabla 102 las valoraciones de los cinco registros, siendo el de la última fila el registro más reciente.

Tabla 102. Valores históricos de la relación proceso-recurso p_{33r22} , los cinco únicos registros - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r22}	0,80	0,40	0,50	0,70	0,60	0,90	0,90	0,90
p_{33r22}	0,30	0,90	0,90	1,00	0,90	0,40	0,40	0,50
p_{33r22}	0,50	0,30	0,60	0,30	0,50	1,00	0,60	0,20
p_{33r22}	0,80	0,10	0,90	0,40	0,30	0,60	0,60	0,50
p_{33r22}	0,20	0,20	0,90	0,60	0,40	0,20	0,90	0,60

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Aplicando la imputación de datos con un histórico de valores correspondientes a cinco ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de p_{33r22} , como se puede ver en la Tabla 103.

Tabla 103. Valores imputados con la media ponderada en la relación proceso-recurso p_{33r22} , aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{33r22}	0,42	0,23	0,86	0,53	0,41	0,43	0,76	0,53

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- p_{34r33} : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *cuatro registros*, se puede apreciar en la Tabla 104 las valoraciones de los cuatros registros, siendo el de la última fila el registro más reciente.

Tabla 104. Valores históricos de la relación proceso-recurso p_{34r33} , los cuatro únicos registros - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{34r33}	0,30	0,80	0,80	0,40	1,00	0,40	0,30	0,50
p_{34r33}	0,90	0,70	0,40	0,70	0,60	0,70	0,50	0,90
p_{34r33}	0,90	0,20	0,30	1,00	0,40	0,60	0,50	0,20
p_{34r33}	0,50	0,40	0,90	0,60	0,50	0,90	1,00	0,60

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Aplicando la imputación de datos con un histórico de valores correspondientes a cuatro ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p34r33**, como se puede ver en la Tabla 105.

Tabla 105. Valores imputados con la media ponderada en la relación proceso-recurso **p34r33**, aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p34r33	0,79	0,67	0,58	0,54	0,66	0,91	0,86	0,40

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- **p35r12**: como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *un registro*, se puede apreciar en la Tabla 106 las valoraciones del único registro.

Tabla 106. Valores históricos de la relación proceso-recurso **p35r12**, registro único (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p35r12	0,60	0,20	0,90	0,60	0,40	0,80	0,90	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019)

Aplicando la imputación de datos con un histórico de valores correspondientes a un ciclo de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p35r12**, como se puede ver en la Tabla 107.

Tabla 107. Valores imputados con la media ponderada en la relación proceso-recurso **p35r12**, aplicando la pauta uno - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p35r12	0,60	0,20	0,90	0,60	0,40	0,80	0,90	0,20

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

- **$p_{35}r_{24}$** : como se muestra en la clasificación de la Tabla 85, para esta relación de proceso-recurso no se cuenta con información histórica, por ello se aplica la segunda pauta asumiendo un valor por defecto en función a las características del nodo, determinando las prestaciones del nodo en base a las características conocidas en la Tabla 18 del **Capítulo II**. Se determina que el nodo 3 es de prestaciones Alta, el valor **0,70** es asignado a cada criterio, como se puede observar en la Tabla 108.

Tabla 108. Valores asignados según las prestaciones del nodo para los criterios de **$p_{35}r_{24}$** , aplicando la pauta dos - (E2).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{35}r_{24}$	0,70	0,70	0,70	0,70	0,70	0,70	0,70	0,70

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Se obtienen los valores de los criterios mediante la imputación o asignación, como se observa en la Tabla 109, que serán utilizados para calcular la prioridad o preferencia de procesos para cada nodo.

Tabla 109. Valoraciones asignadas a los criterios para calcular la prioridad nodal, faltando la totalidad de información de control en algunas relaciones de proceso-recurso - (E2).

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{11}$	0,70	0,50	0,70	0,90	0,80	0,20	0,30	0,40
$p_{11}r_{12}$	0,80	0,70	0,40	0,50	0,30	0,70	0,20	0,40
$p_{11}r_{21}$	0,30	0,40	0,50	0,20	0,90	0,20	0,50	0,70
$p_{11}r_{22}$	0,50	0,50	0,70	0,40	0,80	0,30	0,50	0,60
$p_{11}r_{23}$	0,57	0,49	0,39	0,80	0,58	0,67	0,77	0,43
$p_{11}r_{24}$	0,30	0,50	0,90	0,20	0,60	0,60	0,70	0,40
$p_{12}r_{11}$	0,40	0,70	0,50	0,90	1,00	0,90	0,80	0,80
$p_{12}r_{12}$	0,20	0,70	0,30	0,70	0,80	0,30	0,80	0,90
$p_{12}r_{21}$	0,51	0,35	0,41	0,53	0,75	0,90	0,78	0,48
$p_{12}r_{22}$	0,90	0,60	0,70	0,70	0,80	0,20	0,50	0,40
$p_{12}r_{31}$	0,20	0,50	0,70	0,70	0,30	0,20	0,70	0,80
$p_{12}r_{33}$	0,40	0,50	0,70	0,90	0,30	0,40	0,50	0,80

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p_{13r11}</i>	0,65	0,51	0,24	0,44	0,46	0,71	0,42	0,25
<i>p_{13r12}</i>	0,70	0,80	0,70	0,40	0,90	0,20	0,90	0,70
<i>p_{13r13}</i>	0,70	0,60	0,70	0,80	0,90	0,40	0,80	0,70
<i>p_{13r21}</i>	0,70	0,40	0,90	0,30	0,50	0,70	0,20	0,30
<i>p_{13r22}</i>	0,50	0,90	0,80	0,30	0,50	0,40	0,90	0,30
<i>p_{13r31}</i>	0,50	0,70	0,30	0,60	0,80	0,90	0,90	0,50
<i>p_{13r32}</i>	0,60	0,90	0,30	0,60	0,40	0,80	0,70	0,80
<i>p_{13r33}</i>	0,60	0,20	0,40	0,60	0,90	0,80	0,50	0,80
<i>p_{21r12}</i>	0,20	0,10	0,30	0,80	0,70	0,70	0,50	0,80
<i>p_{21r13}</i>	0,20	0,40	0,80	0,90	0,70	0,40	0,50	0,20
<i>p_{21r22}</i>	0,30	0,50	0,80	0,90	0,50	0,60	0,20	0,40
<i>p_{21r23}</i>	0,70	0,50	0,60	0,80	0,50	0,20	0,90	0,40
<i>p_{21r31}</i>	0,70	0,50	0,80	0,60	0,90	0,40	0,90	0,90
<i>p_{21r33}</i>	0,70	0,40	0,90	0,80	0,40	0,80	0,60	0,60
<i>p_{22r21}</i>	0,50	0,40	0,70	0,80	0,60	0,40	0,60	0,90
<i>p_{22r22}</i>	0,69	0,80	0,77	0,47	0,48	0,90	0,47	0,43
<i>p_{22r31}</i>	0,50	0,40	0,50	0,20	0,40	0,80	0,20	0,90
<i>p_{22r33}</i>	0,80	0,40	0,30	0,20	0,80	0,90	0,50	0,80
<i>p_{23r12}</i>	0,50	0,60	0,80	0,30	0,50	0,70	0,50	0,50
<i>p_{23r24}</i>	0,60	0,20	0,10	0,70	0,30	0,80	0,90	0,40
<i>p_{23r31}</i>	0,30	0,10	0,40	0,80	0,70	0,90	0,40	0,60
<i>p_{23r32}</i>	0,49	0,64	0,83	0,51	0,51	0,51	0,82	0,32
<i>p_{23r33}</i>	0,30	0,60	0,80	0,20	0,90	0,60	0,40	0,20
<i>p_{24r11}</i>	0,40	0,60	0,70	0,90	0,50	0,70	0,90	0,50
<i>p_{24r12}</i>	0,69	0,64	0,57	0,61	0,33	0,58	0,70	0,55
<i>p_{24r23}</i>	0,40	0,90	0,40	0,80	0,80	0,70	0,60	0,30
<i>p_{24r24}</i>	0,50	0,80	0,70	0,90	0,30	0,40	0,80	0,70
<i>p_{25r21}</i>	0,47	0,93	0,50	0,87	0,30	0,60	0,70	0,60
<i>p_{31r13}</i>	0,60	0,90	0,60	0,90	0,70	0,40	0,90	0,80
<i>p_{31r31}</i>	0,80	0,30	0,90	0,50	0,70	0,30	0,90	0,70
<i>p_{31r33}</i>	0,40	0,70	0,70	0,50	0,90	0,90	0,70	0,70
<i>p_{32r11}</i>	0,60	0,90	0,80	0,50	0,90	0,70	0,30	0,70
<i>p_{32r12}</i>	0,70	0,40	0,60	0,90	0,80	1,00	0,90	0,70
<i>p_{32r13}</i>	0,80	0,90	1,00	0,70	0,90	0,30	0,50	0,90
<i>p_{32r23}</i>	0,80	0,80	1,00	0,90	0,60	0,80	0,40	0,90
<i>p_{33r11}</i>	0,20	0,70	0,90	0,80	0,60	0,90	1,00	0,30
<i>p_{33r12}</i>	0,90	1,00	0,90	0,70	0,30	0,50	0,70	0,30
<i>p_{33r13}</i>	0,48	0,60	0,62	0,76	0,81	0,65	0,84	0,34
<i>p_{33r21}</i>	0,40	0,60	0,70	0,80	0,80	0,40	0,80	0,80
<i>p_{33r22}</i>	0,42	0,23	0,86	0,53	0,41	0,43	0,76	0,53
<i>p_{33r23}</i>	0,40	0,60	0,90	1,00	0,60	0,40	0,70	0,80
<i>p_{33r32}</i>	0,60	0,60	0,70	0,80	0,40	0,50	0,80	0,80
<i>p_{33r33}</i>	0,60	1,00	0,70	0,40	0,60	0,80	0,80	0,90
<i>p_{34r12}</i>	0,80	1,00	0,30	0,50	0,70	0,30	0,80	0,70
<i>p_{34r13}</i>	0,20	1,00	0,80	0,50	0,80	0,60	0,90	0,70

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p_{34r22}</i>	0,90	0,80	0,60	0,80	0,90	0,80	0,50	0,60
<i>p_{34r23}</i>	0,40	0,60	0,70	0,80	0,90	0,50	0,90	0,60
<i>p_{34r24}</i>	0,90	0,40	0,70	0,40	0,70	0,50	0,90	0,70
<i>p_{34r31}</i>	0,50	0,60	0,70	0,90	0,70	0,90	0,60	0,70
<i>p_{34r32}</i>	0,80	0,60	0,90	0,50	0,60	0,90	0,30	0,90
<i>p_{34r33}</i>	0,79	0,67	0,58	0,54	0,66	0,91	0,86	0,40
<i>p_{35r12}</i>	0,60	0,20	0,90	0,60	0,40	0,80	0,90	0,20
<i>p_{35r13}</i>	0,80	0,90	0,70	0,30	0,90	0,80	0,70	0,40
<i>p_{35r22}</i>	0,30	0,80	0,90	0,40	0,80	0,60	0,30	0,60
<i>p_{35r24}</i>	0,70	0,70	0,70	0,70	0,70	0,70	0,70	0,70
<i>p_{35r31}</i>	0,90	0,80	0,70	0,40	0,80	0,30	0,40	0,80
<i>p_{35r32}</i>	0,90	0,70	0,80	0,60	0,40	0,90	0,40	0,70
<i>p_{35r33}</i>	0,50	0,70	0,60	0,90	0,80	0,50	0,90	0,70
<i>p_{36r11}</i>	0,80	0,90	0,60	0,50	0,80	0,90	0,90	0,60
<i>p_{36r12}</i>	0,70	0,90	0,40	0,90	0,80	0,90	0,40	0,70
<i>p_{36r13}</i>	0,90	0,90	0,80	0,70	0,80	0,70	0,90	0,70
<i>p_{36r21}</i>	0,80	0,90	0,50	0,30	0,70	0,70	0,90	0,90
<i>p_{36r22}</i>	0,80	0,20	0,10	0,30	0,80	0,70	0,60	0,90
<i>p_{36r24}</i>	0,40	0,70	0,90	0,30	0,80	0,50	0,80	0,90
<i>p_{36r31}</i>	0,50	0,90	0,90	0,70	0,80	0,80	0,80	0,70
<i>p_{36r32}</i>	0,90	0,80	0,60	0,70	0,80	0,50	0,60	0,70
<i>p_{36r33}</i>	0,20	0,70	0,40	0,80	0,80	0,90	0,40	0,50
<i>p_{37r11}</i>	0,90	0,60	0,80	0,60	0,90	0,70	0,90	0,80
<i>p_{37r12}</i>	0,70	0,90	0,80	0,70	0,50	0,80	0,80	0,60
<i>p_{37r21}</i>	0,90	0,70	0,80	0,60	0,60	0,90	0,50	0,40
<i>p_{37r32}</i>	0,80	0,90	0,50	0,30	0,80	0,70	0,40	0,60
<i>p_{37r33}</i>	0,80	0,40	0,60	0,80	0,80	0,60	0,90	0,30

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Se obtienen los valores de los criterios mediante la imputación o asignación, que serán utilizados para calcular la prioridad o preferencia de procesos para cada nodo. Se multiplica escalarmente cada vector de valoraciones de cada requerimiento por el vector de pesos correspondiente a la categoría de carga actual del nodo para obtener la prioridad según cada criterio y la prioridad nodal otorgada a cada requerimiento; esto se puede observar en la Tabla 110.

Tabla 110. Valores asignados a los criterios para calcular la prioridad o preferencia nodal que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo, con valores imputados/asignados correspondientes al escenario E2.

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
$p_{11}r_{11}$	0,7000	0,5000	0,7000	0,9000	0,8000	0,2000	0,3000	0,4000	0,6575
$Pri(p_{11}r_{11})$	0,0700	0,1000	0,2100	0,0900	0,1600	0,0100	0,0075	0,0100	
$p_{11}r_{12}$	0,8000	0,7000	0,4000	0,5000	0,3000	0,7000	0,2000	0,4000	0,5000
$Pri(p_{11}r_{12})$	0,0800	0,1400	0,1200	0,0500	0,0600	0,0350	0,0050	0,0100	
$p_{11}r_{21}$	0,3000	0,4000	0,5000	0,2000	0,9000	0,2000	0,5000	0,7000	0,5000
$Pri(p_{11}r_{21})$	0,0300	0,0800	0,1500	0,0200	0,1800	0,0100	0,0125	0,0175	
$p_{11}r_{22}$	0,5000	0,5000	0,7000	0,4000	0,8000	0,3000	0,5000	0,6000	0,6025
$Pri(p_{11}r_{22})$	0,0500	0,1000	0,2100	0,0400	0,1600	0,0150	0,0125	0,0150	
$p_{11}r_{23}$	0,5700	0,4890	0,3920	0,8020	0,5800	0,6660	0,7650	0,4320	0,5318
$Pri(p_{11}r_{23})$	0,0570	0,0978	0,1176	0,0802	0,1160	0,0333	0,0191	0,0108	
$p_{11}r_{24}$	0,3000	0,5000	0,9000	0,2000	0,6000	0,6000	0,7000	0,4000	0,5975
$Pri(p_{11}r_{24})$	0,0300	0,1000	0,2700	0,0200	0,1200	0,0300	0,0175	0,0100	
$p_{12}r_{11}$	0,4000	0,7000	0,5000	0,9000	1,0000	0,9000	0,8000	0,8000	0,7050
$Pri(p_{12}r_{11})$	0,0400	0,1400	0,1500	0,0900	0,2000	0,0450	0,0200	0,0200	
$p_{12}r_{12}$	0,2000	0,7000	0,3000	0,7000	0,8000	0,3000	0,8000	0,9000	0,5375
$Pri(p_{12}r_{12})$	0,0200	0,1400	0,0900	0,0700	0,1600	0,0150	0,0200	0,0225	
$p_{12}r_{21}$	0,5094	0,3531	0,4104	0,5333	0,7521	0,8969	0,7823	0,4750	0,5247
$Pri(p_{12}r_{21})$	0,0509	0,0706	0,1231	0,0533	0,1504	0,0448	0,0196	0,0119	
$p_{12}r_{22}$	0,9000	0,6000	0,7000	0,7000	0,8000	0,2000	0,5000	0,4000	0,6825
$Pri(p_{12}r_{22})$	0,0900	0,1200	0,2100	0,0700	0,1600	0,0100	0,0125	0,0100	
$p_{12}r_{31}$	0,2000	0,5000	0,7000	0,7000	0,3000	0,2000	0,7000	0,8000	0,5075
$Pri(p_{12}r_{31})$	0,0200	0,1000	0,2100	0,0700	0,0600	0,0100	0,0175	0,0200	
$p_{12}r_{33}$	0,4000	0,5000	0,7000	0,9000	0,3000	0,4000	0,5000	0,8000	0,5525
$Pri(p_{12}r_{33})$	0,0400	0,1000	0,2100	0,0900	0,0600	0,0200	0,0125	0,0200	
$p_{13}r_{11}$	0,6515	0,5071	0,2414	0,4394	0,4646	0,7081	0,4172	0,2535	0,4280
$Pri(p_{13}r_{11})$	0,0652	0,1014	0,0724	0,0439	0,0929	0,0354	0,0104	0,0063	
$p_{13}r_{12}$	0,7000	0,8000	0,7000	0,4000	0,9000	0,2000	0,9000	0,7000	0,7100
$Pri(p_{13}r_{12})$	0,0700	0,1600	0,2100	0,0400	0,1800	0,0100	0,0225	0,0175	
$p_{13}r_{13}$	0,7000	0,6000	0,7000	0,8000	0,9000	0,4000	0,8000	0,7000	0,7175
$Pri(p_{13}r_{13})$	0,0700	0,1200	0,2100	0,0800	0,1800	0,0200	0,0200	0,0175	
$p_{13}r_{21}$	0,7000	0,4000	0,9000	0,3000	0,5000	0,7000	0,2000	0,3000	0,5975
$Pri(p_{13}r_{21})$	0,0700	0,0800	0,2700	0,0300	0,1000	0,0350	0,0050	0,0075	
$p_{13}r_{22}$	0,5000	0,9000	0,8000	0,3000	0,5000	0,4000	0,9000	0,3000	0,6500
$Pri(p_{13}r_{22})$	0,0500	0,1800	0,2400	0,0300	0,1000	0,0200	0,0225	0,0075	
$p_{13}r_{31}$	0,5000	0,7000	0,3000	0,6000	0,8000	0,9000	0,9000	0,5000	0,5800
$Pri(p_{13}r_{31})$	0,0500	0,1400	0,0900	0,0600	0,1600	0,0450	0,0225	0,0125	
$p_{13}r_{32}$	0,6000	0,9000	0,3000	0,6000	0,4000	0,8000	0,7000	0,8000	0,5475
$Pri(p_{13}r_{32})$	0,0600	0,1800	0,0900	0,0600	0,0800	0,0400	0,0175	0,0200	
$p_{13}r_{33}$	0,6000	0,2000	0,4000	0,6000	0,9000	0,8000	0,5000	0,8000	0,5325
$Pri(p_{13}r_{33})$	0,0600	0,0400	0,1200	0,0600	0,1800	0,0400	0,0125	0,0200	
$p_{21}r_{12}$	0,2000	0,1000	0,3000	0,8000	0,7000	0,7000	0,5000	0,8000	

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p₂₁r₁₂)</i>	0,0200	0,0200	0,0900	0,0800	0,1400	0,0350	0,0125	0,0200	0,4175
<i>p₂₁r₁₃</i>	0,2000	0,4000	0,8000	0,9000	0,7000	0,4000	0,5000	0,2000	
<i>Pri(p₂₁r₁₃)</i>	0,0200	0,0800	0,2400	0,0900	0,1400	0,0200	0,0125	0,0050	0,6075
<i>p₂₁r₂₂</i>	0,3000	0,5000	0,8000	0,9000	0,5000	0,6000	0,2000	0,4000	
<i>Pri(p₂₁r₂₂)</i>	0,0300	0,1000	0,2400	0,0900	0,1000	0,0300	0,0050	0,0100	0,6050
<i>p₂₁r₂₃</i>	0,7000	0,5000	0,6000	0,8000	0,5000	0,2000	0,9000	0,4000	
<i>Pri(p₂₁r₂₃)</i>	0,0700	0,1000	0,1800	0,0800	0,1000	0,0100	0,0225	0,0100	0,5725
<i>p₂₁r₃₁</i>	0,7000	0,5000	0,8000	0,6000	0,9000	0,4000	0,9000	0,9000	
<i>Pri(p₂₁r₃₁)</i>	0,0700	0,1000	0,2400	0,0600	0,1800	0,0200	0,0225	0,0225	0,7150
<i>p₂₁r₃₃</i>	0,7000	0,4000	0,9000	0,8000	0,4000	0,8000	0,6000	0,6000	
<i>Pri(p₂₁r₃₃)</i>	0,0700	0,0800	0,2700	0,0800	0,0800	0,0400	0,0150	0,0150	0,6500
<i>p₂₂r₂₁</i>	0,5000	0,4000	0,7000	0,8000	0,6000	0,4000	0,6000	0,9000	
<i>Pri(p₂₂r₂₁)</i>	0,0500	0,0800	0,2100	0,0800	0,1200	0,0200	0,0150	0,0225	0,5975
<i>p₂₂r₂₂</i>	0,6886	0,7977	0,7705	0,4693	0,4795	0,9000	0,4727	0,4341	
<i>Pri(p₂₂r₂₂)</i>	0,0689	0,1595	0,2311	0,0469	0,0959	0,0450	0,0118	0,0109	0,6701
<i>p₂₂r₃₁</i>	0,5000	0,4000	0,5000	0,2000	0,4000	0,8000	0,2000	0,9000	
<i>Pri(p₂₂r₃₁)</i>	0,0500	0,0800	0,1500	0,0200	0,0800	0,0400	0,0050	0,0225	0,4475
<i>p₂₂r₃₃</i>	0,8000	0,4000	0,3000	0,2000	0,8000	0,9000	0,5000	0,8000	
<i>Pri(p₂₂r₃₃)</i>	0,0800	0,0800	0,0900	0,0200	0,1600	0,0450	0,0125	0,0200	0,5075
<i>p₂₃r₁₂</i>	0,5000	0,6000	0,8000	0,3000	0,5000	0,7000	0,5000	0,5000	
<i>Pri(p₂₃r₁₂)</i>	0,0500	0,1200	0,2400	0,0300	0,1000	0,0350	0,0125	0,0125	0,6000
<i>p₂₃r₂₄</i>	0,6000	0,2000	0,1000	0,7000	0,3000	0,8000	0,9000	0,4000	
<i>Pri(p₂₃r₂₄)</i>	0,0600	0,0400	0,0300	0,0700	0,0600	0,0400	0,0225	0,0100	0,3325
<i>p₂₃r₃₁</i>	0,3000	0,1000	0,4000	0,8000	0,7000	0,9000	0,4000	0,6000	
<i>Pri(p₂₃r₃₁)</i>	0,0300	0,0200	0,1200	0,0800	0,1400	0,0450	0,0100	0,0150	0,4600
<i>p₂₃r₃₂</i>	0,4918	0,6398	0,8337	0,5122	0,5143	0,5122	0,8235	0,3194	
<i>Pri(p₂₃r₃₂)</i>	0,0492	0,1280	0,2501	0,0512	0,1029	0,0256	0,0206	0,0080	0,6355
<i>p₂₃r₃₃</i>	0,3000	0,6000	0,8000	0,2000	0,9000	0,6000	0,4000	0,2000	
<i>Pri(p₂₃r₃₃)</i>	0,0300	0,1200	0,2400	0,0200	0,1800	0,0300	0,0100	0,0050	0,6350
<i>p₂₄r₁₁</i>	0,4000	0,6000	0,7000	0,9000	0,5000	0,7000	0,9000	0,5000	
<i>Pri(p₂₄r₁₁)</i>	0,0400	0,1200	0,2100	0,0900	0,1000	0,0350	0,0225	0,0125	0,6300
<i>p₂₄r₁₂</i>	0,6910	0,6400	0,5700	0,6070	0,3310	0,5810	0,7000	0,5450	
<i>Pri(p₂₄r₁₂)</i>	0,0691	0,1280	0,1710	0,0607	0,0662	0,0291	0,0175	0,0136	0,5552
<i>p₂₄r₂₃</i>	0,4000	0,9000	0,4000	0,8000	0,8000	0,7000	0,6000	0,3000	
<i>Pri(p₂₄r₂₃)</i>	0,0400	0,1800	0,1200	0,0800	0,1600	0,0350	0,0150	0,0075	0,6375
<i>p₂₄r₂₄</i>	0,5000	0,8000	0,7000	0,9000	0,3000	0,4000	0,8000	0,7000	
<i>Pri(p₂₄r₂₄)</i>	0,0500	0,1600	0,2100	0,0900	0,0600	0,0200	0,0200	0,0175	0,6275
<i>p₂₅r₂₁</i>	0,4675	0,9325	0,4974	0,8662	0,3000	0,6013	0,6974	0,5987	
<i>Pri(p₂₅r₂₁)</i>	0,0468	0,1865	0,1492	0,0866	0,0600	0,0301	0,0174	0,0150	0,5916
<i>p₃₁r₁₃</i>	0,6000	0,9000	0,6000	0,9000	0,7000	0,4000	0,9000	0,8000	
<i>Pri(p₃₁r₁₃)</i>	0,0600	0,2700	0,1200	0,1800	0,0700	0,0100	0,0225	0,0400	0,7725
<i>p₃₁r₃₁</i>	0,8000	0,3000	0,9000	0,5000	0,7000	0,3000	0,9000	0,7000	
<i>Pri(p₃₁r₃₁)</i>	0,0800	0,0900	0,1800	0,1000	0,0700	0,0075	0,0225	0,0350	0,5850
<i>p₃₁r₃₃</i>	0,4000	0,7000	0,7000	0,5000	0,9000	0,9000	0,7000	0,7000	
<i>Pri(p₃₁r₃₃)</i>	0,0400	0,2100	0,1400	0,1000	0,0900	0,0225	0,0175	0,0350	0,6550

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
$p_{32r_{11}}$	0,6000	0,9000	0,8000	0,5000	0,9000	0,7000	0,3000	0,7000	0,7400
$Pri(p_{32r_{11}})$	0,0600	0,2700	0,1600	0,1000	0,0900	0,0175	0,0075	0,0350	
$p_{32r_{12}}$	0,7000	0,4000	0,6000	0,9000	0,8000	1,0000	0,9000	0,7000	0,6525
$Pri(p_{32r_{12}})$	0,0700	0,1200	0,1200	0,1800	0,0800	0,0250	0,0225	0,0350	
$p_{32r_{13}}$	0,8000	0,9000	1,0000	0,7000	0,9000	0,3000	0,5000	0,9000	0,8450
$Pri(p_{32r_{13}})$	0,0800	0,2700	0,2000	0,1400	0,0900	0,0075	0,0125	0,0450	
$p_{32r_{23}}$	0,8000	0,8000	1,0000	0,9000	0,6000	0,8000	0,4000	0,9000	0,8350
$Pri(p_{32r_{23}})$	0,0800	0,2400	0,2000	0,1800	0,0600	0,0200	0,0100	0,0450	
$p_{33r_{11}}$	0,2000	0,7000	0,9000	0,8000	0,6000	0,9000	1,0000	0,3000	0,6925
$Pri(p_{33r_{11}})$	0,0200	0,2100	0,1800	0,1600	0,0600	0,0225	0,0250	0,0150	
$p_{33r_{12}}$	0,9000	1,0000	0,9000	0,7000	0,3000	0,5000	0,7000	0,3000	0,7850
$Pri(p_{33r_{12}})$	0,0900	0,3000	0,1800	0,1400	0,0300	0,0125	0,0175	0,0150	
$p_{33r_{13}}$	0,4804	0,6010	0,6196	0,7598	0,8134	0,6515	0,8423	0,3412	0,6400
$Pri(p_{33r_{13}})$	0,0480	0,1803	0,1239	0,1520	0,0813	0,0163	0,0211	0,0171	
$p_{33r_{21}}$	0,4000	0,6000	0,7000	0,8000	0,8000	0,4000	0,8000	0,8000	0,6700
$Pri(p_{33r_{21}})$	0,0400	0,1800	0,1400	0,1600	0,0800	0,0100	0,0200	0,0400	
$p_{33r_{22}}$	0,4168	0,2253	0,8568	0,5337	0,4147	0,4274	0,7568	0,5274	0,4848
$Pri(p_{33r_{22}})$	0,0417	0,0676	0,1714	0,1067	0,0415	0,0107	0,0189	0,0264	
$p_{33r_{23}}$	0,4000	0,6000	0,9000	1,0000	0,6000	0,4000	0,7000	0,8000	0,7275
$Pri(p_{33r_{23}})$	0,0400	0,1800	0,1800	0,2000	0,0600	0,0100	0,0175	0,0400	
$p_{33r_{32}}$	0,6000	0,6000	0,7000	0,8000	0,4000	0,5000	0,8000	0,8000	0,6525
$Pri(p_{33r_{32}})$	0,0600	0,1800	0,1400	0,1600	0,0400	0,0125	0,0200	0,0400	
$p_{33r_{33}}$	0,6000	1,0000	0,7000	0,4000	0,6000	0,8000	0,8000	0,9000	0,7250
$Pri(p_{33r_{33}})$	0,0600	0,3000	0,1400	0,0800	0,0600	0,0200	0,0200	0,0450	
$p_{34r_{12}}$	0,8000	1,0000	0,3000	0,5000	0,7000	0,3000	0,8000	0,7000	0,6725
$Pri(p_{34r_{12}})$	0,0800	0,3000	0,0600	0,1000	0,0700	0,0075	0,0200	0,0350	
$p_{34r_{13}}$	0,2000	1,0000	0,8000	0,5000	0,8000	0,6000	0,9000	0,7000	0,7325
$Pri(p_{34r_{13}})$	0,0200	0,3000	0,1600	0,1000	0,0800	0,0150	0,0225	0,0350	
$p_{34r_{22}}$	0,9000	0,8000	0,6000	0,8000	0,9000	0,8000	0,5000	0,6000	0,7625
$Pri(p_{34r_{22}})$	0,0900	0,2400	0,1200	0,1600	0,0900	0,0200	0,0125	0,0300	
$p_{34r_{23}}$	0,4000	0,6000	0,7000	0,8000	0,9000	0,5000	0,9000	0,6000	0,6750
$Pri(p_{34r_{23}})$	0,0400	0,1800	0,1400	0,1600	0,0900	0,0125	0,0225	0,0300	
$p_{34r_{24}}$	0,9000	0,4000	0,7000	0,4000	0,7000	0,5000	0,9000	0,7000	0,5700
$Pri(p_{34r_{24}})$	0,0900	0,1200	0,1400	0,0800	0,0700	0,0125	0,0225	0,0350	
$p_{34r_{31}}$	0,5000	0,6000	0,7000	0,9000	0,7000	0,9000	0,6000	0,7000	0,6925
$Pri(p_{34r_{31}})$	0,0500	0,1800	0,1400	0,1800	0,0700	0,0225	0,0150	0,0350	
$p_{34r_{32}}$	0,8000	0,6000	0,9000	0,5000	0,6000	0,9000	0,3000	0,9000	0,6750
$Pri(p_{34r_{32}})$	0,0800	0,1800	0,1800	0,1000	0,0600	0,0225	0,0075	0,0450	
$p_{34r_{33}}$	0,7882	0,6667	0,5828	0,5430	0,6581	0,9086	0,8591	0,4043	0,6342
$Pri(p_{34r_{33}})$	0,0788	0,2000	0,1166	0,1086	0,0658	0,0227	0,0215	0,0202	
$p_{35r_{12}}$	0,6000	0,2000	0,9000	0,6000	0,4000	0,8000	0,9000	0,2000	0,5125
$Pri(p_{35r_{12}})$	0,0600	0,0600	0,1800	0,1200	0,0400	0,0200	0,0225	0,0100	
$p_{35r_{13}}$	0,8000	0,9000	0,7000	0,3000	0,9000	0,8000	0,7000	0,4000	0,6975
$Pri(p_{35r_{13}})$	0,0800	0,2700	0,1400	0,0600	0,0900	0,0200	0,0175	0,0200	
$p_{35r_{22}}$	0,3000	0,8000	0,9000	0,4000	0,8000	0,6000	0,3000	0,6000	

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p_{35r22})</i>	0,0300	0,2400	0,1800	0,0800	0,0800	0,0150	0,0075	0,0300	0,6625
<i>p_{35r24}</i>	0,7000	0,7000	0,7000	0,7000	0,7000	0,7000	0,7000	0,7000	
<i>Pri(p_{35r24})</i>	0,0700	0,2100	0,1400	0,1400	0,0700	0,0175	0,0175	0,0350	0,7000
<i>p_{35r31}</i>	0,9000	0,8000	0,7000	0,4000	0,8000	0,3000	0,4000	0,8000	
<i>Pri(p_{35r31})</i>	0,0900	0,2400	0,1400	0,0800	0,0800	0,0075	0,0100	0,0400	0,6875
<i>p_{35r32}</i>	0,9000	0,7000	0,8000	0,6000	0,4000	0,9000	0,4000	0,7000	
<i>Pri(p_{35r32})</i>	0,0900	0,2100	0,1600	0,1200	0,0400	0,0225	0,0100	0,0350	0,6875
<i>p_{35r33}</i>	0,5000	0,7000	0,6000	0,9000	0,8000	0,5000	0,9000	0,7000	
<i>Pri(p_{35r33})</i>	0,0500	0,2100	0,1200	0,1800	0,0800	0,0125	0,0225	0,0350	0,7100
<i>p_{36r11}</i>	0,8000	0,9000	0,6000	0,5000	0,8000	0,9000	0,9000	0,6000	
<i>Pri(p_{36r11})</i>	0,0800	0,2700	0,1200	0,1000	0,0800	0,0225	0,0225	0,0300	0,7250
<i>p_{36r12}</i>	0,7000	0,9000	0,4000	0,9000	0,8000	0,9000	0,4000	0,7000	
<i>Pri(p_{36r12})</i>	0,0700	0,2700	0,0800	0,1800	0,0800	0,0225	0,0100	0,0350	0,7475
<i>p_{36r13}</i>	0,9000	0,9000	0,8000	0,7000	0,8000	0,7000	0,9000	0,7000	
<i>Pri(p_{36r13})</i>	0,0900	0,2700	0,1600	0,1400	0,0800	0,0175	0,0225	0,0350	0,8150
<i>p_{36r21}</i>	0,8000	0,9000	0,5000	0,3000	0,7000	0,7000	0,9000	0,9000	
<i>Pri(p_{36r21})</i>	0,0800	0,2700	0,1000	0,0600	0,0700	0,0175	0,0225	0,0450	0,6650
<i>p_{36r22}</i>	0,8000	0,2000	0,1000	0,3000	0,8000	0,7000	0,6000	0,9000	
<i>Pri(p_{36r22})</i>	0,0800	0,0600	0,0200	0,0600	0,0800	0,0175	0,0150	0,0450	0,3775
<i>p_{36r24}</i>	0,4000	0,7000	0,9000	0,3000	0,8000	0,5000	0,8000	0,9000	
<i>Pri(p_{36r24})</i>	0,0400	0,2100	0,1800	0,0600	0,0800	0,0125	0,0200	0,0450	0,6475
<i>p_{36r31}</i>	0,5000	0,9000	0,9000	0,7000	0,8000	0,8000	0,8000	0,7000	
<i>Pri(p_{36r31})</i>	0,0500	0,2700	0,1800	0,1400	0,0800	0,0200	0,0200	0,0350	0,7950
<i>p_{36r32}</i>	0,9000	0,8000	0,6000	0,7000	0,8000	0,5000	0,6000	0,7000	
<i>Pri(p_{36r32})</i>	0,0900	0,2400	0,1200	0,1400	0,0800	0,0125	0,0150	0,0350	0,7325
<i>p_{36r33}</i>	0,2000	0,7000	0,4000	0,8000	0,8000	0,9000	0,4000	0,5000	
<i>Pri(p_{36r33})</i>	0,0200	0,2100	0,0800	0,1600	0,0800	0,0225	0,0100	0,0250	0,6075
<i>p_{37r11}</i>	0,9000	0,6000	0,8000	0,6000	0,9000	0,7000	0,9000	0,8000	
<i>Pri(p_{37r11})</i>	0,0900	0,1800	0,1600	0,1200	0,0900	0,0175	0,0225	0,0400	0,7200
<i>p_{37r12}</i>	0,7000	0,9000	0,8000	0,7000	0,5000	0,8000	0,8000	0,6000	
<i>Pri(p_{37r12})</i>	0,0700	0,2700	0,1600	0,1400	0,0500	0,0200	0,0200	0,0300	0,7600
<i>p_{37r21}</i>	0,9000	0,7000	0,8000	0,6000	0,6000	0,9000	0,5000	0,4000	
<i>Pri(p_{37r21})</i>	0,0900	0,2100	0,1600	0,1200	0,0600	0,0225	0,0125	0,0200	0,6950
<i>p_{37r32}</i>	0,8000	0,9000	0,5000	0,3000	0,8000	0,7000	0,4000	0,6000	
<i>Pri(p_{37r32})</i>	0,0800	0,2700	0,1000	0,0600	0,0800	0,0175	0,0100	0,0300	0,6475
<i>p_{37r33}</i>	0,8000	0,4000	0,6000	0,8000	0,8000	0,6000	0,9000	0,3000	
<i>Pri(p_{37r33})</i>	0,0800	0,1200	0,1200	0,1600	0,0800	0,0150	0,0225	0,0150	0,6125

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Cálculo de las prioridades o preferencias de los procesos para acceder a los recursos compartidos disponibles

Una vez calculado la prioridad según cada criterio y la prioridad nodal otorgada a cada

requerimiento se proceden a calcular las prioridades globales o finales, es decir, con qué prioridad, o sea en qué orden, los recursos solicitados serán otorgados y a qué procesos se hará dicho otorgamiento. Para el cálculo de las prioridades finales se utilizan la Tabla 111 y la Tabla 112 (los valores marcados son valores obtenidos con métodos de imputación o asignación, y reemplazan los valores faltantes).

Tabla 111. Prioridades nodales de los procesos para acceder a cada recurso p_{11} al p_{25} - (E2).

Recursos	Prioridades Nodales de los Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{11}	0,6575	0,7050	0,4280	0,0000	0,0000	0,0000	0,6300	0,0000
r_{12}	0,5000	0,5375	0,7100	0,4175	0,0000	0,6000	0,5552	0,0000
r_{13}	0,0000	0,0000	0,7175	0,6075	0,0000	0,0000	0,0000	0,0000
r_{21}	0,5000	0,5247	0,5975	0,0000	0,5975	0,0000	0,0000	0,5916
r_{22}	0,6025	0,6825	0,6500	0,6050	0,6701	0,0000	0,0000	0,0000
r_{23}	0,5318	0,0000	0,0000	0,5725	0,0000	0,0000	0,6375	0,0000
r_{24}	0,5975	0,0000	0,0000	0,0000	0,0000	0,3325	0,6275	0,0000
r_{31}	0,0000	0,5075	0,5800	0,7150	0,4475	0,4600	0,0000	0,0000
r_{32}	0,0000	0,0000	0,5475	0,0000	0,0000	0,6355	0,0000	0,0000
r_{33}	0,0000	0,5525	0,5325	0,6500	0,5075	0,6350	0,0000	0,0000

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Nota: los valores marcados indican resultados de imputación.

Tabla 112. Prioridades nodales de los procesos para acceder a cada recurso p_{31} al p_{37} - (E2).

Recursos	Prioridades Nodales de los Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{11}	0,0000	0,7400	0,6925	0,0000	0,0000	0,7250	0,7200
r_{12}	0,0000	0,6525	0,7850	0,6725	0,5125	0,7475	0,7600
r_{13}	0,7725	0,8450	0,6400	0,7325	0,6975	0,8150	0,0000
r_{21}	0,0000	0,0000	0,6700	0,0000	0,0000	0,6650	0,6950
r_{22}	0,0000	0,0000	0,4848	0,7625	0,6625	0,3775	0,0000
r_{23}	0,0000	0,8350	0,7275	0,6750	0,0000	0,0000	0,0000
r_{24}	0,0000	0,0000	0,0000	0,5700	0,7000	0,6475	0,0000
r_{31}	0,5850	0,0000	0,0000	0,6925	0,6875	0,7950	0,0000
r_{32}	0,0000	0,0000	0,6525	0,6750	0,6875	0,7325	0,6475
r_{33}	0,6550	0,0000	0,7250	0,6342	0,7100	0,6075	0,6125

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019)

Posteriormente se calcula el vector de pesos finales que se utilizará en el proceso final de agregación para determinar el orden o prioridad de acceso a los recursos; esto se muestran en la Tabla 113 y Tabla 114.

Tabla 113. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos p_{11} al p_{25} - (E2).

Pesos	Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
Finales	0,2000	0,1330	0,2000	0,1330	0,1330	0,2000	0,0670	0,2000
Finales Normalizados	0,0970	0,0650	0,0970	0,0650	0,0650	0,0970	0,0320	0,0970

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 114. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos p_{31} al p_{37} - (E2).

Pesos	Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
Finales	0,1330	0,0670	0,0670	0,2000	0,0670	0,0670	0,2000
Finales Normalizados	0,0650	0,0320	0,0320	0,0970	0,0320	0,0320	0,0970

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Las prioridades nodales indicadas en la Tabla 111 y Tabla 112 tomadas fila por fila, es decir, respecto de cada recurso, se multiplicarán escalarmente por el vector de pesos finales normalizados indicados en la Tabla 113 y Tabla 114, para obtener las prioridades globales finales de acceso de cada proceso a cada recurso y de allí, el orden o prioridad con que se asignarán los recursos y a qué proceso se asignará cada uno de ellos; esto se indican en la Tabla 115 y Tabla 116.

El mayor de estos productos hechos para los distintos procesos en relación al mismo recurso indicará cuál de los procesos tendrá acceso al recurso (en caso de empates se podría optar por dar ganador al proceso identificado con el número menor); esto se muestran en rojo en la Tabla 115 y Tabla 116.

Tabla 115. Prioridades globales finales de los procesos para acceder a cada recurso p_{11} al p_{25} - (E2).

Recursos	Prioridades globales finales de los procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{11}	0,0638	0,0458	0,0415	0,0000	0,0000	0,0000	0,0202	0,0000
r_{12}	0,0485	0,0349	0,0689	0,0271	0,0000	0,0582	0,0178	0,0000
r_{13}	0,0000	0,0000	0,0696	0,0395	0,0000	0,0000	0,0000	0,0000
r_{21}	0,0485	0,0341	0,0580	0,0000	0,0388	0,0000	0,0000	0,0574
r_{22}	0,0584	0,0444	0,0631	0,0393	0,0436	0,0000	0,0000	0,0000
r_{23}	0,0516	0,0000	0,0000	0,0372	0,0000	0,0000	0,0204	0,0000
r_{24}	0,0580	0,0000	0,0000	0,0000	0,0000	0,0323	0,0201	0,0000
r_{31}	0,0000	0,0330	0,0563	0,0465	0,0291	0,0446	0,0000	0,0000
r_{32}	0,0000	0,0000	0,0531	0,0000	0,0000	0,0616	0,0000	0,0000
r_{33}	0,0000	0,0359	0,0517	0,0423	0,0330	0,0616	0,0000	0,0000

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 116. Prioridades globales finales de los procesos para acceder a cada recurso p_{31} al p_{37} - (E2).

Recursos	Prioridades globales finales de los procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{11}	0,0000	0,0237	0,0222	0,0000	0,0000	0,0232	0,0698
r_{12}	0,0000	0,0209	0,0251	0,0652	0,0164	0,0239	0,0737
r_{13}	0,0502	0,0270	0,0205	0,0711	0,0223	0,0261	0,0000
r_{21}	0,0000	0,0000	0,0214	0,0000	0,0000	0,0213	0,0674
r_{22}	0,0000	0,0000	0,0155	0,0740	0,0212	0,0121	0,0000
r_{23}	0,0000	0,0267	0,0233	0,0655	0,0000	0,0000	0,0000
r_{24}	0,0000	0,0000	0,0000	0,0553	0,0224	0,0207	0,0000
r_{31}	0,0380	0,0000	0,0000	0,0672	0,0220	0,0254	0,0000
r_{32}	0,0000	0,0000	0,0209	0,0655	0,0220	0,0234	0,0628
r_{33}	0,0426	0,0000	0,0232	0,0615	0,0227	0,0194	0,0594

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

La sumatoria de todos estos productos en relación al mismo recurso indicará la prioridad que tendrá dicho recurso para ser asignado. Esto constituye la Función de Asignación para Sistemas Distribuidos (**FASD**), que se muestra en la Tabla 117.

La suma de todos estos productos en relación con el mismo recurso indicará la prioridad que deberá asignarse a este recurso, en relación con los demás recursos que también deberán asignarse.

Tabla 117. Prioridades globales finales para asignar los recursos, constituye la Función de Asignación para Sistemas Distribuidos (FASD) - (E2).

FASD	Prioridad Global Final Para Asignar el Recurso	Proceso
r_{11}	0,3102	r_{11} al p_{37}
r_{12}	0,4807	r_{12} al p_{37}
r_{13}	0,3263	r_{13} al p_{34}
r_{21}	0,3469	r_{21} al p_{37}
r_{22}	0,3715	r_{22} al p_{34}
r_{23}	0,2247	r_{23} al p_{34}
r_{24}	0,2087	r_{24} al p_{11}
r_{31}	0,3621	r_{31} al p_{34}
r_{32}	0,3094	r_{32} al p_{34}
r_{33}	0,4533	r_{33} al p_{23}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El orden final de asignación de los recursos y los procesos destinatarios de los mismos se obtiene ordenando la Tabla 117 (FASD), que se muestra en la Tabla 118.

Tabla 118. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso, constituye la Función de Asignación para Sistemas Distribuidos Ordenada (FASDO) - (E2).

FASDO	Prioridad Global Final para Asignar el Recurso	Proceso
r_{12}	0,4807	r_{12} al p_{37}
r_{33}	0,4533	r_{33} al p_{23}
r_{22}	0,3715	r_{22} al p_{34}
r_{31}	0,3621	r_{31} al p_{34}
r_{21}	0,3469	r_{21} al p_{37}
r_{13}	0,3263	r_{13} al p_{34}
r_{11}	0,3102	r_{11} al p_{37}
r_{32}	0,3094	r_{32} al p_{34}
r_{23}	0,2247	r_{23} al p_{34}
r_{24}	0,2087	r_{24} al p_{11}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El siguiente paso es reiterar el procedimiento como lo establece (La Red Martínez, 2017), retirando de las solicitudes de recursos las asignaciones ya hechas; también debe tenerse en cuenta que los recursos asignados quedarán disponibles cuando los procesos los hayan liberado, pudiendo por lo tanto ser asignados a otros procesos. El resultado de las sucesivas iteraciones finaliza una vez que se han atendido todas las solicitudes de recursos de todos los procesos respetando la exclusión mutua y las prioridades de los procesos, las prioridades nodales y las prioridades finales; el resultado de las sucesivas iteraciones se muestra en la Tabla 119, Tabla 120, Tabla 121, Tabla 122, Tabla 123, Tabla 124, Tabla 125, Tabla 126, Tabla 127, Tabla 128 y Tabla 129.

Tabla 119. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (segunda iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,4070	r_{12} al p_{13}
0,3917	r_{33} al p_{34}
0,2975	r_{22} al p_{13}
0,2949	r_{31} al p_{13}
0,2795	r_{21} al p_{13}
0,2552	r_{13} al p_{13}
0,2439	r_{32} al p_{37}
0,2403	r_{11} al p_{11}
0,1592	r_{23} al p_{11}
0,1507	r_{24} al p_{34}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 120. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (tercera iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,3381	r_{12} al p_{34}
0,3302	r_{33} al p_{37}
0,2386	r_{31} al p_{21}
0,2345	r_{22} al p_{11}
0,2215	r_{21} al p_{25}
0,1856	r_{13} al p_{31}
0,1811	r_{32} al p_{23}
0,1765	r_{11} al p_{12}
0,1076	r_{23} al p_{21}
0,0955	r_{24} al p_{23}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 121. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (cuarta iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,2799	r_{12} al p_{23}
0,2707	r_{33} al p_{13}
0,1922	r_{31} al p_{23}
0,1760	r_{22} al p_{12}
0,1642	r_{21} al p_{11}
0,1354	r_{13} al p_{21}
0,1307	r_{11} al p_{13}
0,1194	r_{32} al p_{13}
0,0704	r_{23} al p_{32}
0,0632	r_{24} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 122. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (quinta iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,2191	r_{33} al p_{31}
0,2147	r_{12} al p_{11}
0,1475	r_{31} al p_{31}
0,1317	r_{22} al p_{22}
0,1157	r_{21} al p_{22}
0,0959	r_{13} al p_{36}
0,0892	r_{11} al p_{32}
0,0663	r_{32} al p_{36}
0,0437	r_{23} al p_{33}
0,0408	r_{24} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 123. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (sexta iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,1765	r_{33} al p_{21}
0,1662	r_{12} al p_{12}
0,1095	r_{31} al p_{12}
0,0881	r_{22} al p_{21}
0,0768	r_{21} al p_{12}
0,0698	r_{13} al p_{32}
0,0655	r_{11} al p_{36}
0,0429	r_{32} al p_{35}
0,0204	r_{23} al p_{24}
0,0201	r_{24} al p_{24}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 124. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (séptima iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,1343	r_{33} al p_{12}
0,1312	r_{12} al p_{21}
0,0765	r_{31} al p_{22}
0,0488	r_{22} al p_{35}
0,0428	r_{13} al p_{35}
0,0427	r_{21} al p_{33}
0,0423	r_{11} al p_{33}
0,0209	r_{32} al p_{33}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 125. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (octava iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,1041	r_{12} al p_{33}
0,0983	r_{33} al p_{22}
0,0474	r_{31} al p_{36}
0,0276	r_{22} al p_{33}
0,0213	r_{21} al p_{36}
0,0205	r_{13} al p_{33}
0,0202	r_{11} al p_{24}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 126. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (novena iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,0790	r_{12} al p_{36}
0,0654	r_{33} al p_{33}
0,0220	r_{31} al p_{35}
0,0121	r_{22} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 127. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décima iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,0550	r_{12} al p_{32}
0,0422	r_{33} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 128. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo primera iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,0342	r_{12} al p_{24}
0,0194	r_{33} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 129. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo segunda iteración) - (E2).

Prioridad Global Final Ordenada	Asignación
0,0164	r_{12} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

De esta manera se han atendido todas las solicitudes de recursos de todos los procesos respetando la exclusión mutua y las prioridades de los procesos, las prioridades nodales y las prioridades finales.

3.5. Evaluación

El modelo propuesto en el modelo decisión fue efectivo, permitió completar los valores faltantes para que se establezca las asignaciones de los recursos a los procesos que han solicitado. La diferencia entre los valores originales en (La Red Martínez, 2017) y los valores imputados están dentro de lo razonablemente esperable.

En esta propuesta se consideran como parte de la información histórica para aplicar una nueva imputación/asignación, los resultados de las imputaciones/asignaciones en los diferentes ciclos de recolección.

3.6. Discusiones y comentarios

El modelo de decisión propuesto considera la imputación de datos cuando la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos no llega en su

totalidad, reemplazando esos valores faltantes por valores estimados, para que finalmente el modelo de decisión establezca un correcto orden de asignación de recursos a los procesos, respetando la modalidad de exclusión mutua.

Esta propuesta considera la información estimada (resultado de las imputaciones/asignaciones en los diferentes ciclos) como parte de la información histórica para aplicar una nueva imputación/asignación.

El mecanismo no solo se aplica cuando falta la totalidad de información de control de los criterios establecidos para calcular la carga computacional y/ o prioridad o preferencia de los procesos respecto del nodo, también es necesario cuando falte más del 30% de los valores de las variables que se recibe en un ciclo de recepción de información de control que envía los nodos al nodo central y cuando se disponga de pocos registros históricos para aplicar la imputación.

El contenido de este capítulo ha servido para armar un artículo que será enviado a una revista de alto impacto.

Capítulo IV – Modelo de decisión

4.1. Introducción a los modelos de decisión

El Modelo de Decisión permite definir objetivos claros, basándose en la recolección de datos y su relevancia, estudiando las posibles alternativas y evaluando las consecuencias. La tarea consiste en seleccionar la mejor alternativa para lograr la mejor solución.

Los sistemas distribuidos tienen como objetivo principal el de compartir recursos, sea cual fuera el lugar donde sea solicitado, pero teniendo en cuenta la disponibilidad de los mismos, ya sea para las aplicaciones, equipos y datos. El acceso a los recursos compartidos se realiza respetando la modalidad de exclusión mutua distribuida y mediante los requisitos de sincronización. La decisión será en función a la demanda de recursos solicitados por parte de los procesos y las prioridades que tendrán cada proceso para acceder al recurso. Se considerará el valor de agregación de las prioridades de los diferentes nodos en cuanto a conceder los recursos compartidos. Se tendrán en cuenta las prioridades de los nodos para la asignación de todos los recursos compartidos de acuerdo con las necesidades de recursos de cada proceso.

Las premisas, estructuras de datos y el operador mencionado en (La Red Martínez, 2017), (Agostini & La Red Martínez, 2019), (Agostini et al., 2018) y (Agostini, 2019) son aplicadas para resolver el escenario que se describe a continuación. Se aplicarán métodos de imputación cuando se presenten valores faltantes en la información de control que llega al nodo central por parte de los nodos, para poder completar dichos valores. Los mecanismos a utilizar para imputar son el *K-Means* y la *Media Ponderada*.

4.2. Escenario propuesto

Se considera una macro imagen del sistema, que consiste en resolver un conjunto de

solicitudes de recursos por partes de los procesos, que ocasiona un cambio constante e iterativo del estado del sistema, manteniendolo actualizado, de acuerdo a la concreción de los distintos requerimientos. Como ya se ha mencionado, los nodos que componen el sistema y que alojan recursos y procesos, pueden tener distintos estados. Esta información la debe gestionar y mantener actualizada el Nodo Central, como se muestra en la Figura 8.

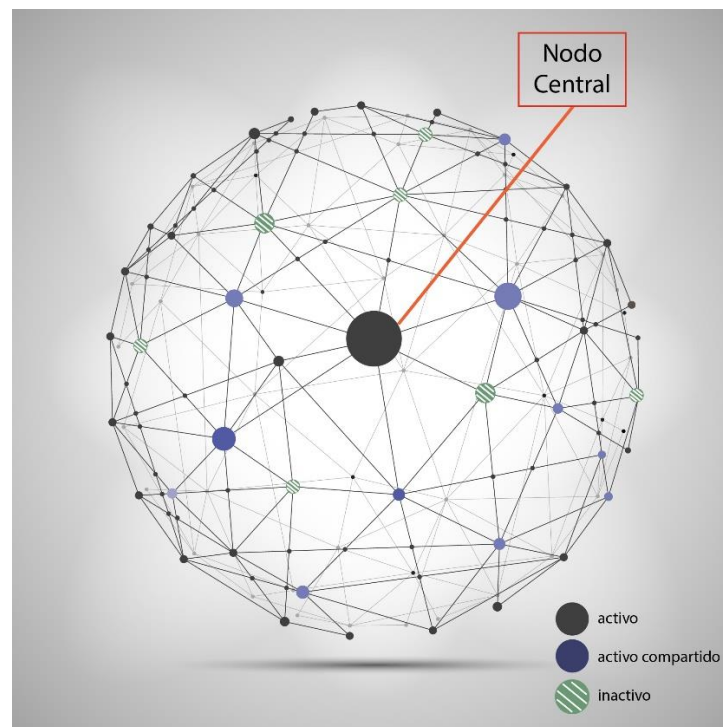


Figura 8. Estados de los nodos.

4.3. Operador de agregación utilizado por el Modelo de Decisión

En (La Red Martínez, 2017), las características de los operadores de agregación descritos permiten considerar que el método propuesto pertenece a la familia de operadores de agregación Neat-OWA. La propuesta desarrollada en el trabajo citado consiste en generar un modelo de decisión y su correspondiente operador de agregación para la gestión de grupos de procesos que comparten recursos, realizando modificaciones de los operadores de la familia OWA

mencionados.

La propuesta mencionada anteriormente no considera la información de control faltante, por lo que en este trabajo se propone incorporar capas de imputación/asignación, teniendo en cuenta los siguientes pasos:

- a) Identificar el propósito del problema;
- b) Identificar las alternativas;
- c) Listar los criterios que van a ser tenidos en cuenta a elegir la mejor alternativa;
 - Indicadores de prestaciones de cada nodo.
 - Cálculo de la carga computacional actual de los nodos.
 - Establecimiento de categorías de la carga computacional de cada nodo y los vectores de pesos asociados a ellos.
 - Cálculo de las prioridades o preferencias de los procesos teniendo en cuenta el estado del nodo.
 - Cálculo de las prioridades o preferencias de los procesos para acceder a los recursos compartidos disponibles.

a) Identificar el propósito del problema

En este trabajo lo que se resuelve son las solicitudes de recursos por parte de los procesos que se encuentran en nodos distribuidos. A partir de un nodo central que aloja al runtime, se recibe la información de control respecto del estado de los nodos y de los requerimientos de acceso de

procesos a recursos, proporcionada por todos los nodos que están compartiendo información, y a través de un operador de agregación determina y soluciona la siguiente premisa: “que los procesos accedan a recursos compartidos en la modalidad de exclusión mutua, incorporando mecanismos de imputación de datos para los casos en que la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos sea incompleta y/o inexistente”.

El nodo central puede recibir información incompleta o inexistente sobre alguno de los nodos; es decir, algunos de los criterios de evaluación de carga puede no estar disponible (ej.: CPU, memoria virtual, prioridad de proceso, etc.). Como solución a este problema se mejora un modelo de decisión y su correspondiente operador de agregación agregando una capa de imputación de datos, para que el modelo de decisión pueda establecer un orden de asignación de recursos. Cuando la información suministrada por los nodos esté completa, el modelo de decisión se ejecuta sin necesidad de la imputación de datos.

b) Identificar las alternativas

Son las opciones que ayudan lograr el propósito que se trazó en el paso anterior (Identificar el propósito del problema). El Runtime alojado en cada nodo gestiona los procesos y recursos compartidos, para que luego se defina el escenario correspondiente. En un determinado ciclo de recolección de información proporcionada por todos los nodos, el nodo central puede recibir información incompleta y/o inexistente de uno o algunos de los criterios utilizados para la carga computacional de cada nodo.

Además, para cada nodo, se consideran distintos criterios para calcular la prioridad o preferencia que tiene cada uno, en las asignaciones de recursos a procesos. Al igual que el caso anterior, algunos de los valores de estos criterios podrían no estar disponibles en un ciclo de

recolección.

Para obtener un indicador de las prestaciones de cada nodo, se tendrá en cuenta sus características, por ejemplo, la velocidad de procesamiento, capacidad de memoria, velocidad de transmisión de datos, velocidad de entrada de salida, entre otros. Estas características permitirán identificar nodos de altas, media y bajas prestaciones. El objetivo principal será favorecer a los nodos de mejores prestaciones por sobre los demás, por considerar que esos nodos estarán en mejores condiciones para resolver las diferentes asignaciones con diferentes cargas de trabajo. Cuando un nodo arranca, informará a través de su runtime, el detalle de sus características al runtime del nodo central, permitiendo que este último pueda clasificarlos según sus prestaciones.

La imputación de datos será la herramienta a utilizar en el modelo decisión propuesto para completar los valores faltantes con valores estimados, y que finalmente se pueda establecer un correcto orden de asignación de los recursos a los procesos que solicitan, en sistemas distribuidos. Se opta por utilizar por lo ya expuesto en el **Capítulo II** y **Capítulo III** dos mecanismos de imputación teniendo en cuenta las siguientes consideraciones:

Escenario 1; se aplicará el método K-Means cuando la información de control de las variables que indican el estado de carga de alguno o algunos de los nodos sea *incompleta*, utilizando registros históricos de ciclos de recolecciones anteriores alojados en el nodo central.

Escenario 2; se aplicará el mecanismo de medias ponderadas cuando *no se reciba la totalidad* de información de control de las variables que indican el estado de carga de alguno o algunos de los nodos, utilizando registros históricos de ciclos de recolecciones anteriores alojados en el nodo central.

Si no se cuenta con registros históricos de ciclos de recolecciones anteriores, se asignará

un valor por defecto en función a las prestaciones del nodo.

En caso de que la información suministrada por los nodos esté completa, el modelo de decisión se ejecuta sin necesidad de la imputación de datos, como lo establecen (La Red Martínez, 2017) y (Agostini, 2019).

c) Listar los criterios que van a ser tenidos en cuenta para elegir la mejor alternativa

- Indicadores de prestaciones de cada nodo.

El runtime del nodo que arranca estará informando al runtime del nodo central sus características, por ejemplo, la velocidad de procesamiento, capacidad de memoria, velocidad de transmisión de datos, velocidad de entrada de salida, entre otros. Para obtener un indicador de las prestaciones del nodo, se tienen en cuenta sus características, por ejemplo, la velocidad de procesamiento, capacidad de memoria, velocidad de transmisión de datos, velocidad de entrada de salida, entre otros.

Para la carga nodal se asumirán valores para los indicadores de prestaciones de los nodos, para los nodos de prestaciones Altas se asumirá un valor de 30%, los nodos de prestaciones Medias un valor de 50% y los nodos de prestaciones Bajas un valor de 70%.

En cuanto a la prioridad de los procesos respecto del nodo, se asumirá un valor de 0,30 para los nodos de prestaciones Altas, valor de 0,50 para los nodos de prestaciones Medias y un valor de 0,70 para los nodos de prestaciones Bajas.

- **Cálculo de la carga computacional actual de los nodos.**

Se contemplará el cálculo de la carga actual de los nodos. Para obtener un indicador de la carga computacional actual de cada nodo, se pueden adoptar diferentes criterios; en esta propuesta los criterios serán el porcentaje de uso de la CPU, el porcentaje de uso de la memoria y el porcentaje de uso de la operación de entrada/salida. La carga computacional de cada nodo, el número de criterios para determinar la carga de los nodos, los criterios que se aplican y el cálculo de la carga computacional de cada nodo, son los mencionados en (La Red Martínez, 2017) y (Agostini, 2019).

En un ciclo de recolección de información, el nodo central puede información de control con valores faltantes de uno o algunos de los criterios utilizados para el cálculo de la carga computacional, entonces se aplica la imputación de datos para completar los valores faltantes, utilizando información histórica de ciclos de recolección anteriores, si no se cuenta con dicha información se asumirá un valor por defecto en función a la carga computacional informada y/o las características del nodo.

- **Establecimiento de categorías de la carga computacional de cada nodo y los vectores de pesos asociados a ellos.**

Para este ejemplo, al igual que en en (La Red Martínez, 2017) y (Agostini, 2019), las categorías serán: Alta (si la carga es mayor al 70%), Media (si la carga está entre el 40% y el 70% inclusive) y Baja (si la carga es menor al 40%).

La asignación de pesos a los distintos criterios y los vectores de pesos de cada categoría de carga son los utilizados en investigaciones previas mencionadas anteriormente.

- **Cálculo de las prioridades o preferencias de los procesos teniendo en cuenta el estado del nodo**

Las prioridades se calculan en cada nodo para cada solicitud de recursos originada en cada proceso; el cálculo considera el vector de pesos correspondiente según la carga actual del nodo y el vector de los valores concedidos por el nodo según los criterios de evaluación de la solicitud.

En un ciclo de recolección de información de control el nodo central puede que reciba de alguno o algunos nodos información de control con valores faltantes sobre los criterios de evaluación de carga, utilizados para calcular las prioridades o preferencias de los procesos, entonces para completar los valores faltantes, se aplica el método de imputación utilizando información histórica de ciclos de recolecciones anteriores, disponibles en el nodo central, o si no se dispone de dicha información se asignará un valor por defecto.

- **Cálculo de las prioridades o preferencias de los procesos para acceder a los recursos compartidos disponibles**

Para el cálculo de las prioridades finales se utiliza los valores obtenidos de las prioridades o preferencias nodales calculadas en la etapa anterior; donde cada fila contiene la información de las prioridades nodales de los distintos procesos para acceder a un determinado recurso.

Posteriormente se ve la necesidad de calcular el vector de pesos finales que se utilizará en el proceso de agregación para determinar el orden o la prioridad de acceso a los recursos.

A continuación, se normalizan los pesos obtenidos recientemente dividiendo cada uno de ellos por la sumatoria de todos, obteniendo un vector de peso normalizado (en el rango de 0 a 1 inclusive) y con la restricción de que la suma de los elementos del vector debe dar 1.

Las prioridades nodales tomadas fila por fila para cada recurso se multiplicarán escalarmente por el vector de pesos finales normalizados, y así obtener las prioridades globales finales de acceso de cada proceso a cada recurso. El orden o la prioridad con que se asignarán los recursos y a qué proceso se asignará cada uno, estará determinado por el cálculo de la Prioridad global final.

El mayor de estos productos realizados para los diferentes procesos en relación con el mismo recurso indicará cuál de los procesos tendrá acceso al recurso.

La sumatoria de todos estos productos en relación al mismo recurso indicará la prioridad que tendrá dicho recurso para ser asignado, en relación a los demás recursos que también tendrán que asignarse. Esto es lo que se denominará Función de Asignación de Sistemas Distribuidos Lingüística (**FASD**).

Seguidamente se repite el procedimiento, pero eliminando las solicitudes de asignaciones ya realizadas; cabe señalar que los recursos asignados estarán disponibles una vez que los procesos los liberen y, por lo tanto, podrán asignarse a otros procesos.

4.4. Modelo de decisión propuesto

El objetivo principal del modelo de decisión propuesto en esta investigación es la de considerar un entorno de ejecución distribuida de procesos, donde los mismos pueden estar agrupados, con solicitudes diferentes para poder acceder a los recursos compartidos, para que sea seleccionado el método de agregación más adecuado al escenario posible, para que los recursos solicitados por los procesos se generen por una secuencia de asignación, respetando la modalidad de exclusión mutua distribuida para el acceso a los recursos por parte de los procesos.

Es considerada por el modelo propuesto la información que caracterizan a los procesos (si

los procesos pertenecen a grupos de procesos o si los mismos son independientes, y si la aplicación a la que pertenecen requiere algún tipo de consenso respecto del acceso a los recursos), las prioridades nodales, estado de los nodos, recursos, procesos y del sistema para producir una solución global satisfactoria. Este método evalúa esta información, para determinar cuál es el escenario según la carga de trabajo proporcionada, y en base a ello elige el operador de agregación correspondiente, como se indica en la Figura 9.

En el modelo de decisión propuesto se ejecuta el operador de agregación, que corresponde al escenario general (Función de Asignación en Sistemas Distribuidos, FASD) considerando la imputación de datos cuando se presentan valores faltantes en la información de control de alguno de los nodos que llega al nodo central en un Sistema Distribuido.

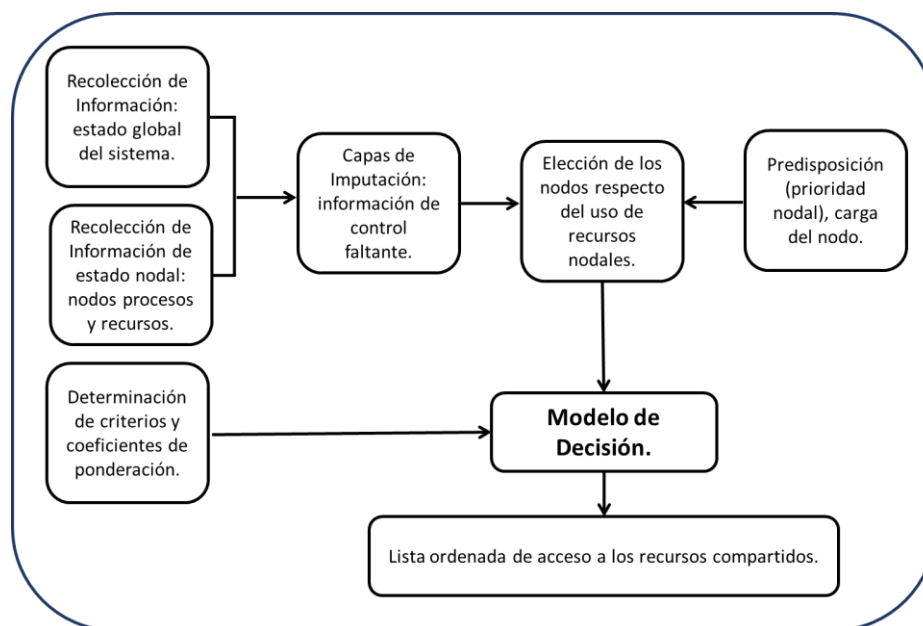


Figura 9. Modelo global de decisión para acceso a recursos compartidos, considerando la imputación de datos

Fuente: Elaboración propia.

De acuerdo a la información ingresada, se podrá seguir con el escenario general o hacer uso de uno de los otros operadores para los siguientes escenarios, la Función de Asignación en

Sistemas Distribuidos FASD, aplicando la imputación de datos con K-Menas cuando la información de control de alguno de los nodos del SD llega incompleta al nodo central (ver **Capítulo II**, correspondiente al escenario E1) y utilizando un mecanismo de Medias Ponderadas cuando la información de control de alguno de los nodos del SD que llega al nodo central es inexistente (ver **Capítulo III**, correspondiente al escenario E2). Posteriormente se deberá aplicar el proceso de resolución, como se indica en la Figura 10.

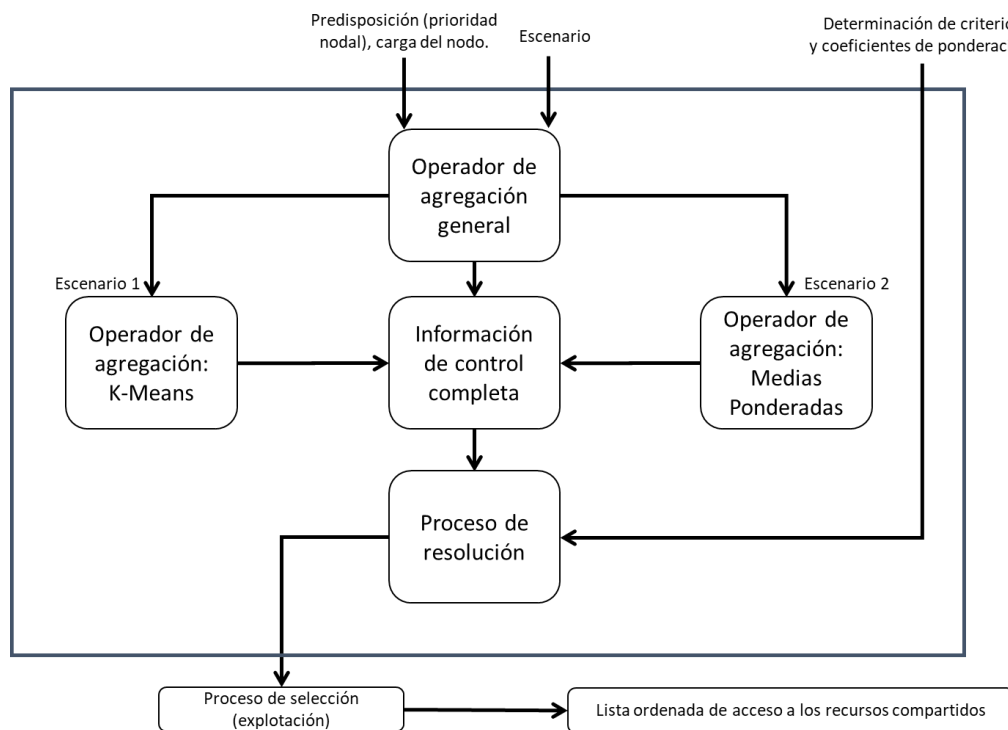


Figura 10. Modelo de decisión propuesto

Fuente: Elaboración propia.

4.5. Pautas para aplicar el método de imputación cuando se reciba de algún nodo información de control incompleta o inexistente

En función a la información de control histórica (disponible en el nodo central y que corresponde a la información de carga computacional y preferencias de los procesos en ciclos de recolecciones anteriores) disponible para los métodos de imputación/asignación, se describirán

```

graph TD
    Inicio([Inicio a]) --> D1{¿Falta Información de control? b}
    D1 -- no --> MD[MD sin Imputación d]
    D1 -- si --> D2{¿Hay registros históricos? c}
    D2 -- no --> D3{¿Hay varios criterios de ponderación? j}
    D2 -- si --> D4{¿Información de control faltante parcial? e}
    D3 -- si --> Def1[Defecto 1 k]
    D3 -- no --> Def2[Defecto 2 l]
    D4 -- si --> D5{¿Información de control nodal? f}
    D4 -- no --> D6{¿Información de control nodal? m}
    D5 -- si --> KM[K-Means g]
    D5 -- no --> D7{¿Hay históricos p_kl r_ij? h}
    D6 -- si --> MP[Medias Ponderadas n]
    D6 -- no --> D8{¿Hay históricos p_kl r_ij? o}
    D7 -- si --> KM
    D7 -- no --> D9{¿Considera históricos p_kl r_xx? i}
    D9 -- si --> KM
    D9 -- no --> Def2
    D8 -- si --> MP
    D8 -- no --> Def2
    MD --> Def1
    Def1 --> Def2
    Def2 --> Def2
  
```

El diagrama de flujo describe la selección de algoritmo de aprendizaje automático basado en la disponibilidad de información de control y registros históricos. El proceso comienza en el nodo 'Inicio' (a) y se dirige a la decisión '¿Falta Información de control?' (b). Si la respuesta es 'no', se selecciona 'MD sin Imputación' (d). Si la respuesta es 'si', se dirige a la decisión '¿Hay registros históricos?' (c). Si la respuesta es 'no', se dirige a la decisión '¿Hay varios criterios de ponderación?' (j). Si la respuesta es 'si', se selecciona 'Defecto (1)' (k). Si la respuesta es 'no', se selecciona 'Defecto (2)' (l). Si la respuesta es 'si', se dirige a la decisión '¿Información de control faltante parcial?' (e). Si la respuesta es 'si', se dirige a la decisión '¿Información de control nodal?' (f). Si la respuesta es 'no', se dirige a la decisión '¿Información de control nodal?' (m). Si la respuesta es 'si', se selecciona 'Medias Ponderadas' (n). Si la respuesta es 'no', se dirige a la decisión '¿Hay históricos' (o). Si la respuesta es 'si', se selecciona 'Medias Ponderadas' (n). Si la respuesta es 'no', se dirige a la decisión '¿Hay históricos' (h). Si la respuesta es 'si', se selecciona 'K-Means' (g). Si la respuesta es 'no', se dirige a la decisión '¿Considera históricos' (i). Si la respuesta es 'si', se selecciona 'K-Means' (g). Si la respuesta es 'no', se selecciona 'Defecto (2)' (l). Si la respuesta es 'si', se selecciona 'K-Means' (g). Si la respuesta es 'no', se selecciona 'Defecto (2)' (l). Si la respuesta es 'si', se selecciona 'K-Means' (g). Si la respuesta es 'no', se selecciona 'Defecto (2)' (l).

En la Tabla 130 se describen las pautas que se tienen en cuenta para aplicar la imputación los diferentes escenarios, tomando como referencia el diagrama de flujo de la Figura 11.

Tabla 130. Descripción de las pautas establecidas para los diferentes escenarios.

Escenarios	Información de control	Pautas para aplicar la imputación o asignación			
E – Completa en (b)	Pertenece a los nodos y a los procesos	Aplicar el modelo de decisión sin necesidad de la imputación de datos en (d).	-----	-----	-----
E1 – Incompleta en (e)	Información de control pertenece a los nodos	Primero se analiza si se cuenta con información histórica en (c), y si la información de control es de los nodos en (f), se aplica la imputación con K-Means en (g).	La segunda opción es, si no se cuenta con información histórica en (c), se realiza la asignación de un valor por defecto (1) en (k) (se considera cuando se tiene más de un criterio de ponderación en (j), en función de la carga computacional informada y las características del nodo.	La tercera opción es, si no se cuenta con información histórica en (c), se realiza la asignación de un valor por defecto (2) en (l) (se considera cuando se tiene un criterio de ponderación en (j), en función de las características del nodo.	-----
	Información de control pertenece a los procesos	Primero se analiza si se cuenta con información histórica en (c), si la información de control es de los procesos en (f), si los registros históricos son de $r_{ij} p_{kl}$ en (h), se aplica la imputación con K-Means en (g).	La segunda opción si se cuenta con información histórica en (c), si la información de control es de los procesos en (f), pero no se tienen registros históricos de $r_{ij} p_{kl}$ en (h), utilizaremos los registros históricos de $r_{xx} p_{kl}$ en (i), para aplicar la imputación con K-Means en (g).	La tercera opción es, si no se cuenta con registros históricos de $r_{ij} p_{kl}$ en (h), o de $r_{xx} p_{kl}$ en (i), se realiza la asignación de un valor por defecto (1) en (k), en función de la carga computacional informada y las características del nodo.	La cuarta opción es, si no se cuenta con registros históricos de $r_{ij} p_{kl}$ en (h), o de $r_{xx} p_{kl}$ en (i), se realiza la asignación de un valor por defecto (2) en (l), en función de las características del nodo.
E2 – Inexistente en (e)	Información de control pertenece a los nodos	Se analiza si se cuenta con información histórica en (b), y si la información de control es de los nodos en (m), se puede realizar la imputación con Media Ponderadas en (d).	La segunda opción es, si no se cuenta con información histórica en (c), se realiza la asignación de un valor por defecto (1) en (k), en función de la carga computacional informada y las características del nodo.	La tercera opción es, si no se cuenta con información histórica en (c), se realiza la asignación de un valor por defecto (2) en (l), en función de las características del nodo.	-----
	Información de control pertenece a los procesos	Primero se analiza si se cuenta con información histórica en (c), si la información de control es de los procesos en (m), si los registros históricos son de $r_{ij} p_{kl}$ en (o), se aplica la imputación con Medias Ponderadas en (d).	La segunda opción es, si no se cuenta con registros históricos de $r_{ij} p_{kl}$ en (h), o de $r_{xx} p_{kl}$ en (i), se realiza la asignación de un valor por defecto (1) en (k), en función de la carga computacional informada y las características del nodo.	La tercera opción es, si no se cuenta con registros históricos de $r_{ij} p_{kl}$ en (h), o de $r_{xx} p_{kl}$ en (i), se realiza la asignación de un valor por defecto (2) en (l), en función de las características del nodo.	-----

Fuente: Elaboración propia.

4.6. Consideraciones del modelo de decisión propuesto

En el modelo propuesto, en cada nodo se define una interfaz entre las aplicaciones y el sistema operativo, que a través de un runtime (software en tiempo de ejecución complementario del sistema operativo) incluido en esa interfaz, gestiona los procesos y recursos compartidos y define el escenario correspondiente. Además, los runtime interactúan entre sí para intercambiar información y existe un runtime coordinador global en uno de los nodos que evalúa y ejecuta el modelo de decisión y el operador de agregación correspondiente, se agrega la capa de imputación/asignación previa a la gestión de recursos y procesos, para disponer con la información de control completa en caso de ser necesario, como se indica en la Figura 12.

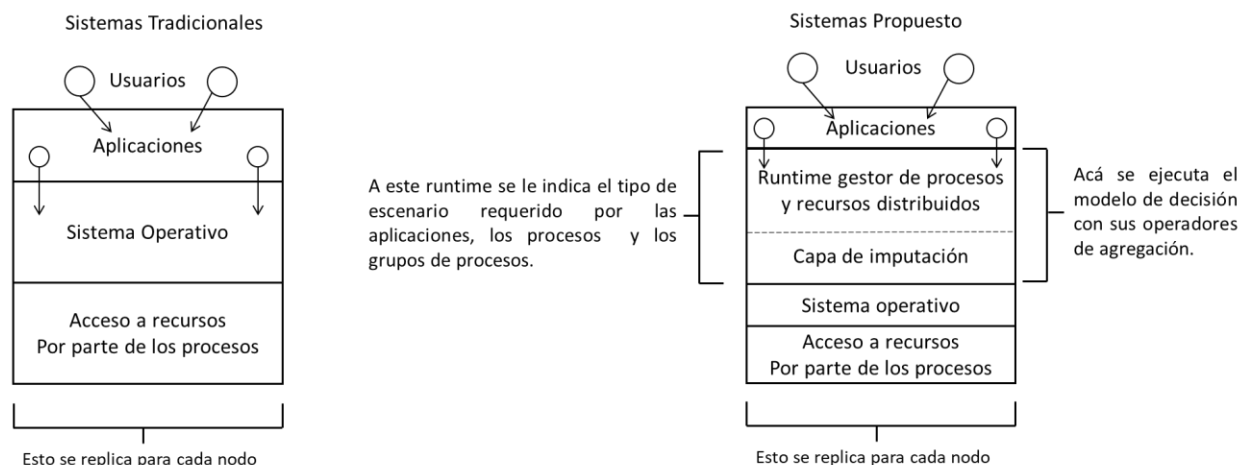


Figura 12. Comparación del Modelo propuesto y los modelos tradicionales

Fuente: Elaboración propia

Este método innovador considera diferentes opciones para aplicar la imputación/asignación para resolver la problemática de valores faltantes en la información de control, utilizando diferentes mecanismos para distintos escenarios (mencionados en el **Capítulo II** y **Capítulo III**), utilizando información de control de ciclos de asignaciones anteriores (disponibles en el nodo central para la imputación), las prestaciones de los nodos y el promedio de la carga computacional informada.

4.7. Ejemplo de modelo de decisión aplicado a alguno de los algoritmos tradicionales

En este trabajo se pretende que en el modelo de decisión propuesto se pueda observar un caso particular de alguno de los métodos que son considerados tradicionales, considerando el valor faltante en el criterio prioridad de proceso, utilizado para calcular de las prioridades o preferencias de los procesos.

El cálculo se realizará únicamente con procesos independientes, por tal motivo se inhabilita la columna que considera si un proceso integra un grupo de procesos, como se puede observar en la Tabla 131.

Tabla 131. Pesos asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos (método tradicional).

Procesos	Pasos Finales	
	Si integra un grupo de procesos	Si es independiente
p_{11}	-	$w_{f11}=1/np$
...	-	...
p_{kl}	-	$w_{kl}=1/np$
...	-	...
p_{np}	-	$w_{fnp}=1/np$

Fuente: Elaboración propia.

Los métodos considerados tradicionales no consideran la totalidad de los criterios establecidos en el modelo propuesto, generalmente sólo utilizan el criterio “**Prioridad Proceso**”, por lo tanto, solo se tendrá en cuenta para el criterio en cuestión para calcular la prioridad global, cuyas valoraciones de los pesos puede observarse en la Tabla 132 donde han sido inhabilitados los demás criterios.

Tabla 132. Pesos asignados a los criterios para calcular la prioridad o preferencia que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo - (método tradicional).

Categorías	Pesos							
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
Alta	-	-	-	-	0,100	-	-	-
Media	-	-	-	-	0,200	-	-	-
Baja	-	-	-	-	0,100	-	-	-

Fuente: Elaboración propia.

Estas prioridades, al igual que el **Capítulo II y Capítulo III**, se calculan en cada nodo para cada solicitud de recursos originada en cada proceso; el cálculo considera el vector de peso correspondiente según la carga actual del nodo y el vector de los valores concedidos por el nodo según los criterios de evaluación de la solicitud.

Puede darse la situación de que, en un determinado ciclo de recolección de información de control, el nodo central no reciba la totalidad de información de control del criterio ***“Prioridad Proceso”***, que en los métodos tradicionales son los valores que permiten calcular las prioridades o preferencias de los procesos, entonces se aplica el mecanismo de imputación/asignación teniendo en cuenta para la variable dos pautas.

- Primera pauta; se analiza si se cuenta con información histórica, y si corresponde a la misma relación recurso-proceso, de ser afirmativo se puede realizar la imputación con la *Media Ponderada*, teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación recurso-proceso, ponderando con mayor peso w a la información más reciente y con menor peso w a los más antiguos, se obtiene entonces, los valores faltantes del criterio ***“Prioridad Proceso”***.
- Segunda pauta; si no se cuenta con información histórica para la misma relación recurso-proceso, pero si se cuenta con información histórica de proceso con otros recursos, se

realiza la imputación con la *Media Ponderada*, teniendo en cuenta un histórico de valores correspondientes a cierta cantidad de ciclos de recolección de información de la relación recurso-proceso, ponderando con mayor peso w a la información más reciente y con menor peso w a los más antiguos, se obtiene entonces, los valores faltantes del criterio “***Prioridad Proceso***”.

- Tercera pauta; si no se cuenta con información histórica específica de la relación recurso-proceso y al no considerar para esta clasificación información histórica del proceso con relación a otros recursos, se realizará la **asignación de un valor por defecto**, en función de las características del nodo.

Para un ciclo de recolección de información, que para este ejemplo son utilizado la información de control y resultados de (La Red Martínez, 2017) y (Agostini, 2019), donde no consideran valores faltantes en la información de los criterios establecidos para calcular las prioridades o preferencias de los procesos para cada nodo, y teniendo en cuenta que para los métodos tradicionales no se consideran la totalidad de los criterios establecidos en los trabajos mencionados, sólo se tendrán en cuenta en la prioridad, el criterio “***Prioridad Proceso***”.

A partir de esos datos completos del criterio mencionado se aplica la amputación para simular faltantes de datos de control, seleccionando para este ejemplo el mecanismo MCAR parametrizado en un 10%, para finalmente imputarlos y verificar si los resultados obtenidos con los métodos de imputación/asignación funcionan adecuadamente. En una situación real, no se hace amputación, sólo imputación de datos faltantes.

Obtenidos los valores faltantes en algunos criterios, se aplica el método de imputación/asignación teniendo en cuenta las pautas establecidas anteriormente. Los valores

faltantes de los criterios para la relación proceso-recurso se pueden observar en la Tabla 133.

Tabla 133. Valoraciones asignadas a los criterios para calcular la prioridad, amputados con el mecanismo MCAR al 10% - (método tradicional).

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{11}$	-	-	-	-	0,80	-	-	-
$p_{11}r_{12}$	-	-	-	-	0,30	-	-	-
$p_{11}r_{21}$	-	-	-	-	0,90	-	-	-
$p_{11}r_{22}$	-	-	-	-	0,80	-	-	-
$p_{11}r_{23}$	-	-	-	-	X	-	-	-
$p_{11}r_{24}$	-	-	-	-	0,60	-	-	-
$p_{12}r_{11}$	-	-	-	-	1,00	-	-	-
$p_{12}r_{12}$	-	-	-	-	0,80	-	-	-
$p_{12}r_{21}$	-	-	-	-	X	-	-	-
$p_{12}r_{22}$	-	-	-	-	0,80	-	-	-
$p_{12}r_{31}$	-	-	-	-	0,30	-	-	-
$p_{12}r_{33}$	-	-	-	-	0,30	-	-	-
$p_{13}r_{11}$	-	-	-	-	X	-	-	-
$p_{13}r_{12}$	-	-	-	-	0,90	-	-	-
$p_{13}r_{13}$	-	-	-	-	0,90	-	-	-
$p_{13}r_{21}$	-	-	-	-	0,50	-	-	-
$p_{13}r_{22}$	-	-	-	-	0,50	-	-	-
$p_{13}r_{31}$	-	-	-	-	0,80	-	-	-
$p_{13}r_{32}$	-	-	-	-	0,40	-	-	-
$p_{13}r_{33}$	-	-	-	-	0,90	-	-	-
$p_{21}r_{12}$	-	-	-	-	0,70	-	-	-
$p_{21}r_{13}$	-	-	-	-	0,70	-	-	-
$p_{21}r_{22}$	-	-	-	-	0,50	-	-	-
$p_{21}r_{23}$	-	-	-	-	0,50	-	-	-
$p_{21}r_{31}$	-	-	-	-	0,90	-	-	-
$p_{21}r_{33}$	-	-	-	-	0,40	-	-	-
$p_{22}r_{21}$	-	-	-	-	0,60	-	-	-
$p_{22}r_{22}$	-	-	-	-	X	-	-	-
$p_{22}r_{31}$	-	-	-	-	0,40	-	-	-
$p_{22}r_{33}$	-	-	-	-	0,80	-	-	-
$p_{23}r_{12}$	-	-	-	-	0,50	-	-	-
$p_{23}r_{24}$	-	-	-	-	0,30	-	-	-
$p_{23}r_{31}$	-	-	-	-	0,70	-	-	-
$p_{23}r_{32}$	-	-	-	-	X	-	-	-
$p_{23}r_{33}$	-	-	-	-	0,90	-	-	-
$p_{24}r_{11}$	-	-	-	-	0,50	-	-	-
$p_{24}r_{12}$	-	-	-	-	X	-	-	-
$p_{24}r_{23}$	-	-	-	-	0,80	-	-	-
$p_{24}r_{24}$	-	-	-	-	0,30	-	-	-
$p_{25}r_{21}$	-	-	-	-	X	-	-	-
$p_{31}r_{13}$	-	-	-	-	0,70	-	-	-

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p</i> ₃₁ <i>r</i> ₃₁	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₁ <i>r</i> ₃₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₂ <i>r</i> ₁₁	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₂ <i>r</i> ₁₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₂ <i>r</i> ₁₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₂ <i>r</i> ₂₃	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₃ <i>r</i> ₁₁	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₃ <i>r</i> ₁₂	-	-	-	-	0,30	-	-	-
<i>p</i> ₃₃ <i>r</i> ₁₃	-	-	-	-	X	-	-	-
<i>p</i> ₃₃ <i>r</i> ₂₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₃ <i>r</i> ₂₂	-	-	-	-	X	-	-	-
<i>p</i> ₃₃ <i>r</i> ₂₃	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₃ <i>r</i> ₃₂	-	-	-	-	0,40	-	-	-
<i>p</i> ₃₃ <i>r</i> ₃₃	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₄ <i>r</i> ₁₂	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₄ <i>r</i> ₁₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₄ <i>r</i> ₂₂	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₄ <i>r</i> ₂₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₄ <i>r</i> ₂₄	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₄ <i>r</i> ₃₁	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₄ <i>r</i> ₃₂	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₄ <i>r</i> ₃₃	-	-	-	-	X	-	-	-
<i>p</i> ₃₅ <i>r</i> ₁₂	-	-	-	-	X	-	-	-
<i>p</i> ₃₅ <i>r</i> ₁₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₅ <i>r</i> ₂₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₅ <i>r</i> ₂₄	-	-	-	-	X	-	-	-
<i>p</i> ₃₅ <i>r</i> ₃₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₅ <i>r</i> ₃₂	-	-	-	-	0,40	-	-	-
<i>p</i> ₃₅ <i>r</i> ₃₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₁₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₁₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₁₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₂₁	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₆ <i>r</i> ₂₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₂₄	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₃₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₃₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>r</i> ₃₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₇ <i>r</i> ₁₁	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₇ <i>r</i> ₁₂	-	-	-	-	0,50	-	-	-
<i>p</i> ₃₇ <i>r</i> ₂₁	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₇ <i>r</i> ₃₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₇ <i>r</i> ₃₃	-	-	-	-	0,80	-	-	-

Fuente: Elaboración propia.

Imputación de las valoraciones faltantes de los criterios de prioridad o preferencia de los procesos - (método tradicional)

Se procede a imputar o asignar dichos valores para completar los datos faltantes, teniendo en cuenta las pautas definidas anteriormente.

Al igual que los **Capítulo II** y **Capítulo III**, cuando se asigne un valor por defecto se tendrán en cuenta los siguientes detalles de implementación:

- Nodo de prestaciones altas se asume un valor de 0,30.
- Nodo de prestaciones medias se asume un valor de 0,50.
- Nodo de prestaciones bajas se asume un valor 0,70.

Después de haber analizado la información con que se cuenta sobre las variables de estados principales de los nodos donde existen valores faltantes, y teniendo en cuenta las pautas para la imputación o asignación de valores, la relación de proceso-recurso queda clasificada como se muestra en la Tabla 134.

Tabla 134. Clasificación de valores faltantes de acuerdo a las pautas establecidas, para calcular la prioridad o preferencia de los procesos - (método tradicional).

Valor Faltante	Pauta 1: Existe información histórica misa relación proceso-recurso		Pauta 2: Existe información histórica de proceso con otros recursos.		Pauta 3: No hay información histórica.
	Cant. Registro histórico $p_{kl} r_{ij}$	Aplica	Cant. Registro histórico $p_{kl} r_{xx}$	Aplica	
$p_{11}r_{23}$	10	SI	---	---	---
$p_{12}r_{21}$	6	SI	---	---	---
$p_{13}r_{11}$	9	SI	---	---	---
$p_{22}r_{22}$	3	SI	---	---	---
$p_{23}r_{32}$	8	SI	---	---	---
$p_{24}r_{12}$	10	SI	---	---	---
$p_{25}r_{21}$	2	SI	---	---	---
$p_{33}r_{13}$	7	SI	---	---	---
$p_{33}r_{22}$	5	SI	---	---	---
$p_{34}r_{33}$	4	SI	---	---	---
$p_{35}r_{12}$	1	SI	---	---	---
$p_{35}r_{24}$	0	NO	0	NO	SI

Fuente: Elaboración propia.

Es importante que se entienda que al tener información histórica de la misma relación proceso-recurso, se aplica la imputación teniendo en cuenta la pauta 1, únicamente de no tener dicha información se tendrán en cuenta la información de proceso con otros recursos para así aplicar la pauta 2, y de no tener ambas informaciones recién se puede aplicar la pauta 3.

- **$p_{11}r_{23}$** : como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *diez registros*, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 135.

Tabla 135. Valores históricos de la relación proceso-recurso **$p_{11}r_{23}$** , últimos cinco registros de diez - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{23}$	-	-	-	-	1,00	-	-	-
$p_{11}r_{23}$	-	-	-	-	0,30	-	-	-
$p_{11}r_{23}$	-	-	-	-	0,70	-	-	-
$p_{11}r_{23}$	-	-	-	-	0,80	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a diez ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **$p_{11}r_{23}$** , como se puede ver en la Tabla 136.

Tabla 136. Valor imputado con la media ponderada en la relación proceso-recurso **$p_{11}r_{23}$** , aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}r_{23}$	-	-	-	-	0,58	-	-	-

Fuente: Elaboración propia.

- $p_{12}r_{21}$: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con seis registros, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 137.

Tabla 137. Valores históricos de la relación proceso-recurso $p_{12}r_{21}$, últimos cinco registros de seis (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{12}r_{21}$	-	-	-	-	0,70	-	-	-
$p_{12}r_{21}$	-	-	-	-	1,00	-	-	-
$p_{12}r_{21}$	-	-	-	-	0,80	-	-	-
$p_{12}r_{21}$	-	-	-	-	0,40	-	-	-
$p_{12}r_{21}$	-	-	-	-	0,90	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a seis ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{12}r_{21}$, como se puede ver en la Tabla 138.

Tabla 138. Valor imputado con la media ponderada en la relación proceso-recurso $p_{12}r_{21}$, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{12}r_{21}$	-	-	-	-	0,75	-	-	-

Fuente: Elaboración propia.

- $p_{13}r_{11}$: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con nueve registros, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 139.

Tabla 139. Valores históricos de la relación proceso-recurso $p_{13}r_{11}$, últimos cinco registros de nueve - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{13}r_{11}$	-	-	-	-	0,40	-	-	-
$p_{13}r_{11}$	-	-	-	-	0,60	-	-	-
$p_{13}r_{11}$	-	-	-	-	0,60	-	-	-
$p_{13}r_{11}$	-	-	-	-	0,30	-	-	-
$p_{13}r_{11}$	-	-	-	-	0,50	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a seis ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{13}r_{11}$, como se puede ver en la Tabla 140.

Tabla 140. Valor imputado con la media ponderada en la relación proceso-recurso $p_{13}r_{11}$, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{13}r_{11}$	-	-	-	-	0,46	-	-	-

Fuente: Elaboración propia.

- $p_{22}r_{22}$: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *tres registros*, las valoraciones se pueden apreciar en la Tabla 141, siendo el de la última fila el registro más reciente.

Tabla 141. Valores históricos de la relación proceso-recurso $p_{22}r_{22}$, tres únicos registros - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{22}r_{22}$	-	-	-	-	0,80	-	-	-
$p_{22}r_{22}$	-	-	-	-	0,50	-	-	-
$p_{22}r_{22}$	-	-	-	-	0,40	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a tres ciclos

de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{22}r_{22}$, como se puede ver en la Tabla 142.

Tabla 142. Valor imputado con la media ponderada en la relación proceso-recurso $p_{22}r_{22}$, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{22}r_{22}$	-	-	-	-	0,48	-	-	-

Fuente: Elaboración propia.

- $p_{23}r_{32}$: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con ocho registros, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 143.

Tabla 143. Valores históricos de la relación proceso-recurso $p_{23}r_{32}$, últimos cinco registros de ocho - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{23}r_{32}$	-	-	-	-	0,40	-	-	-
$p_{23}r_{32}$	-	-	-	-	0,70	-	-	-
$p_{23}r_{32}$	-	-	-	-	1,00	-	-	-
$p_{23}r_{32}$	-	-	-	-	0,70	-	-	-
$p_{23}r_{32}$	-	-	-	-	0,30	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a ocho ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{23}r_{32}$, como se puede ver en la Tabla 144.

Tabla 144. Valor imputado con la media ponderada en la relación proceso-recurso $p_{23}r_{32}$, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{23}r_{32}$	-	-	-	-	0,51	-	-	-

Fuente: Elaboración propia.

- **p_{24r12}** : como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con diez registros, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 145.

Tabla 145. Valores históricos de la relación proceso-recurso **p_{24r12}** , últimos cinco registros de diez - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	-	-	-	-	1,00	-	-	-
p_{24r12}	-	-	-	-	0,30	-	-	-
p_{24r12}	-	-	-	-	0,30	-	-	-
p_{24r12}	-	-	-	-	0,30	-	-	-
p_{24r12}	-	-	-	-	0,30	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a diez ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p_{24r12}** , como se puede ver en la Tabla 146.

Tabla 146. Valor imputado con la media ponderada en la relación proceso-recurso **p_{24r12}** , aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{24r12}	-	-	-	-	0,33	-	-	-

Fuente: Elaboración propia.

- **p_{25r21}** : como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *dos registros*, las valoraciones se pueden apreciar en la Tabla 147, siendo el de la última fila el registro más reciente.

Tabla 147. Valores históricos de la relación proceso-recurso $p_{25}r_{21}$, dos únicos registros - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{25}r_{21}$	-	-	-	-	0,30	-	-	-
$p_{25}r_{21}$	-	-	-	-	0,30	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a dos ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de $p_{25}r_{21}$, como se puede ver en la Tabla 148.

Tabla 148. Valor imputado con la media ponderada en la relación proceso-recurso $p_{25}r_{21}$, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{25}r_{21}$	-	-	-	-	0,30	-	-	-

Fuente: Elaboración propia.

- $p_{33}r_{13}$: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con *siete registros*, se puede apreciar en una tabla reducida las valoraciones de los cinco últimos registros, siendo el de la última fila el registro más reciente, como lo indica la Tabla 149.

Tabla 149. Valores históricos de la relación proceso-recurso $p_{33}r_{13}$, los cinco últimos registros de siete - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{33}r_{13}$	-	-	-	-	0,80	-	-	-
$p_{33}r_{13}$	-	-	-	-	0,80	-	-	-
$p_{33}r_{13}$	-	-	-	-	0,50	-	-	-
$p_{33}r_{13}$	-	-	-	-	1,00	-	-	-
$p_{33}r_{13}$	-	-	-	-	0,80	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a siete ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de

valores faltantes para los criterios de **p33r13**, como se puede ver en la Tabla 150.

Tabla 150. Valor imputado con la media ponderada en la relación proceso-recurso **p33r13**, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p33r13	-	-	-	-	0,81	-	-	-

Fuente: Elaboración propia.

- **p33r22**: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con cinco registros, se puede apreciar en la Tabla 151 las valoraciones de los cinco registros, siendo el de la última fila el registro más reciente.

Tabla 151. Valores históricos de la relación proceso-recurso **p33r22**, los cinco únicos registros - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p33r22	-	-	-	-	0,60	-	-	-
p33r22	-	-	-	-	0,90	-	-	-
p33r22	-	-	-	-	0,50	-	-	-
p33r22	-	-	-	-	0,30	-	-	-
p33r22	-	-	-	-	0,40	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a cinco ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de **p33r22**, como se puede ver en la Tabla 152.

Tabla 152. Valor imputado con la media ponderada en la relación proceso-recurso **p33r22**, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p33r22	-	-	-	-	0,41	-	-	-

Fuente: Elaboración propia.

- **p34r33**: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso

se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con cuatro registros, se puede apreciar en la Tabla 153 las valoraciones de los cuatro registros, siendo el de la última fila el registro más reciente.

Tabla 153. Valores históricos de la relación proceso-recurso *p₃₄r₃₃*, los cuatro únicos registros - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p₃₄r₃₃</i>	-	-	-	-	1,00	-	-	-
<i>p₃₄r₃₃</i>	-	-	-	-	0,60	-	-	-
<i>p₃₄r₃₃</i>	-	-	-	-	0,40	-	-	-
<i>p₃₄r₃₃</i>	-	-	-	-	0,50	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a cuatro ciclos de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de *p₃₄r₃₃*, como se puede ver en la Tabla 154.

Tabla 154. Valor imputado con la media ponderada en la relación proceso-recurso *p₃₄r₃₃*, aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p₃₄r₃₃</i>	-	-	-	-	0,66	-	-	-

Fuente: Elaboración propia.

- *p₃₅r₁₂*: como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso se cuenta con información histórica, esto permite hacer la imputación con la media ponderada como lo establece la primera pauta. Para este ejemplo, la tabla de valores históricos cuenta con un registro, se puede apreciar en la Tabla 155 las valoraciones del único registro.

Tabla 155. Valor histórico de la relación proceso-recurso p_{35r12} , registro único - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{35r12}	-	-	-	-	0,40	-	-	-

Fuente: Elaboración propia.

Aplicando la imputación de datos con un histórico de valores correspondientes a un ciclo de recolección de información de la relación de proceso-recurso, se obtienen la totalidad de valores faltantes para los criterios de p_{35r12} , como se puede ver en la Tabla 156.

Tabla 156. Valor imputado con la media ponderada en la relación proceso-recurso p_{35r12} , aplicando la pauta uno - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{35r12}	-	-	-	-	0,40	-	-	-

Fuente: Elaboración propia.

- p_{35r24} : como se muestra en la clasificación de la Tabla 134, para esta relación de proceso-recurso no se cuenta con información histórica, por ello se aplica la segunda pauta asumiendo un valor por defecto en función a las características del nodo, determinando las prestaciones del nodo en base a las características conocidas en la Tabla 18, del **Capítulo II**. Se determina que el nodo 3 es de prestaciones Alta, el valor **0,70** es asignado a cada criterio, como se puede observar en la Tabla 157.

Tabla 157. Valor asignado según las prestaciones del nodo para los criterios de p_{35r24} , aplicando la pauta dos - (método tradicional).

Proceso Recurso	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
p_{35r24}	-	-	-	-	0,70	-	-	-

Fuente: Elaboración propia.

Se obtienen los valores de los criterios mediante la imputación o asignación, como se observa en la Tabla 158, que serán utilizados para calcular la prioridad o preferencia de procesos para cada nodo.

Tabla 158. Valoraciones asignadas a los criterios para calcular la prioridad nodal, con valores imputados y asignado en algunas relaciones de proceso-recurso - (método tradicional).

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
$p_{11}f_{11}$	-	-	-	-	0,80	-	-	-
$p_{11}f_{12}$	-	-	-	-	0,30	-	-	-
$p_{11}f_{21}$	-	-	-	-	0,90	-	-	-
$p_{11}f_{22}$	-	-	-	-	0,80	-	-	-
$p_{11}f_{23}$	-	-	-	-	0,58	-	-	-
$p_{11}f_{24}$	-	-	-	-	0,60	-	-	-
$p_{12}f_{11}$	-	-	-	-	1,00	-	-	-
$p_{12}f_{12}$	-	-	-	-	0,8	-	-	-
$p_{12}f_{21}$	-	-	-	-	0,75	-	-	-
$p_{12}f_{22}$	-	-	-	-	0,80	-	-	-
$p_{12}f_{31}$	-	-	-	-	0,30	-	-	-
$p_{12}f_{33}$	-	-	-	-	0,30	-	-	-
$p_{13}f_{11}$	-	-	-	-	0,46	-	-	-
$p_{13}f_{12}$	-	-	-	-	0,90	-	-	-
$p_{13}f_{13}$	-	-	-	-	0,90	-	-	-
$p_{13}f_{21}$	-	-	-	-	0,50	-	-	-
$p_{13}f_{22}$	-	-	-	-	0,50	-	-	-
$p_{13}f_{31}$	-	-	-	-	0,80	-	-	-
$p_{13}f_{32}$	-	-	-	-	0,40	-	-	-
$p_{13}f_{33}$	-	-	-	-	0,90	-	-	-
$p_{21}f_{12}$	-	-	-	-	0,70	-	-	-
$p_{21}f_{13}$	-	-	-	-	0,70	-	-	-
$p_{21}f_{22}$	-	-	-	-	0,50	-	-	-
$p_{21}f_{23}$	-	-	-	-	0,50	-	-	-
$p_{21}f_{31}$	-	-	-	-	0,90	-	-	-
$p_{21}f_{33}$	-	-	-	-	0,40	-	-	-
$p_{22}f_{21}$	-	-	-	-	0,60	-	-	-
$p_{22}f_{22}$	-	-	-	-	0,48	-	-	-
$p_{22}f_{31}$	-	-	-	-	0,40	-	-	-
$p_{22}f_{33}$	-	-	-	-	0,80	-	-	-
$p_{23}f_{12}$	-	-	-	-	0,50	-	-	-
$p_{23}f_{24}$	-	-	-	-	0,30	-	-	-
$p_{23}f_{31}$	-	-	-	-	0,70	-	-	-
$p_{23}f_{32}$	-	-	-	-	0,51	-	-	-
$p_{23}f_{33}$	-	-	-	-	0,90	-	-	-
$p_{24}f_{11}$	-	-	-	-	0,50	-	-	-
$p_{24}f_{12}$	-	-	-	-	0,33	-	-	-
$p_{24}f_{23}$	-	-	-	-	0,80	-	-	-
$p_{24}f_{24}$	-	-	-	-	0,30	-	-	-
$p_{25}f_{21}$	-	-	-	-	0,30	-	-	-
$p_{31}f_{13}$	-	-	-	-	0,70	-	-	-
$p_{31}f_{31}$	-	-	-	-	0,70	-	-	-

Procesos Recursos	Criterios							
	Nº Proc.	% CPU	% Mem.	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S
<i>p</i> ₃₁ <i>F</i> ₃₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₂ <i>F</i> ₁₁	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₂ <i>F</i> ₁₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₂ <i>F</i> ₁₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₂ <i>F</i> ₂₃	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₃ <i>F</i> ₁₁	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₃ <i>F</i> ₁₂	-	-	-	-	0,30	-	-	-
<i>p</i> ₃₃ <i>F</i> ₁₃	-	-	-	-	0,81	-	-	-
<i>p</i> ₃₃ <i>F</i> ₂₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₃ <i>F</i> ₂₂	-	-	-	-	0,41	-	-	-
<i>p</i> ₃₃ <i>F</i> ₂₃	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₃ <i>F</i> ₃₂	-	-	-	-	0,40	-	-	-
<i>p</i> ₃₃ <i>F</i> ₃₃	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₄ <i>F</i> ₁₂	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₄ <i>F</i> ₁₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₄ <i>F</i> ₂₂	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₄ <i>F</i> ₂₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₄ <i>F</i> ₂₄	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₄ <i>F</i> ₃₁	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₄ <i>F</i> ₃₂	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₄ <i>F</i> ₃₃	-	-	-	-	0,66	-	-	-
<i>p</i> ₃₅ <i>F</i> ₁₂	-	-	-	-	0,40	-	-	-
<i>p</i> ₃₅ <i>F</i> ₁₃	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₅ <i>F</i> ₂₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₅ <i>F</i> ₂₄	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₅ <i>F</i> ₃₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₅ <i>F</i> ₃₂	-	-	-	-	0,40	-	-	-
<i>p</i> ₃₅ <i>F</i> ₃₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₁₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₁₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₁₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₂₁	-	-	-	-	0,70	-	-	-
<i>p</i> ₃₆ <i>F</i> ₂₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₂₄	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₃₁	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₃₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₆ <i>F</i> ₃₃	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₇ <i>F</i> ₁₁	-	-	-	-	0,90	-	-	-
<i>p</i> ₃₇ <i>F</i> ₁₂	-	-	-	-	0,50	-	-	-
<i>p</i> ₃₇ <i>F</i> ₂₁	-	-	-	-	0,60	-	-	-
<i>p</i> ₃₇ <i>F</i> ₃₂	-	-	-	-	0,80	-	-	-
<i>p</i> ₃₇ <i>F</i> ₃₃	-	-	-	-	0,80	-	-	-

Fuente: Elaboración propia.

Se obtienen los valores para el criterio “**Prioridad Proceso**” (donde se presentaban valores

faltantes) mediante la imputación o asignación, que serán utilizados para calcular la prioridad o preferencia de procesos para cada nodo. Se multiplica escalarmente cada vector de valoraciones de cada requerimiento por el vector de pesos correspondiente a la categoría de carga actual del nodo para obtener la prioridad según cada criterio y la prioridad nodal otorgada a cada requerimiento; esto se puede observar en la Tabla 159.

Si bien el modelo de decisión obtiene la información de todos los criterios, cabe destacar que, para los métodos tradicionales, el vector de peso de la Tabla 162 y Tabla 163, sólo va a influir en la prioridad o preferencia de los procesos respecto del nodo, y para el criterio **“Prioridad Proceso”**.

Tabla 159. Valoraciones asignadas a los criterios para calcular la prioridad o preferencia nodal que cada nodo otorgará a cada requerimiento de cada proceso según la carga del nodo, con valores imputados/asignados - (método tradicional).

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
$p_{11}r_{11}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{11}r_{11})$	-	-	-	-	0,1600	-	-	-	
$p_{11}r_{12}$	-	-	-	-	0,3000	-	-	-	0,0300
$Pri(p_{11}r_{12})$	-	-	-	-	0,0600	-	-	-	
$p_{11}r_{21}$	-	-	-	-	0,9000	-	-	-	0,0900
$Pri(p_{11}r_{21})$	-	-	-	-	0,1800	-	-	-	
$p_{11}r_{22}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{11}r_{22})$	-	-	-	-	0,1600	-	-	-	
$p_{11}r_{23}$	-	-	-	-	0,5800	-	-	-	0,0580
$Pri(p_{11}r_{23})$	-	-	-	-	0,1160	-	-	-	
$p_{11}r_{24}$	-	-	-	-	0,6000	-	-	-	0,0600
$Pri(p_{11}r_{24})$	-	-	-	-	0,1200	-	-	-	
$p_{12}r_{11}$	-	-	-	-	1,0000	-	-	-	0,1000
$Pri(p_{12}r_{11})$	-	-	-	-	0,2000	-	-	-	
$p_{12}r_{12}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{12}r_{12})$	-	-	-	-	0,1600	-	-	-	
$p_{12}r_{21}$	-	-	-	-	0,7521	-	-	-	0,0752
$Pri(p_{12}r_{21})$	-	-	-	-	0,1504	-	-	-	
$p_{12}r_{22}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{12}r_{22})$	-	-	-	-	0,1600	-	-	-	
$p_{12}r_{31}$	-	-	-	-	0,3000	-	-	-	

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p_{12r31})</i>	-	-	-	-	0,0600	-	-	-	0,0300
<i>p_{12r33}</i>	-	-	-	-	0,3000	-	-	-	
<i>Pri(p_{12r33})</i>	-	-	-	-	0,0600	-	-	-	0,0300
<i>p_{13r11}</i>	-	-	-	-	0,4646	-	-	-	
<i>Pri(p_{13r11})</i>	-	-	-	-	0,0929	-	-	-	0,0465
<i>p_{13r12}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{13r12})</i>	-	-	-	-	0,1800	-	-	-	0,0900
<i>p_{13r13}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{13r13})</i>	-	-	-	-	0,1800	-	-	-	0,0900
<i>p_{13r21}</i>	-	-	-	-	0,5000	-	-	-	
<i>Pri(p_{13r21})</i>	-	-	-	-	0,1000	-	-	-	0,0500
<i>p_{13r22}</i>	-	-	-	-	0,5000	-	-	-	
<i>Pri(p_{13r22})</i>	-	-	-	-	0,1000	-	-	-	0,0500
<i>p_{13r31}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{13r31})</i>	-	-	-	-	0,1600	-	-	-	0,0800
<i>p_{13r32}</i>	-	-	-	-	0,4000	-	-	-	
<i>Pri(p_{13r32})</i>	-	-	-	-	0,0800	-	-	-	0,0400
<i>p_{13r33}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{13r33})</i>	-	-	-	-	0,1800	-	-	-	0,0900
<i>p_{21r12}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{21r12})</i>	-	-	-	-	0,1400	-	-	-	0,1400
<i>p_{21r13}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{21r13})</i>	-	-	-	-	0,1400	-	-	-	0,1400
<i>p_{21r22}</i>	-	-	-	-	0,5000	-	-	-	
<i>Pri(p_{21r22})</i>	-	-	-	-	0,1000	-	-	-	0,1000
<i>p_{21r23}</i>	-	-	-	-	0,5000	-	-	-	
<i>Pri(p_{21r23})</i>	-	-	-	-	0,1000	-	-	-	0,1000
<i>p_{21r31}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{21r31})</i>	-	-	-	-	0,1800	-	-	-	0,1800
<i>p_{21r33}</i>	-	-	-	-	0,4000	-	-	-	
<i>Pri(p_{21r33})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{22r21}</i>	-	-	-	-	0,6000	-	-	-	
<i>Pri(p_{22r21})</i>	-	-	-	-	0,1200	-	-	-	0,1200
<i>p_{22r22}</i>	-	-	-	-	0,4795	-	-	-	
<i>Pri(p_{22r22})</i>	-	-	-	-	0,0959	-	-	-	0,0959
<i>p_{22r31}</i>	-	-	-	-	0,4000	-	-	-	
<i>Pri(p_{22r31})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{22r33}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{22r33})</i>	-	-	-	-	0,1600	-	-	-	0,1600
<i>p_{23r12}</i>	-	-	-	-	0,5000	-	-	-	
<i>Pri(p_{23r12})</i>	-	-	-	-	0,1000	-	-	-	0,1000
<i>p_{23r24}</i>	-	-	-	-	0,3000	-	-	-	
<i>Pri(p_{23r24})</i>	-	-	-	-	0,0600	-	-	-	0,0600
<i>p_{23r31}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{23r31})</i>	-	-	-	-	0,1400	-	-	-	0,1400

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>p</i> _{23r32}	-	-	-	-	0,5143	-	-	-	
<i>Pri</i> (<i>p</i> _{23r32})	-	-	-	-	0,1029	-	-	-	0,1029
<i>p</i> _{23r33}	-	-	-	-	0,9000	-	-	-	
<i>Pri</i> (<i>p</i> _{23r33})	-	-	-	-	0,1800	-	-	-	0,1800
<i>p</i> _{24r11}	-	-	-	-	0,5000	-	-	-	
<i>Pri</i> (<i>p</i> _{24r11})	-	-	-	-	0,1000	-	-	-	0,1000
<i>p</i> _{24r12}	-	-	-	-	0,3310	-	-	-	
<i>Pri</i> (<i>p</i> _{24r12})	-	-	-	-	0,0662	-	-	-	0,0662
<i>p</i> _{24r23}	-	-	-	-	0,8000	-	-	-	
<i>Pri</i> (<i>p</i> _{24r23})	-	-	-	-	0,1600	-	-	-	0,1600
<i>p</i> _{24r24}	-	-	-	-	0,3000	-	-	-	
<i>Pri</i> (<i>p</i> _{24r24})	-	-	-	-	0,0600	-	-	-	0,0600
<i>p</i> _{25r21}	-	-	-	-	0,3000	-	-	-	
<i>Pri</i> (<i>p</i> _{25r21})	-	-	-	-	0,0600	-	-	-	0,0600
<i>p</i> _{31r13}	-	-	-	-	0,7000	-	-	-	
<i>Pri</i> (<i>p</i> _{31r13})	-	-	-	-	0,0700	-	-	-	0,0700
<i>p</i> _{31r31}	-	-	-	-	0,7000	-	-	-	
<i>Pri</i> (<i>p</i> _{31r31})	-	-	-	-	0,0700	-	-	-	0,0700
<i>p</i> _{31r33}	-	-	-	-	0,9000	-	-	-	
<i>Pri</i> (<i>p</i> _{31r33})	-	-	-	-	0,0900	-	-	-	0,0900
<i>p</i> _{32r11}	-	-	-	-	0,9000	-	-	-	
<i>Pri</i> (<i>p</i> _{32r11})	-	-	-	-	0,0900	-	-	-	0,0900
<i>p</i> _{32r12}	-	-	-	-	0,8000	-	-	-	
<i>Pri</i> (<i>p</i> _{32r12})	-	-	-	-	0,0800	-	-	-	0,0800
<i>p</i> _{32r13}	-	-	-	-	0,9000	-	-	-	
<i>Pri</i> (<i>p</i> _{32r13})	-	-	-	-	0,0900	-	-	-	0,0900
<i>p</i> _{32r23}	-	-	-	-	0,6000	-	-	-	
<i>Pri</i> (<i>p</i> _{32r23})	-	-	-	-	0,0600	-	-	-	0,0600
<i>p</i> _{33r11}	-	-	-	-	0,6000	-	-	-	
<i>Pri</i> (<i>p</i> _{33r11})	-	-	-	-	0,0600	-	-	-	0,0600
<i>p</i> _{33r12}	-	-	-	-	0,3000	-	-	-	
<i>Pri</i> (<i>p</i> _{33r12})	-	-	-	-	0,0300	-	-	-	0,0300
<i>p</i> _{33r13}	-	-	-	-	0,8134	-	-	-	
<i>Pri</i> (<i>p</i> _{33r13})	-	-	-	-	0,0813	-	-	-	0,0813
<i>p</i> _{33r21}	-	-	-	-	0,8000	-	-	-	
<i>Pri</i> (<i>p</i> _{33r21})	-	-	-	-	0,0800	-	-	-	0,0800
<i>p</i> _{33r22}	-	-	-	-	0,4147	-	-	-	
<i>Pri</i> (<i>p</i> _{33r22})	-	-	-	-	0,0415	-	-	-	0,0415
<i>p</i> _{33r23}	-	-	-	-	0,6000	-	-	-	
<i>Pri</i> (<i>p</i> _{33r23})	-	-	-	-	0,0600	-	-	-	0,0600
<i>p</i> _{33r32}	-	-	-	-	0,4000	-	-	-	
<i>Pri</i> (<i>p</i> _{33r32})	-	-	-	-	0,0400	-	-	-	0,0400
<i>p</i> _{33r33}	-	-	-	-	0,6000	-	-	-	
<i>Pri</i> (<i>p</i> _{33r33})	-	-	-	-	0,0600	-	-	-	0,0600
<i>p</i> _{34r12}	-	-	-	-	0,7000	-	-	-	

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
<i>Pri(p_{34r12})</i>	-	-	-	-	0,0700	-	-	-	0,0700
<i>p_{34r13}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{34r13})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{34r22}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{34r22})</i>	-	-	-	-	0,0900	-	-	-	0,0900
<i>p_{34r23}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{34r23})</i>	-	-	-	-	0,0900	-	-	-	0,0900
<i>p_{34r24}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{34r24})</i>	-	-	-	-	0,0700	-	-	-	0,0700
<i>p_{34r31}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{34r31})</i>	-	-	-	-	0,0700	-	-	-	0,0700
<i>p_{34r32}</i>	-	-	-	-	0,6000	-	-	-	
<i>Pri(p_{34r32})</i>	-	-	-	-	0,0600	-	-	-	0,0600
<i>p_{34r33}</i>	-	-	-	-	0,6581	-	-	-	
<i>Pri(p_{34r33})</i>	-	-	-	-	0,0658	-	-	-	0,0658
<i>p_{35r12}</i>	-	-	-	-	0,4000	-	-	-	
<i>Pri(p_{35r12})</i>	-	-	-	-	0,0400	-	-	-	0,0400
<i>p_{35r13}</i>	-	-	-	-	0,9000	-	-	-	
<i>Pri(p_{35r13})</i>	-	-	-	-	0,0900	-	-	-	0,0900
<i>p_{35r22}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{35r22})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{35r24}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{35r24})</i>	-	-	-	-	0,0700	-	-	-	0,0700
<i>p_{35r31}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{35r31})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{35r32}</i>	-	-	-	-	0,4000	-	-	-	
<i>Pri(p_{35r32})</i>	-	-	-	-	0,0400	-	-	-	0,0400
<i>p_{35r33}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{35r33})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r11}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r11})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r12}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r12})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r13}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r13})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r21}</i>	-	-	-	-	0,7000	-	-	-	
<i>Pri(p_{36r21})</i>	-	-	-	-	0,0700	-	-	-	0,0700
<i>p_{36r22}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r22})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r24}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r24})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r31}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r31})</i>	-	-	-	-	0,0800	-	-	-	0,0800
<i>p_{36r32}</i>	-	-	-	-	0,8000	-	-	-	
<i>Pri(p_{36r32})</i>	-	-	-	-	0,0800	-	-	-	0,0800

Procesos Recursos	Criterios								Prioridades Nodales
	Nº Proc.	% CPU	% Mem	% MV	Prioridad Proc.	Sobrec. Mem.	Sobrec. Proc.	Sobrec. E/S	
$p_{36}r_{33}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{36}r_{33})$	-	-	-	-	0,0800	-	-	-	
$p_{37}r_{11}$	-	-	-	-	0,9000	-	-	-	0,0900
$Pri(p_{37}r_{11})$	-	-	-	-	0,0900	-	-	-	
$p_{37}r_{12}$	-	-	-	-	0,5000	-	-	-	0,0500
$Pri(p_{37}r_{12})$	-	-	-	-	0,0500	-	-	-	
$p_{37}r_{21}$	-	-	-	-	0,6000	-	-	-	0,0600
$Pri(p_{37}r_{21})$	-	-	-	-	0,0600	-	-	-	
$p_{37}r_{32}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{37}r_{32})$	-	-	-	-	0,0800	-	-	-	
$p_{37}r_{33}$	-	-	-	-	0,8000	-	-	-	0,0800
$Pri(p_{37}r_{33})$	-	-	-	-	0,0800	-	-	-	

Fuente: Elaboración propia.

Nota: los valores marcados corresponden a resultados de imputación.

Se debe calcular las prioridades nodales, pesos finales y prioridades globales de los procesos para acceder a los recursos.

La Tabla 160 y la Tabla 161, se construyen a partir de los valores de prioridades nodales calculadas de la Tabla 159, que para este ejemplo coincide con el criterio “**Prioridad Proceso**”.

Tabla 160. Prioridades nodales de los procesos para acceder a cada recurso p_{11} a p_{25} - (método tradicional).

Recursos	Prioridades Nodales de los Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{11}	0,0800	0,1000	0,0465	-	-	-	0,1000	-
r_{12}	0,0300	0,0800	0,0900	0,1400	-	0,1000	0,0662	-
r_{13}	-	-	0,0900	0,1400	-	-	-	-
r_{21}	0,0900	0,0752	0,0500	-	0,1200	-	-	0,0600
r_{22}	0,0800	0,0800	0,0500	0,1000	0,0959	-	-	-
r_{23}	0,0580	-	-	0,1000	-	-	0,1600	-
r_{24}	0,0600	-	-	-	-	0,0600	0,0600	-
r_{31}	-	0,0300	0,0800	0,1800	0,0800	0,1400	-	-
r_{32}	-	-	0,0400	-	-	0,1029	-	-
r_{33}	-	0,0300	0,0900	0,0800	0,1600	0,1800	-	-

Fuente: Elaboración propia.

Nota: los valores marcados corresponden a resultados de imputación.

Tabla 161. Prioridades nodales de los procesos para acceder a cada recurso p_{31} a p_{37} - (método tradicional).

Recursos	Prioridades Nodales de los Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{11}	-	0,0900	0,0600	-	-	0,0800	0,0900
r_{12}	-	0,0800	0,0300	0,0700	0,0400	0,0800	0,0500
r_{13}	0,0700	0,0900	0,0813	0,0800	0,0900	0,0800	-
r_{21}	-	-	0,0800	-	-	0,0700	0,0600
r_{22}	-	-	0,0415	0,0900	0,0800	0,0800	-
r_{23}	-	0,0600	0,0600	0,0900	0,0700	-	-
r_{24}	-	-	-	0,0700	-	0,0800	-
r_{31}	0,0700	-	-	0,0700	0,0800	0,0800	-
r_{32}	-	-	0,0400	0,0600	0,0400	0,0800	0,0800
r_{33}	0,0900	-	0,0600	0,0658	0,0800	0,0800	0,0800

Fuente: Elaboración propia.

Nota: los valores marcados corresponden a resultados de imputación.

Posteriormente se calcula el vector de pesos finales que se utilizará en el proceso final de agregación para determinar el orden o prioridad de acceso a los recursos; esto se muestran en la Tabla 162 y en la Tabla 163.

Tabla 162. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, de p_{11} al p_{25} - (método tradicional).

Pesos	Procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
Finales	0,067	0,067	0,067	0,067	0,067	0,067	0,067	0,067
Finales Normalizados	0,067	0,067	0,067	0,067	0,067	0,067	0,067	0,067

Fuente: Elaboración propia.

Tabla 163. Pesos finales y pesos finales normalizados asignados a los procesos para calcular la prioridad o preferencia final de acceso a los recursos, de p_{31} al p_{37} - (método tradicional).

Pesos	Procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
Finales	0,067	0,067	0,067	0,067	0,067	0,067	0,067
Finales Normalizados	0,067	0,067	0,067	0,067	0,067	0,067	0,067

Fuente: Elaboración propia.

Como se había mencionado, los métodos tradicionales no consideran los grupos de procesos, el cálculo sólo tendrá en cuenta que los procesos son independientes. Dado que en el ejemplo existen 15 procesos, y que el cálculo de los pesos es wf_{ij} igual a $1/np$ para procesos independientes, donde np es el número de procesos en el sistema (15), el cálculo para los pesos de cada proceso (wp_{ij}) es igual a $1/15$. Para el cálculo de los pesos normalizados (nwp_{ij}) se divide cada valor de wp_{ij} por la sumatoria de todos los wp_{ij} . Se debe calcular el vector de peso final que se utilizará en el proceso de agregación final para determinar el orden o la prioridad de acceso a los recursos. Además, los pesos recientemente obtenidos deberán normalizarse dividiendo cada uno de ellos por la suma de todos ellos, esto se pueden observar en la Tabla 162 y en la Tabla 163.

Las prioridades nodales indicadas en la Tabla 160 y Tabla 161 tomadas fila por fila, es decir, respecto de cada recurso, se multiplicarán escalarmente por el vector de pesos finales normalizados indicado en la Tabla 162 y en la Tabla 163, para obtener las prioridades globales finales de acceso de cada proceso a cada recurso y de allí, el orden o prioridad con que se asignarán los recursos y a qué proceso se asignará cada uno de ellos; esto se indica en la Tabla 164 y Tabla 165.

Tabla 164. Prioridades globales finales de los procesos para acceder a cada recurso, de p_{11} al p_{25} - (método tradicional).

Recursos	Prioridades globales finales de los procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{11}	0,0054	0,0067	0,0031	-	-	-	0,0067	-
r_{12}	0,0020	0,0054	0,0060	0,0094	-	0,0067	0,0044	-
r_{13}	-	-	0,0060	0,0094	-	-	-	-
r_{21}	0,0060	0,0050	0,0034	-	0,0080	-	-	0,0040
r_{22}	0,0054	0,0054	0,0034	0,0067	0,0064	-	-	-
r_{23}	0,0039	-	-	0,0067	-	-	0,0107	-

Recursos	Prioridades globales finales de los procesos							
	p_{11}	p_{12}	p_{13}	p_{21}	p_{22}	p_{23}	p_{24}	p_{25}
r_{24}	0,0040	-	-	-	-	0,0040	0,0040	-
r_{31}	-	0,0020	0,0054	0,0121	0,0054	0,0094	-	-
r_{32}	-	-	0,0027	-	-	0,0069	-	-
r_{33}	-	0,0020	0,0060	0,0054	0,0107	0,0121	-	-

Fuente: Elaboración propia.

Nota: los valores marcados corresponden a resultados de imputación.

Tabla 165. Prioridades globales finales de los procesos para acceder a cada recurso, de p_{31} al p_{37} - (método tradicional).

Recursos	Prioridades globales finales de los procesos						
	p_{31}	p_{32}	p_{33}	p_{34}	p_{35}	p_{36}	p_{37}
r_{11}	-	0,0060	0,0040	-	-	0,0054	0,0060
r_{12}	-	0,0054	0,0020	0,0047	0,0027	0,0054	0,0034
r_{13}	0,0047	0,0060	0,0054	0,0054	0,0060	0,0054	-
r_{21}	-	-	0,0054	-	-	0,0047	0,0040
r_{22}	-	-	0,0028	0,0060	0,0054	0,0054	-
r_{23}	-	0,0040	0,0040	0,0060	0,0047	-	-
r_{24}	-	-	-	0,0047	-	0,0054	-
r_{31}	0,0047	-	-	0,0047	0,0054	0,0054	-
r_{32}	-	-	0,0027	0,0040	0,0027	0,0054	0,0054
r_{33}	0,0060	-	0,0040	0,0044	0,0054	0,0054	0,0054

Fuente: Elaboración propia.

Nota: los valores marcados corresponden a resultados de imputación.

El mayor de estos productos hechos para los distintos procesos en relación al mismo recurso indicará cuál de los procesos tendrá acceso al recurso (en caso de empates se podría optar por dar ganador al proceso identificado con el número menor); esto se muestran en rojo en la Tabla 164 y Tabla 165.

La sumatoria de todos estos productos en relación al mismo recurso indicará la prioridad que tendrá dicho recurso para ser asignado. Esto constituye la Función de Asignación para Sistemas Distribuidos (**FASD**), que se muestra en la Tabla 166.

La suma de todos estos productos en relación con el mismo recurso indicará la prioridad que deberá asignarse a este recurso, en relación con los demás recursos que también deberán asignarse.

Tabla 166. Prioridades globales finales para asignar los recursos, constituye la Función de Asignación para Sistemas Distribuidos (FASD), del método tradicional (MT).

FASD	Prioridad Global Final para Asignar el Recurso	Proceso
r_{11}	0,0433	r_{11} al p_{12}
r_{12}	0,0574	r_{12} al p_{21}
r_{13}	0,0483	r_{13} al p_{21}
r_{21}	0,0405	r_{21} al p_{22}
r_{22}	0,0467	r_{22} al p_{21}
r_{23}	0,0354	r_{23} al p_{24}
r_{24}	0,0268	r_{24} al p_{36}
r_{31}	0,0543	r_{31} al p_{21}
r_{32}	0,0297	r_{32} al p_{23}
r_{33}	0,0667	r_{33} al p_{23}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El orden final de asignación de los recursos y los procesos destinatarios de los mismos se obtiene ordenando la (FASD), que se muestra en la Tabla 167.

Tabla 167. Prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso, constituye la Función de Asignación para Sistemas Distribuidos Ordenada (FASDO) - MT.

FASDO	Prioridad Global Final para Asignar el Recurso	Proceso
r_{33}	0,0667	r_{33} al p_{23}
r_{12}	0,0574	r_{12} al p_{21}
r_{31}	0,0543	r_{31} al p_{21}
r_{13}	0,0483	r_{13} al p_{21}
r_{22}	0,0467	r_{22} al p_{21}
r_{11}	0,0433	r_{11} al p_{12}
r_{21}	0,0405	r_{21} al p_{22}
r_{23}	0,0354	r_{23} al p_{24}
r_{32}	0,0297	r_{32} al p_{23}
r_{24}	0,0268	r_{24} al p_{36}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

El siguiente paso es reiterar el procedimiento retirando de las solicitudes de recursos las asignaciones ya hechas; también debe tenerse en cuenta que los recursos asignados quedarán disponibles cuando los procesos los hayan liberado, pudiendo por lo tanto ser asignados a otros procesos. El resultado de las sucesivas iteraciones finaliza una vez que se han atendido todas las solicitudes de recursos de todos los procesos respetando la exclusión mutua y las prioridades de los procesos, las prioridades nodales y las prioridades finales; el resultado de las sucesivas iteraciones se muestra en la Tabla 168, Tabla 169, Tabla 170, Tabla 171, Tabla 172, Tabla 173, Tabla 174, Tabla 175, Tabla 176, Tabla 177 y Tabla 178.

Tabla 168. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (segunda iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0547	r_{33} al p_{22}
0,0480	r_{12} al p_{23}
0,0422	r_{31} al p_{23}
0,0400	r_{22} al p_{22}
0,0389	r_{13} al p_{13}
0,0366	r_{11} al p_{24}
0,0325	r_{21} al p_{11}
0,0247	r_{23} al p_{21}
0,0228	r_{32} al p_{36}
0,0215	r_{24} al p_{34}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 169. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (tercera iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0439	r_{33} al p_{31}
0,0413	r_{12} al p_{13}
0,0336	r_{22} al p_{34}
0,0329	r_{13} al p_{32}
0,0328	r_{31} al p_{13}
0,0299	r_{11} al p_{32}
0,0265	r_{21} al p_{33}
0,0180	r_{23} al p_{34}
0,0174	r_{32} al p_{37}
0,0168	r_{24} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 170. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (cuarta iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0379	r_{33} al p_{13}
0,0353	r_{12} al p_{12}
0,0276	r_{22} al p_{11}
0,0275	r_{31} al p_{22}
0,0269	r_{13} al p_{35}
0,0239	r_{11} al p_{37}
0,0211	r_{21} al p_{12}
0,0121	r_{24} al p_{11}
0,0121	r_{32} al p_{34}
0,0119	r_{23} al p_{32}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 171. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (quinta iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0319	r_{33} al p_{21}
0,0299	r_{12} al p_{32}
0,0222	r_{22} al p_{12}
0,0221	r_{31} al p_{35}
0,0209	r_{13} al p_{33}
0,0179	r_{11} al p_{11}
0,0161	r_{21} al p_{36}
0,0080	r_{24} al p_{23}
0,0080	r_{32} al p_{13}
0,0079	r_{23} al p_{33}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 172. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (sexta iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0265	r_{33} al p_{35}
0,0245	r_{12} al p_{36}
0,0168	r_{22} al p_{35}
0,0168	r_{31} al p_{36}
0,0154	r_{13} al p_{34}
0,0125	r_{11} al p_{36}
0,0114	r_{21} al p_{25}
0,0054	r_{32} al p_{33}
0,0040	r_{24} al p_{24}
0,0039	r_{23} al p_{11}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 173. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (séptima iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0212	r_{33} al p_{36}
0,0192	r_{12} al p_{34}
0,0115	r_{22} al p_{36}
0,0114	r_{31} al p_{31}
0,0101	r_{13} al p_{36}
0,0074	r_{21} al p_{37}
0,0071	r_{11} al p_{33}
0,0027	r_{32} al p_{35}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 174. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (octava iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0158	r_{33} al p_{37}
0,0145	r_{12} al p_{24}
0,0067	r_{31} al p_{34}
0,0061	r_{22} al p_{13}
0,0047	r_{13} al p_{31}
0,0034	r_{21} al p_{13}
0,0031	r_{11} al p_{13}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 175. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (novena iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0104	r_{33} al p_{34}
0,0101	r_{12} al p_{37}
0,0028	r_{22} al p_{33}
0,0020	r_{31} al p_{12}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 176. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décima iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0067	r_{12} al p_{35}
0,0060	r_{33} al p_{33}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 177. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo primera iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0040	r_{12} al p_{11}
0,0020	r_{33} al p_{12}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

Tabla 178. Orden o prioridad final de asignación de los recursos y proceso al cual se asigna cada recurso (décimo segunda iteración) - MT.

Prioridad Global Final Ordenada	Asignación
0,0020	r_{12} al p_{33}

Fuente: Elaboración propia. Basado en las propuestas de (La Red Martínez, 2017) y (Agostini, 2019).

De esta manera se han atendido todas las solicitudes de recursos de todos los procesos respetando la exclusión mutua y las prioridades de los procesos, las prioridades nodales y las prioridades finales.

4.8. Comparación de los resultados obtenidos con los métodos propuestos y los métodos tradicionales

Se puede observar en la Tabla 179, los resultados obtenidos para las asignaciones de los recursos a procesos del paper (La Red Martínez, 2017) que dio inicio al trabajo, los dos escenarios propuestos considerando la imputación/asignación y el método tradicional, detallados con la iteración, la posición de 84 asignaciones que se encuentran y el valor de prioridad global.

Tabla 179. Asignaciones obtenidas en (La Red Martínez, 2017), del escenario E1, del escenario E2 y del método tradicional (MT).

Asig. Recurso Proceso	(La Red Martínez, 2017)			Escenario E1			Escenario E2			Método Tradicional		
	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global
r_{12} al p_{37}	1	1	0,4731	1	1	0,4717	1	1	0,4807	9	77	0,0101
r_{33} al p_{23}	1	2	0,4680	1	2	0,4690	1	2	0,4533	1	1	0,0667
r_{31} al p_{34}	1	3	0,3705	1	3	0,3726	1	4	0,3621	8	71	0,0067
r_{11} al p_{37}	1	4	0,3512	1	4	0,3500	1	7	0,3102	4	36	0,0239
r_{22} al p_{34}	1	5	0,3440	1	5	0,3455	1	3	0,3715	3	23	0,0336
r_{21} al p_{37}	1	6	0,3300	1	7	0,3251	1	5	0,3469	7	66	0,0074
r_{13} al p_{13}	1	7	0,3286	1	6	0,3290	2	16	0,2552	2	15	0,0389
r_{32} al p_{34}	1	8	0,3032	1	8	0,3033	1	8	0,3094	4	39	0,0121

Asig. Recurso Proceso	(La Red Martínez, 2017)			Escenario E1			Escenario E2			Método Tradicional		
	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global
r23 al p11	1	9	0,2492	1	9	0,2496	2	19	0,1592	6	60	0,0039
r24 al p34	1	10	0,1895	1	10	0,1864	2	20	0,1507	2	20	0,0215
r33 al p34	2	11	0,4065	4	31	0,2874	2	12	0,3917	9	76	0,0104
r12 al p34	2	12	0,3995	2	12	0,3977	3	21	0,3381	7	62	0,0192
r31 al p13	2	13	0,3035	2	13	0,3054	2	14	0,2949	3	25	0,0328
r11 al p11	2	14	0,2815	2	14	0,2800	2	18	0,2403	5	46	0,0179
r22 al p11	2	15	0,2702	2	15	0,2715	3	24	0,2345	4	33	0,0276
r21 al p25	2	16	0,2627	2	16	0,2577	3	25	0,2215	6	57	0,0114
r13 al p34	2	17	0,2570	2	17	0,2570	1	6	0,3263	6	55	0,0154
r32 al p37	2	18	0,2379	2	18	0,2378	2	17	0,2439	3	29	0,0174
r23 al p34	2	19	0,1732	2	19	0,1740	6	59	0,0204	3	28	0,0180
r24 al p11	2	20	0,1344	2	20	0,1311	1	10	0,2087	4	38	0,0121
r33 al p13	3	21	0,3468	2	11	0,4074	4	32	0,2707	4	31	0,0379
r12 al p23	3	22	0,3344	3	22	0,3325	4	31	0,2799	2	12	0,0480
r31 al p21	3	23	0,2425	3	23	0,2443	3	23	0,2386	1	3	0,0543
r22 al p13	3	24	0,2233	4	34	0,1794	2	13	0,2975	8	72	0,0061
r11 al p13	3	25	0,2123	3	25	0,2106	4	37	0,1307	8	75	0,0031
r21 al p13	3	26	0,1998	4	35	0,1597	2	15	0,2795	8	74	0,0034
r13 al p31	3	27	0,1861	3	27	0,1859	3	26	0,1856	8	73	0,0047
r32 al p13	3	28	0,1752	3	28	0,1748	4	38	0,1194	5	49	0,0080
r23 al p21	3	29	0,1079	3	29	0,1085	3	29	0,1076	2	18	0,0247
r24 al p23	3	30	0,0952	3	30	0,0918	3	30	0,0955	5	48	0,0080
r33 al p37	4	31	0,2873	3	21	0,3474	3	22	0,3302	8	69	0,0158
r12 al p13	4	32	0,2764	4	32	0,2743	2	11	0,4070	3	22	0,0413
r31 al p23	4	33	0,1964	4	33	0,1978	4	33	0,1922	2	13	0,0422
r22 al p12	4	34	0,1798	5	44	0,1354	4	34	0,1760	5	43	0,0222
r21 al p12	4	35	0,1573	3	26	0,2019	6	55	0,0768	4	37	0,0211
r11 al p12	4	36	0,1431	4	36	0,1432	3	28	0,1765	1	6	0,0433
r13 al p21	4	37	0,1363	4	37	0,1357	4	36	0,1354	1	4	0,0483
r32 al p23	4	38	0,1172	4	38	0,1166	3	27	0,1811	1	9	0,0297
r23 al p32	4	39	0,0710	4	39	0,0713	4	39	0,0704	4	40	0,0119
r24 al p35	4	40	0,0630	6	60	0,0184	4	40	0,0632	3	30	0,0168
r33 al p12	5	41	0,2280	5	41	0,2296	7	61	0,1343	11	83	0,0020
r12 al p11	5	42	0,2246	5	42	0,2223	5	42	0,2147	11	82	0,0040
r31 al p31	5	43	0,1519	5	43	0,1528	5	43	0,1475	7	64	0,0114
r22 al p21	5	44	0,1385	6	54	0,0938	6	54	0,0881	1	5	0,0467
r21 al p22	5	45	0,1160	5	45	0,1167	5	45	0,1157	1	7	0,0405
r13 al p32	5	46	0,0971	5	46	0,0962	6	56	0,0698	3	24	0,0329
r11 al p32	5	47	0,0899	5	47	0,0896	5	47	0,0892	3	26	0,0299
r32 al p36	5	48	0,0669	5	48	0,0670	5	48	0,0663	2	19	0,0228
r23 al p33	5	49	0,0440	5	49	0,0443	5	49	0,0437	5	50	0,0079
r24 al p36	5	50	0,0411	4	40	0,0595	5	50	0,0408	1	10	0,0268
r33 al p31	6	51	0,1828	6	51	0,1841	5	41	0,2191	3	21	0,0439
r12 al p12	6	52	0,1767	6	52	0,1743	6	52	0,1662	4	32	0,0353
r31 al p12	6	53	0,1141	6	53	0,1148	6	53	0,1095	9	79	0,0020
r22 al p22	6	54	0,0994	3	24	0,2245	5	44	0,1317	2	14	0,0400
r21 al p11	6	55	0,0774	6	55	0,0777	4	35	0,1642	2	17	0,0325
r13 al p36	6	56	0,0698	6	56	0,0692	5	46	0,0959	7	65	0,0101
r11 al p36	6	57	0,0660	6	57	0,0656	6	57	0,0655	6	56	0,0125
r32 al p35	6	58	0,0432	6	58	0,0430	6	58	0,0429	7	68	0,0027
r23 al p24	6	59	0,0206	6	59	0,0210	1	9	0,2247	1	8	0,0354
r24 al p24	6	60	0,0202	5	50	0,0385	6	60	0,0201	6	59	0,0040
r33 al p21	7	61	0,1406	7	61	0,1415	6	51	0,1765	5	41	0,0319
r12 al p21	7	62	0,1367	7	62	0,1340	7	62	0,1312	1	2	0,0574
r31 al p22	7	63	0,0767	7	63	0,0771	7	63	0,0765	4	34	0,0275
r22 al p35	7	64	0,0544	8	72	0,0333	7	64	0,0488	6	53	0,0168
r13 al p35	7	65	0,0435	7	66	0,0431	7	65	0,0428	4	35	0,0269
r21 al p33	7	66	0,0431	7	65	0,0433	7	66	0,0427	3	27	0,0265
r11 al p33	7	67	0,0427	7	67	0,0424	7	67	0,0423	7	67	0,0071

Asig. Recurso Proceso	(La Red Martínez, 2017)			Escenario E1			Escenario E2			Método Tradicional		
	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global	Iter.	Posición Global	Prior. Global
<i>r</i> ₃₂ al <i>p</i> ₃₃	7	68	0,0210	7	68	0,0210	7	68	0,0209	6	58	0,0054
<i>r</i> ₁₂ al <i>p</i> ₃₃	8	69	0,1098	8	69	0,1069	8	69	0,1041	12	84	0,0020
<i>r</i> ₃₃ al <i>p</i> ₂₂	8	70	0,0986	8	70	0,0992	8	70	0,0983	2	11	0,0547
<i>r</i> ₃₁ al <i>p</i> ₃₆	8	71	0,0478	8	71	0,0480	8	71	0,0474	6	54	0,0168
<i>r</i>₂₂ al <i>p</i>₃₃	8	72	0,0331	7	64	0,0545	8	72	0,0276	9	78	0,0028
<i>r</i> ₂₁ al <i>p</i> ₃₆	8	73	0,0215	8	73	0,0213	8	73	0,0213	5	47	0,0161
<i>r</i>₁₃ al <i>p</i>₃₃	8	74	0,0210	8	74	0,0208	8	74	0,0205	5	45	0,0209
<i>r</i> ₁₁ al <i>p</i> ₂₄	8	75	0,0203	8	75	0,0202	8	75	0,0202	2	16	0,0366
<i>r</i> ₁₂ al <i>p</i> ₃₆	9	76	0,0844	9	76	0,0818	9	76	0,0790	6	52	0,0245
<i>r</i> ₃₃ al <i>p</i> ₃₃	9	77	0,0659	9	77	0,0662	9	77	0,0654	10	81	0,0060
<i>r</i> ₃₁ al <i>p</i> ₃₅	9	78	0,0222	9	78	0,0220	9	78	0,0220	5	44	0,0221
<i>r</i> ₂₂ al <i>p</i> ₃₆	9	79	0,0122	9	79	0,0121	9	79	0,0121	7	63	0,0115
<i>r</i>₁₂ al <i>p</i>₂₄	10	80	0,0603	11	82	0,0368	11	82	0,0342	8	70	0,0145
<i>r</i> ₃₃ al <i>p</i> ₃₅	10	81	0,0425	10	81	0,0430	10	81	0,0422	6	51	0,0265
<i>r</i> ₁₂ al <i>p</i> ₃₂	11	82	0,0380	10	80	0,0578	10	80	0,0550	5	42	0,0299
<i>r</i> ₃₃ al <i>p</i> ₃₆	11	83	0,0196	11	83	0,0200	11	83	0,0194	7	61	0,0212
<i>r</i>₁₂ al <i>p</i>₃₅	12	84	0,0169	12	84	0,0164	12	84	0,0164	10	80	0,0067

Fuente: Elaboración propia.

Nota: Las relaciones recurso-proceso marcadas tienen valores imputados.

Con los resultados obtenidos de la Tabla 179, se pueden comparar las asignaciones del paper (La Red Martínez, 2017) y los diferentes escenarios (E1, E2 y MT) en cuanto a las iteraciones y las posiciones globales obtenidas.

Se observan que posterior a la imputación/asignación, de las 84 asignaciones se presentan variaciones en las iteraciones para las diferentes asignaciones, como lo indican la Tabla 180, Tabla 181 y Tabla 182.

Tabla 180. Comparativa de asignaciones obtenidas en cada iteración de (La Red Martínez, 2017) con respecto al escenario E1.

Estado	Cantidad asignaciones	Porcentaje
Mantuvieron	68	80,96%
Bajaron	8	9,52%
Subieron	8	9,52%
Total	84	100%

Fuente: Elaboración propia.

- Se observa en la Tabla 180, que para el escenario E1 se mantuvieron dentro de la misma iteración el 80,96% de las asignaciones, el 9,52% bajaron de iteración, y el 9,52% subieron de iteración.

Tabla 181. Comparativa de asignaciones obtenidas en cada iteración de (La Red Martínez, 2017) con respecto al escenario E2.

Estado	Cantidad asignaciones	Porcentaje
Mantuvieron	53	63,10%
Bajaron	16	19,05%
Subieron	15	17,85%
Total	84	100%

Fuente: Elaboración propia.

- Se observa en la Tabla 181, que para el escenario E2 se mantuvieron dentro de la misma iteración el 63,10% de las asignaciones, el 19,07% bajaron de iteración, y el 17,85% subieron de iteración.

Tabla 182. Comparativa de asignaciones obtenidas en cada iteración de (La Red Martínez, 2017) con respecto al método tradicional.

Estado	Cantidad asignaciones	Porcentaje
Mantuvieron	7	8,33%
Bajaron	37	44,05%
Subieron	40	47,62%
Total	84	100%

Fuente: Elaboración propia

- Para los métodos tradicionales se observa en la Tabla 182 que se mantuvieron dentro de la misma iteración el 8,33% de las asignaciones, el 44,05% bajaron de iteración y el 47,62% subieron de iteración.

Las posiciones globales para las asignaciones representa el orden en que los recursos son asignados a los proceso, se observan que posterior a la imputación/asignación de las 84 asignaciones existen variaciones en las posiciones globales para las diferentes situaciones, se puede comparar los resultados obtenidos de la Tabla 179 con respecto al trabajo del paper (La Red Martínez, 2017) en la Tabla 183, Tabla 184 y Tabla 185.

- Se observa en la Tabla 183, que para el escenario E1 se mantuvieron en la misma posición global el 76,20% de las asignaciones, el 11,90% bajaron de posición global (perdieron prioridad), y el

11,90% subieron de posición global (ganaron prioridad).

Tabla 183. Comparativa de las posiciones globales obtenidas en las asignaciones de (La Red Martínez, 2017) con respecto al escenario E1.

Estado	Cantidad asignaciones	Porcentaje
Mantuvieron	64	76,20%
Bajaron	10	11,90%
Subieron	10	11,90%
Total	84	100%

Fuente: Elaboración propia.

Tabla 184. Comparativa de las posiciones globales obtenidas en las asignaciones de (La Red Martínez, 2017) con respecto al escenario E2.

Estado	Cantidad asignaciones	Porcentaje
Mantuvieron	43	51,19%
Bajaron	21	25,00%
Subieron	20	23,81%
Total	84	100%

Fuente: Elaboración propia

- Para el escenario E2 se observa en la Tabla 184, que se mantuvieron en la misma posición el 51,19% de las asignaciones, el 25,00% bajaron de posición y el 23,81% subieron de posición.

Tabla 185. Comparativa de las posiciones globales obtenidas en las asignaciones de (La Red Martínez, 2017) con respecto al método tradicional.

Posición MT	Cantidad asignaciones	Porcentaje
Mantuvieron	1	1,19%
Bajaron	40	47,62%
Subieron	43	51,19%
Total	84	100%

Fuente: Elaboración propia

- Para los métodos tradicionales se observa en la Tabla 185, que se mantuvieron en la misma posición sólo el 1,19% de las asignaciones, el 47,62% bajaron de posición, y el 51,19% subieron de posición, teniendo en cuenta que solo se consideran para este método un solo criterio que es la prioridad de proceso, de los ocho criterios considerados en la propuesta de (La Red Martínez,

2017).

Para aplicar la imputación/asignación se tuvieron en cuenta diferentes pautas que fueron establecidas con anterioridad, tanto para la carga computacional y para las prioridades o preferencias de los procesos, para aplicar la imputación se tuvieron en cuenta información histórica de ciclos de recolección, y para la asignación por defecto las características del nodo.

Carga computacional

Para la carga computacional se puede observar en la Tabla 186 y Tabla 187, la cantidad de nodos que han sido afectados con la imputación/asignación, teniendo en cuenta las pautas establecidas en los diferentes escenarios.

Tabla 186. Cantidad de nodos imputados/asignados con las pautas establecidas para la carga computacional del escenario E1.

Pautas	Cantidad de nodos	Porcentaje
- Imputación con información histórica del nodo	2	66,67%
- Asignación por defecto	1	33,33%
- Ninguna	0	0,00%
Total	3	100%

Fuente: Elaboración propia.

- En el escenario E1 fueron afectados todos los nodos con la falta de información de una variable, se observa en la Tabla 186 que el 66,67% fueron imputados con información histórica, y el 33,33% fueron asignados con un valor por defecto en función a las características del nodo.

Tabla 187. Cantidad de nodos imputados/asignados con las pautas establecidas para la carga computacional del escenario E2.

Pautas	Cantidad de nodos	Porcentaje
- Imputación con información histórica del nodo	1	33,33%
- Asignación por defecto	0	0,00%
- Ninguna	2	66,67%
Total	3	100%

Fuente: Elaboración propia.

- El escenario E2 considera información de control faltante de todas las variables de un nodo

cualquiera, se observa en la Tabla 187 que el 33,33% de los nodos fueron imputados con información histórica, el 0,00% representa la asignación por defecto que no fue aplicado, y el 66,67% los nodos no tuvieron información de control faltante, lo cual indica que no hubo necesidad de aplicar la imputación o la asignación a los mismos.

Prioridad o preferencia de los procesos

Para la prioridad o preferencia de los procesos se pueden observar en la Tabla 188, Tabla 189 y Tabla 190 la cantidad de relaciones recurso-proceso que han sido afectados con la imputación/asignación, teniendo en cuenta las pautas establecidas en los diferentes escenarios.

Tabla 188. Cantidad de relaciones recurso-proceso imputados/asignados con las pautas establecidas para calcular la prioridad o preferencia de los procesos del escenario E1.

Pautas	Cantidad de recurso-proceso	Porcentaje
- Imputación con información histórica $r_{kl} p_{ij}$	8	66,66%
- Imputación con información histórica $r_{xx} p_{ij}$	2	16,67%
- Asignación por defecto	2	16,67%
Total	12	100 %

Fuente: Elaboración propia.

- En el escenario E1 fueron afectados en total 12 relaciones de recurso-proceso con la falta de información de una variable, se observa en la Tabla 188 que el 66,66% fueron imputados con información histórica de la misma relación de recurso-proceso ($r_{kl} p_{ij}$), el 16,67% fueron imputados con información histórica del proceso con otros recursos ($r_{xx} p_{ij}$) y el 16,67% fueron asignados con un valor por defecto en función a las características del nodo.

Tabla 189. Cantidad de relaciones recurso-proceso imputados/asignados con las pautas establecidas para calcular la prioridad o preferencia de los procesos del escenario E2.

Pautas	Cantidad de recurso-proceso	Porcentaje
- Imputación con información histórica $r_{kl} p_{ij}$	11	91,67%
- Imputación con información histórica $r_{xx} p_{ij}$	No considera	No considera
- Asignación por defecto	1	8,33%
Total	12	100 %

Fuente: Elaboración propia.

- En el escenario E2 fueron afectados en total 12 relaciones de recurso-proceso con la falta de

información de todas las variables, se observa en la Tabla 189 que el 91,67 % fueron imputados con información histórica de la misma relación de recurso-proceso ($r_{kl} p_{ij}$) y el 8,33% fueron asignados con un valor por defecto en función a las características del nodo.

Tabla 190. Cantidad de relaciones recurso-proceso imputados/asignados con las pautas establecidas para calcular la prioridad o preferencia de los procesos para el método tradicional.

Pautas	Cantidad de recurso-proceso	Porcentaje
- Imputación con información histórica $r_{kl} p_{ij}$	11	91,67%
- Imputación con información histórica $r_{xx} p_{ij}$	No considera	No considera
- Asignación por defecto	1	8,33%
Total	12	100 %

Fuente: Elaboración propia.

- En el método tradicional fueron afectados en total 12 relaciones de recurso-proceso con la falta de información la única variable considerada que es la prioridad de proceso, se observa en la Tabla 190 que el 91,67 % fueron imputados con información histórica de la misma relación de recurso-proceso ($r_{kl} p_{ij}$) y el 8,33% fueron asignados con un valor por defecto en función a las características del nodo.

4.9. Consideraciones acerca de las operaciones de agregación

En (La Red Martínez, 2017) y (Agostini, 2019) las características de las operaciones de agregación descriptas permiten considerar que el método propuesto pertenece a la familia de operadores de agregación *Neat-OWA*, operadores que se caracterizan por lo siguiente:

La definición de los operadores OWA indica que $f(a_1, a_2, \dots, a_n) = \sum_{j=1}^n w_j \cdot b_j$, donde b_j es

el j -ésimo valor mayor de las a_n , con la restricción para los pesos de satisfacer (1) $w_i \in [0,1]$ y (2)

$$\sum_{i=1}^n w_i = 1.$$

Para la familia de operadores *Neat*-OWA los pesos serán calculados en función de los elementos que se agregan, o más exactamente de los valores a agregar ordenados, los b_j , manteniéndose las condiciones (1) y (2). En este caso los pesos son: $w_i = f_i(b_1, \dots, b_n)$, definiéndose el operador

$$F(a_1, \dots, a_n) = \sum_i f_i(b_1, \dots, b_n) \cdot b_i$$

Para esta familia, donde los pesos dependen de la agregación, no se exige la satisfacción de todas las propiedades de los operadores OWA.

Además, para poder afirmar que un operador de agregación es *Neat*, es necesario que el valor final de agregación sea independiente del orden de los valores. Sea $A = (a_1, \dots, a_n)$ las estradas a agregar, sea $B = (b_1, \dots, b_n)$ las entradas ordenadas y $C = (c_1, \dots, c_n) = \text{Perm}(a_1, \dots, a_n)$ una permutación de las entradas. Formalmente se define un operador OWA como *neat* si

$$F(a_1, a_2, \dots, a_n) = \sum_{i=1}^n w_i \cdot b_i$$

Produce el mismo resultado para cualquier asignación $C = B$.

4.10. Discusiones y comentarios

El modelo propuesto en esta investigación considera la imputación de datos utilizando métodos de imputación para diferentes escenarios, cuando el nodo central recibe la información con valores faltantes de alguno de los nodos del sistema distribuido. El modelo define cuál de los métodos de imputación se utilizará y ejecutará cuando la información de las variables que indican el estado de carga sea incompleta y/o inexistente, y finalmente el modelo de decisión pueda

establecer un correcto orden de asignación de recursos a los procesos que solicitan.

El orden de asignación será determinado por la prioridad promedio general de todas las asignaciones de cada grupo. El sistema distribuido regula y actualiza constantemente el estado local de cada nodo, las decisiones de acceso a los recursos modifican estos estados por lo que debe ser reajustado repetidamente, garantizando la exclusión mutua y reordenando nuevas prioridades. El método debe repetirse siempre que haya grupos de procesos que requieran recursos compartidos.

El contenido de este capítulo ha servido para armar un artículo que será enviado a una revista de alto impacto.

Capítulo V - Conclusiones y futuras líneas de investigación

5.1. Conclusiones

Capítulo I

En el capítulo I se ha hecho la descripción de antecedentes, el marco teórico, se han planteado las hipótesis, se han considerado a un modelo de decisión agregar capas de imputación para cuando la información de control de alguno de los nodos lleguen al nodo central de manera *incompleta* (escenario 1) o *inexistente* (escenario 2), se han contemplado situaciones en las que los procesos accedieron a recursos compartidos en la modalidad de exclusión mutua con o sin constituir grupos de procesos, que podían requerir o no sincronización y que podían tener o no exigencias estrictas de consenso para lograr el acceso.

Capítulo II

En el capítulo II el modelo de decisión propuesto considera la imputación de datos cuando la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos llega de manera incompleta al Runtime central, reemplazando esos valores faltantes por valores estimados, para que finalmente el modelo de decisión establezca un correcto orden de asignación de recursos a los procesos que han solicitado, respetando la modalidad de exclusión mutua distribuida.

El método de imputación utilizado es considerado uno de los más fiables y de menor consumo de recursos computacionales.

Es importante mencionar que la solución presentada es compatible con la gestión de transacciones globales habitualmente implementada en los sistemas de gestión de base de datos.

Es un método general, que permite ser adaptado a diferentes circunstancias y objetivos, de acuerdo a las necesidades de los nodos, cargas de trabajo, prioridades, etc.

Capítulo III

En el capítulo III el modelo de decisión propuesto considera la imputación de datos cuando la información sobre las variables que indican el estado de carga de alguno o algunos de los nodos no llega en su totalidad al Runtime central, reemplazando esos valores faltantes por valores estimados, para que finalmente el modelo de decisión establezca un correcto orden de asignación de recursos a los procesos, respetando la exclusión mutua.

Al igual que el capítulo anterior es un método general, que permite ser adaptado a diferentes circunstancias y objetivos, de acuerdo a las necesidades de los nodos, cargas de trabajo, prioridades, etc.

Capítulo IV

Se han definido las características del Modelo de Decisión. Se ha explicado que el modelo de decisión utiliza un Runtime que gestiona los procesos y recursos compartidos y define el escenario correspondiente.

Los métodos de imputación utilizan información de control de ciclos anteriores que están disponibles en el nodo central, y en caso de que el nodo haya arrancado recientemente y no se disponga de dicha información se consideran las prestaciones de los nodos (valores asumidos en función a las características del nodo, por ejemplo, su velocidad de procesamiento, capacidad de disco, memoria entre otros) para asignar valor por defecto.

5.2. Futuras líneas de investigación

Como futuras líneas de investigación relacionadas al desarrollo de este trabajo podemos mencionar las siguientes:

- Se prevé mejorar la evaluación de la propuesta considerando características como tiempo de ejecución, consumo de memoria y de CPU. Del mismo modo, dicha evaluación, será aplicada y comparada con otras aproximaciones que intenten resolver el mismo problema.
- Se tiene previsto considerar criterios de optimización del consumo de energía por parte de los componentes de los sistemas distribuidos de procesamiento en función de los objetivos de Desarrollo Sostenible de las Naciones Unidas.
- Evaluar la migración controlada de procesos considerando la transferencia de procesos entre nodos de forma automática para hacer mejor uso de los recursos, característica que también está ligada a la distribución y el equilibrio de la carga de trabajo entre los nodos, contribuyendo a aumentar la tolerancia a fallos y así mejorar el rendimiento global del sistema distribuido, aportando a su autorregulación.

Referencias

- Abril, Daniel, Guillermo Navarro-Arribas, and Torra Vicenc. (2014). "Aprendizaje Supervisado Para El Enlace de Registros a Través de La Media Ponderada." *Actas de La XIII Reunión Española Sobre Criptología y Seguridad de La Información (RECSI 2014)* 281–84.
- Agostini, F., D. L. La Red Martínez, and J. C. Acosta. (2018). "Modeling of the Consensus in the Allocation of Resources in Distributed Systems." *International Journal of Advanced Computer Science and Applications* 9(12). doi: 10.14569/IJACSA.2018.091204.
- Agostini, F. (2019). "Nueva Propuesta Para La Administración de Recursos y Procesos En Sistemas Distribuidos." Universidad Nacional del Nordeste.
- Agostini, F., and D. L. La Red Martínez. (2019). "Allocation of Shared Resources." *14th Iberian Conference on Information Systems and Technologies - CISTI 2019* 1–6.
- Agostini, Federico, D. L. la Red Martínez, and Julio C. Acosta. (2019). "Assignment of Resources in Distributed Systems with Strict Consensus Requirements." in *IMCIC 2019 - 10th International Multi-Conference on Complexity, Informatics and Cybernetics, Proceedings*. Vol. 1.
- Agrawal, Divyakant, and Amr El Abbadi. (1991). "An Efficient and Fault-Tolerant Solution for Distributed Mutual Exclusion." *ACM Transactions on Computer Systems (TOCS)* 9(1). doi: 10.1145/103727.103728.
- Alagarsamy, K. (2003). "Some Myths about Famous Mutual Exclusion Algorithms." *ACM SIGACT News* 34(3). doi: 10.1145/945526.945527.
- Aljuaid, Tahani, and Sreela Sasi. (2016). "Proper Imputation Techniques for Missing Values in Data Sets." in *Proceedings of the 2016 International Conference on Data Science and Engineering, ICDSE 2016*.
- Baez, Diego, Maricruz Olazabal, and Jorge Romero. (2019). *Toma de Decisiones Empresariales a Través de La Media Ponderada Ordenada*. Vol. 19.
- Bagrodia, Rajive. (1989). "Process Synchronization: Design and Performance Evaluation of Distributed Algorithms." *IEEE Transactions on Software Engineering* 15(9). doi: 10.1109/32.31364.
- Batista, Gustavo E. A. P. A., and Maria Carolina Monard. (2015). "A Study of K-Nearest Neighbour as a Model-Based Method to Treat Missing Data." *Machine Learning (August)*.
- Berchtold, André, and Joan Carles Surís. (2017). "Imputation of Repeatedly Observed Multinomial Variables in Longitudinal Surveys." *Communications in Statistics: Simulation and Computation* 46(4). doi: 10.1080/03610918.2015.1082588.
- Birman, Kenneth P., and Thomas A. Joseph. (1987). "Reliable Communication in the Presence of Failures." *ACM Transactions on Computer Systems (TOCS)* 5(1). doi: 10.1145/7351.7478.
- Birman, Kenneth, André Schiper, and Pat Stephenson. (1991). "Lightweight Causal and Atomic Group Multicast." *ACM Transactions on Computer Systems (TOCS)* 9(3). doi: 10.1145/128738.128742.
- Birman, Kenneth P. (2005). *Reliable Distributed Systems: Technologies, Web Services, and Applications*.
- Cao, Guohong, and Mukesh Singhal. (2001). "A Delay-Optimal Quorum-Based Mutual Exclusion Algorithm for Distributed Systems." *IEEE Transactions on Parallel and Distributed Systems* 12(12). doi: 10.1109/71.970560.

- Cartwright, M. H., M. J. Shepperd, and Q. Song. (2003). "Dealing with Missing Software Project Data." in *Proceedings - International Software Metrics Symposium*. Vols. 2003-Janua.
- Casleton, Emily, Dave Osthus, and Kendra Van Buren. (2018). "Imputation for Multisource Data with Comparison and Assessment Techniques." in *Applied Stochastic Models in Business and Industry*. Vol. 34.
- Chang, Gang, and Tongmin Ge. (2011). "Comparison of Missing Data Imputation Methods for Traffic Flow." in *Proceedings 2011 International Conference on Transportation, Mechanical, and Electrical Engineering*, TMEE 2011.
- Chao, Xiangrui, Gang Kou, and Yi Peng. (2016). "An Optimization Model Integrating Different Preference Formats." in *2016 6th International Conference on Computers Communications and Control*, ICCCC 2016.
- Chen, J., and J. Shao. (2000). "Nearest Neighbor Imputation for Survey Data." *Journal of Official Statistics* 16(2).
- Chen, Chen T. (2001). "Applying Linguistic Decision-Making Method to Deal with Service Quality Evaluation Problems." *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems* 9(SUPPL.). doi: 10.1142/S0218488501001022.
- Chiclana, F., E. Herrera-Viedma, F. Herrera, and S. Alonso. (2004). "Induced Ordered Weighted Geometric Operators and Their Use in the Aggregation of Multiplicative Preference Relations." *International Journal of Intelligent Systems* 19(3). doi: 10.1002/int.10172.
- Chiclana, Francisco, Francisco Chiclana, Francisco Herrera, and Enrique Herrera-viedma. (2000). "The Ordered Weighted Geometric Operator: Properties and Application in MCDM Problems." *IN PROC. 8TH CONF. INFORM. PROCESSING AND MANAGEMENT OF UNCERTAINTY IN KNOWLEDGEBASED SYSTEMS* 985–91.
- Chiclana, F., F. Herrera, and E. Herrera-Viedma. (2001). "Integrating Multiplicative Preference Relations in a Multipurpose Decision-Making Model Based on Fuzzy Preference Relations." *Fuzzy Sets and Systems* 122(2). doi: 10.1016/S0165-0114(00)00004-X.
- Choudhury, Suvra Jyoti, and Nikhil R. Pal. (2019). "Imputation of Missing Data with Neural Networks for Classification." *Knowledge-Based Systems* 182. doi: 10.1016/j.knosys.2019.07.009.
- Cicirelli, Franco, and Libero Nigro. (2016). "Modelling and Verification of Mutual Exclusion Algorithms." in *Proceedings - IEEE International Symposium on Distributed Simulation and Real-Time Applications, DS-RT*.
- Clarke, M. R. B., Richard O. Duda, and Peter E. Hart. (1974). "Pattern Classification and Scene Analysis." *Journal of the Royal Statistical Society. Series A (General)* 137(3). doi: 10.2307/2344977.
- Cover, T. M., and P. E. Hart. (1967). "Nearest Neighbor Pattern Classification." *IEEE Transactions on Information Theory* 13(1):21–27. doi: 10.1109/TIT.1967.1053964.
- De Leeuw, Edith D. (2001). "Reducing Missing Data in Surveys: An Overview of Methods." *Quality and Quantity* 35(2).
- Dong, Yucheng, Hengjie Zhang, and Enrique Herrera-Viedma. (2016). "Consensus Reaching Model in the Complex and Dynamic MAGDM Problem." *Knowledge-Based Systems* 106. doi: 10.1016/j.knosys.2016.05.046.

- Downey, Ronald G., and Craig V. King. (1998). "Missing Data in Likert Ratings: A Comparison of Replacement Methods." *Journal of General Psychology* 125(2). doi: 10.1080/00221309809595542.
- Dubes, R. C., and A. K. Jain. (1988). "Algorithms for Clustering Data." *Prentice Hall College*.
- Farhangfar, Alireza, Lukasz A. Kurgan, and Witold Pedrycz. (2007). "A Novel Framework for Imputation of Missing Values in Databases." *IEEE Transactions on Systems, Man, and Cybernetics Part A: Systems and Humans* 37(5). doi: 10.1109/TSMCA.2007.902631.
- Fodor, János, and Marc Roubens. (1994). Fuzzy Preference Modelling and Multicriteria Decision Support.
- Forc, Juan I., Miguel Pagola, Barrenechea Edurne, and Humberto Bustince. (2018). "Adaptación Automática Del Operador de Pooling Aprendiendo Pesos de Medias Ponderadas Ordenadas En Redes Neuronales Convolucionales." *XVIII Conferencia de La Asociación Española Para La Inteligencia Artificial (CAEPIA)* 1225–30.
- Freund, John E., and Gary A. Simon. (1994). "Estadística Elemental." Pp. 43–44 in. Naucalpan de Juárez, Edo. de México.
- Fullér, R. (1996). "OWA Operators in Decision Making." *Exploring the Limits of Support Systems*, TUCS General Publications 3.
- García-Melón, M., J. Ferrís-Oñate, J. Aznar-Bellver, and R. Poveda-Bautista. (2006). "Farmland Appraisal: An Analytic Network Process (ANP) Approach."
- García-Laencina, P. J., G. Rodríguez-Bermúdez, J. L. Roca-González, J. Roca-González, and J. Roca-Dorda. (2012). "Técnicas de Vecindad Para La Estimación de Información Incompleta En Problemas de Diagnóstico Médico." *XX Congreso Anual de La Sociedad Española de Ingeniería Biomédica (CASEIB 2012)*.
- Gediga, Günther, and Ivo Düntsch. (2003). "Maximum Consistency of Incomplete Data via Non-Invasive Imputation." *Artificial Intelligence Review* 19(1). doi: 10.1023/A:1022188514489.
- Gómez, Daniel, Karina Rojas, Javier Montero, J. Tinguaro Rodríguez, and Gleb Beliakov. (2014). "CONSISTENCIA Y ESTABILIDAD EN OPERADORES DE AGREGACIÓN: Una Aplicación Al Problema de Datos Perdidos." *ESTYLF 2014 XVII CONGRESO ESPAÑOL SOBRE TECNOLOGÍAS Y LÓGICA FUZZY* 289–96.
- Greco, Salvatore, Benedetto Matarazzo, and Roman Slowinski. (2002). "Rough Sets Methodology for Sorting Problems in Presence of Multiple Attributes and Criteria." *European Journal of Operational Research* 138(2). doi: 10.1016/S0377-2217(01)00244-2.
- Gu, Daqian, and Yang Gao. (2005). "Incremental Gradient Descent Imputation Method for Missing Data in Learning Classifier Systems."
- Huisman, Mark. (2000). "Imputation of Missing Item Responses: Some Simple Techniques." *Quality and Quantity* 34(4). doi: 10.1023/A:1004782230065.
- Huisman, Mark. (2014). "Imputation of Missing Network Data: Some Simple Procedures." in *Encyclopedia of Social Network Analysis and Mining*.
- Husson, François, Julie Josse, Balasubramanian Narasimhan, and Geneviève Robin. (2019). "Imputation of Mixed Data With Multilevel Singular Value Decomposition." *Journal of Computational and Graphical Statistics* 28(3). doi: 10.1080/10618600.2019.1585261.
- Jönsson, Per, and Claes Wohlin. (2004). "An Evaluation of K-Nearest Neighbour Imputation Using Likert Data." in *Proceedings - International Software Metrics Symposium*.

- Joseph, Thomas A., and Kenneth P. Briman. (1989). Reliable Broadcast Protocols, in *Distributed Systems*. Ithaca, NY 14853-7501.
- Joshi, Rajeev, and Gerard J. Holzmann. (2007). "A Mini Challenge: Build a Verifiable Filesystem." *Formal Aspects of Computing* 19(2). doi: 10.1007/s00165-006-0022-3.
- Kaashoek, Marinus Frans. (1992). Group Communication in Distributed Computer Systems. Amsterdam.
- Kalton, Graham, and Daniel Kasprzyk. (1982). IMPUTING FOR MISSING SURVEY RESPONSES. *University of Michigan*.
- Kamei, Sayaka, and Hirotsugu Kakugawa. (2018). "An Asynchronous Message-Passing Distributed Algorithm for the Global Critical Section Problem." in *Proceedings - 2017 5th International Symposium on Computing and Networking, CANDAR 2017*. Vols. 2018-Janua.
- Kaufman, Leonard., and Peter J. Rousseeuw. (1990). Finding Groups in Data: An Introduction to Cluster Analysis (Wiley Series in Probability and Statistics). Vol. 66.
- La Red Martínez, D. L., J. M. Doña, J. I. Peláez, and E. B. Fernández. (2011). "WKC-OWA, a New Neat-OWA Operator to Aggregate Information in Democratic Decision Problems." *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems* 19(5):759–79. doi: 10.1142/S0218488511007222.
- La Red Martinez, David L., and Julio C. Acosta. (2014). "Perspectives of New Decision Making Models of Processes Synchronization in Distributed Systems." *INTERNATIONAL JOURNAL OF MANAGEMENT & INFORMATION TECHNOLOGY* 9(1). doi: 10.24297/ijmit.v9i1.671.
- La Red Martínez, D. L., and N. Pinto. (2015). "Brief Review of Aggregation Operators." *Wulfenia* 22-Nº 4:114–37.
- La Red Martínez, David L., and Julio C. Acosta. (2015a). "Review of Modeling Preferences for Decision Models." *European Scientific Journal* 11(36):1–19.
- La Red Martínez, David L., and Julio C. Acosta. (2015b). "Aggregation Operators Review - Mathematical Properties and Behavioral Measures." *International Journal of Intelligent Systems and Applications* 7(10):63–76. doi: 10.5815/ijisa.2015.10.08.
- La Red Martínez, David L. (2017). "Aggregation Operator for Assignment of Resources in Distributed Systems." *International Journal of Advanced Computer Science and Applications* 8(10). doi: 10.14569/ijacsa.2017.081053.
- La Red Martínez, D. L., F. Agostini, and C. Primorac. (2017). "Modelo de Asignación de Recursos Para La Enseñanza de Los Procesos Distribuidos." *Primer Congreso Latinoamericano de Ingeniería - CLADI*.
- La Red Martínez, David L., Julio C. Acosta, and Federico Agostini. (2018). "Assignment of Resources in Distributed Systems." in *IMCIC 2018 - 9th International Multi-Conference on Complexity, Informatics and Cybernetics, Proceedings*. Vol. 2.
- La Red Martínez, David Luis. (2004). Sistemas Operativos. EUDENE - Argentina.
- Laird, R. J., and Rubin D.B. (1987). "Statistical Analysis with Missing Data."
- Lamport, Leslie. (1978). "Time, Clocks, and the Ordering of Events in a Distributed System." *Communications of the ACM* 21(7). doi: 10.1145/359545.359563.
- Little, Roderick J. A., and Donald B. Rubin. (2002). Statistical Analysis with Missing Data.

- Lodha, Sandeep, and Ajay Kshemkalyani. (2000). "A Fair Distributed Mutual Exclusion Algorithm." *IEEE Transactions on Parallel and Distributed Systems* 11(6). doi: 10.1109/71.862205.
- Lu, J., and D. Ruan. (2007). *Multi-Objective Group Decision Making: Methods, Software and Applications with Fuzzy Set Techniques*. London: Imperial College Press.
- Lu, Cheng Bo, and Ying Mei. (2018). "An Imputation Method for Missing Data Based on an Extreme Learning Machine Auto-Encoder." *IEEE Access* 6. doi: 10.1109/ACCESS.2018.2868729.
- Macedonia, Michael R., Michael J. Zyda, David R. Pratt, Donald P. Brutzman, and Paul T. Barham. (1995). "Exploiting Reality with Multicast Groups: A Network Architecture for Large-Scale Virtual Environments." in *Proceedings - Virtual Reality Annual International Symposium*.
- Madhu, G., and T. V. Rajinikanth. (2012). "A Novel Index Measure Imputation Algorithm for Missing Data Values: A Machine Learning Approach." in *2012 IEEE International Conference on Computational Intelligence and Computing Research, ICCIC 2012*.
- Marakas, G. (2002). "Decision Support Systems." New Jersey, USA.
- Martínez, Luis, Jun Liu, and Jian Bo Yang. (2006). "A Fuzzy Model for Design Evaluation Based on Multiple Criteria Analysis in Engineering Systems." *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems* 14(3). doi: 10.1142/S0218488506004035.
- Martínez, Luis, Jun Liu, Da Ruan, and Jian Bo Yang. (2007). "Dealing with Heterogeneous Information in Engineering Evaluation Processes." *Information Sciences* 177(7). doi: 10.1016/j.ins.2006.07.005.
- Mehala, B., and K. Vivekanandan. (2008). "An Analysis on K-Means Algorithm as an Imputation Method to Deal with Missing Values." *Asian Journal of Information Technology* 7:434–41.
- Multivariate Imputation by Chained Equations. (2020). [Online]. Available: <https://cran.r-project.org/web/packages/mice/index.html>.
- Musil, Carol M., Camille B. Warner, Piyanee Klainin Yobas, and Susan L. Jones. (2002). "A Comparison of Imputation Techniques for Handling Missing Data." *Western Journal of Nursing Research* 24(7). doi: 10.1177/019394502762477004.
- Myrtveit, Ingunn, Erik Stensrud, and Ulf H. Olsson. (2001). "Analyzing Data Sets with Missing Data: An Empirical Evaluation of Imputation Methods and Likelihood-Based Methods." *IEEE Transactions on Software Engineering* 27(11). doi: 10.1109/32.965340.
- Nelwamondo, Fulufhelo Vincent, and Tshilidzi Marwala. (2007). "Rough Sets Computations to Impute Missing Data."
- Ngoc Chau, Vo Thi, and Nguyen Hua Phung. (2018). "An Effective and Robust Self-Training Algorithm Using k-Means and Random Forest Models for Program-Level Student Classification." in *Proceedings of the International Conference on Inventive Communication and Computational Technologies, ICICCT 2018*.
- Nilsson, Petra. (2016). "An Imputation Model for Dropouts in Unemployment Data." *Journal of Official Statistics* 32(3). doi: 10.1515/JOS-2016-0036.
- Patil, Bankat M., Ramesh C. Joshi, and Durga Toshniwal. (2010). "Missing Value Imputation Based on K-Mean Clustering with Weighted Distance." *Communications in Computer and Information Science* 94 CCIS(PART 1):600–609. doi: 10.1007/978-3-642-14834-7_56.
- Peláez, J., J. M. Dona, and D. La Red. (2003). "Analysis of the Majority Process in Group Decision Making Process." *Proceedings of the 7th Joint Conference on Information Sciences*.

- Peláez, J., and J. M. Doña. (2003). "Majority Additive-Ordered Weighting Averaging: A New Neat Ordered Weighting Averaging Operator Based on the Majority Process." *International Journal of Intelligent Systems* 18(4). doi: 10.1002/int.10096.
- Peláez, J., JM Dona, and D. La Red. (2004). "Majority Opinion in Group Decision Making Using the Qma-Owa Operator." 449–54.
- Peláez, J., J. M. Doña, and J. A. Gómez-Ruiz. (2007). "Analysis of OWA Operators in Decision Making for Modelling the Majority Concept." *Applied Mathematics and Computation* 186(2). doi: 10.1016/j.amc.2006.07.161.
- Peláez, J., J. M. Doña, and D. L. La Red Martínez. (2009). "A Mix Model of Discounted Cash-Flow and OWA Operators for Strategic Valuation." *International Journal of Interactive Multimedia and Artificial Intelligence* 1(2).
- Preda, Cristian, Alain Duhamel, Monique Picavet, and Tahar Kechadi. (2005). "Tools for Statistical Analysis with Missing Data: Application to a Large Medical Database." in *Studies in Health Technology and Informatics*. Vol. 116.
- Primorac, Carlos Roberto, David Luís La Red Martínez, and Mirta Eve Giovannini. (2020). "Metodología De Evaluación Del Desempeño De Métodos De Imputación Mediante Una Métrica Tradicional Complementada Con Un Nuevo Indicador." *European Scientific Journal ESJ* 16(18). doi: 10.19044/esj.2020.v16n18p61.
- Raaijmakers, Quinten A. W. 1999. "Effectiveness of Different Missing Data Treatments in Surveys with Likert-Type Data: Introducing the Relative Mean Substitution Approach." *Educational and Psychological Measurement* 59(5).
- Rahman, M. Mostafizur, and D. N. Davis. (2013). Machine Learning-Based Missing Value Imputation Method for Clinical Datasets. Vol. 229 LNEE.
- Ricart, Glenn, and Ashok K. Agrawala. (1981). "An Optimal Algorithm for Mutual Exclusion in Computer Networks." *Communications of the ACM* 24(1). doi: 10.1145/358527.358537.
- Roy, Amarjit, Joyeeta Singha, Lalit Manam, and Rabul Hussain Laskar. (2017). "Combination of Adaptive Vector Median Filter and Weighted Mean Filter for Removal of High-Density Impulse Noise from Colour Images." *IET Image Processing* 11(6):352–61. doi: 10.1049/iet-ipr.2016.0320.
- Rubin, Donald B. (1976). "Inference and Missing Data." *Biometrika* 63(3):581–92.
- Van Renesse, Robbert, Kenneth P. Birman, and Werner Vogels. (2003). "Astrolabe: A Robust and Scalable Technology for Distributed System Monitoring, Management, and Data Mining." *ACM Transactions on Computer Systems* 21(2). doi: 10.1145/762483.762485.
- Saaty, T. L. (1980). "The Analytic Hierarchy Process." *MacGraw Hill*.
- Sande, I. G. (1983). "Hot-Deck Imputation Procedures." *Madow, W. G. and Olkin, I., Eds.* 3:334–50.
- Santos, Miriam Seoane, Ricardo Cardoso Pereira, Adriana Fonseca Costa, Jastin Pompeu Soares, Joao Santos, and Pedro Henriques Abreu. (2019). "Generating Synthetic Missing Data: A Review by Missing Mechanism." *IEEE Access* 7:11651–67. doi: 10.1109/ACCESS.2019.2891360.
- Scheffer, J. (2002). "Dealing with Missing Data." *Madow, W. G. and Olkin, I., Eds., Incomplete Data in Sample Surveys* 3:153–60.
- Schouten, Rianne Margaretha, Peter Lugtig, and Gerko Vink. (2018). "Generating Missing Values for Simulation Purposes: A Multivariate Amputation Procedure." *Journal of Statistical Computation and*

- Simulation* 88(15):2909–30. doi: 10.1080/00949655.2018.1491577.
- Schouten, Rianne Margaretha, and Gerko Vink. (2018). “The Dance of the Mechanisms: How Observed Information Influences the Validity of Missingness Assumptions.” *Sociological Methods and Research*. doi: 10.1177/0049124118799376.
- Si, Yanna, Jiexin Pu, Shaofei Zang, and Lifan Sun. (2021). “Extreme Learning Machine Based on Maximum Weighted Mean Discrepancy for Unsupervised Domain Adaptation.” *IEEE Access* 9:2283–93. doi: 10.1109/ACCESS.2020.3047448.
- Silberschatz, A., P. Baer G, and G. Gagne. (2006). *Fundamentos de Sistemas Operativos. Interamericana de España. S.A.U. España: Mc Graw-Hill.*
- Song, Qinbao, Martin Shepperd, and Michelle Cartwright. (2005). “A Short Note on Safest Default Missingness Mechanism Assumptions.” *Empirical Software Engineering* 10(2). doi: 10.1007/s10664-004-6193-8.
- Stallings, William. (2005). *Sistemas Operativos Aspectos Internos y Principios de Diseño. Vol. Quinta Edi.*
- Strike, Kevin, Khaled El Emam, and Nazim Madhavji. (2001). “Software Cost Estimation with Incomplete Data.” *IEEE Transactions on Software Engineering* 27(10). doi: 10.1109/32.962560.
- Tanenbaum, AS, G. Guerrero, and ÓAP Velasco. (1996). *Sistemas Operativos Distribuidos. México: Prentice-Hall Hispanoamericana S.A.*
- Tanenbaum, A. S., and M. Van Steen. (2008). *Sistemas Distribuidos - Principios y Paradigmas. 2da. México: Pearson Educación S. A.*
- Tanenbaum, Andrew S. (2009). *Sistemas Operativos Modernos Tercera Edición.*
- The R Project for Statistical Computing. (2020). “*The R Project for Statistical Computing.*”
- Todeschini, Roberto. (1990). “Weighted K-Nearest Neighbour Method for the Calculation of Missing Values.” *Chemometrics and Intelligent Laboratory Systems* 9(2). doi: 10.1016/0169-7439(90)80098-Q.
- Troyanskaya, Olga, Michael Cantor, Gavin Sherlock, Pat Brown, Trevor Hastie, Robert Tibshirani, David Botstein, and Russ B. Altman. (2001). “Missing Value Estimation Methods for DNA Microarrays.” *Bioinformatics* 17(6). doi: 10.1093/bioinformatics/17.6.520.
- Twala, Bhikisipho. (2009). “An Empirical Comparison of Techniques for Handling Incomplete Data Using Decision Trees.” *Applied Artificial Intelligence* 23(5):373–405. doi: 10.1080/08839510902872223.
- Vallejos, Oscar, Maria Valesani, and Enzo Rigonatto. (2011). “Técnicas de Minería de Datos Aplicada a Bases de Datos Imputadas.” *Escuela Superior Politécnica Del Litoral.*
- Wilks, S. S. (1932). “Moments and Distributions of Estimates of Population Parameters from Fragmentary Samples.” *The Annals of Mathematical Statistics* 3(3). doi: 10.1214/aoms/1177732885.
- Yager, Ronald. (1988). “On Ordered Weighted Averaging Aggregation Operators in Multicriteria Decisionmaking.” *IEEE Transactions on Systems, Man and Cybernetics* 18(1):183–90. doi: 10.1109/21.87068.
- Yager, Ronald. (1993). “Families of OWA Operators.” *Fuzzy Sets and Systems* 59(2). doi: 10.1016/0165-0114(93)90194-M.

- Yager, Ronald, and Janusz Kacprzyk. (1997). The Ordered Weighted Averaging Operators: *Theory and Applications*.
- Yager, Ronald, and G. Pasi. (2002). "Modeling Majority Opinion in Multi-Agent Decision Making." *International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems*.
- Yoo, Byoung Hyun, Junhwan Kim, Byun Woo Lee, Gerrit Hoogenboom, and Kwang Soo Kim. (2020). A Surrogate Weighted Mean Ensemble Method to Reduce the Uncertainty at a Regional Scale for the Calculation of Potential Evapotranspiration. Vol. 10. Springer US.
- Yugander, P., C. H. Tejaswini, J. Meenakshi, K. Samapath Kumar, B. V. N. Sures. Varma, and M. Jagannath. (2020). "MR Image Enhancement Using Adaptive Weighted Mean Filtering and Homomorphic Filtering." *Procedia Computer Science* 167(2019):677–85. doi: 10.1016/j.procs.2020.03.334.

Apéndice de publicaciones

1. Duré, D., Fornerón, J., Agostini, F., La Red Martínez, D. (2021). Imputación de información de control faltante en la gestión de recursos y procesos. Conferencia Ibero Americana WWW/Internet, (CIAWI 2021). **(PUBLICADO)**
2. Duré, D., Agostini, F., La Red Martínez, D. (2021). Imputación de valores faltantes sobre la información de control en el contexto de los sistemas distribuidos. Congreso Nacional de Ingeniería Informática y Sistemas de la Información, (9° CONAIISI 2021). **(PUBLICADO)**.