

Modelo de optimización difuso multiobjetivo multiperíodo para el diseño de una red logística
inversa sostenible

Jorge Esteban Caballero Rodríguez, Milena Patricia Aguilar Giraldo

Trabajo de Grado para Optar el título de Ingeniero Industrial

Director

Henry Lamos Diaz

Ph.D. Física - Matemática

Universidad Industrial de Santander

Facultad de Ingenierías Físico Mecánicas

Escuela de Estudios Industriales y Empresariales

Bucaramanga

2019

AGRADECIMIENTOS

A mi familia por apoyarme durante mi vida.

A Milena por trabajar conmigo en el proyecto y darme compañía.

A todos los profesores, en especial al profesor Henry Lamos, que me ha guiado

Y a Poe, Grass, Pessoa, Hesse, Baudelaire, Dazai y Kierkegaard por cuidarme.

JORGE ESTEBAN CABALLERO RODRÍGUEZ

A Dios por darme la fortaleza y perseverancia para culminar el ciclo de formación profesional en
pregrado.

A mis padres por siempre estar ahí a pesar de la distancia y apoyarme durante todo el tiempo de
estudio.

A mi hermano José Antonio por acompañarme en las noches de desvelo.

A mi apreciado compañero Jorge Esteban por su apoyo incondicional y valioso aporte en el
desarrollo de este proyecto.

A nuestro director Henry Lamos por la paciencia y guía que nos ofreció.

MILENA PATRICIA AGUILAR GIRALDO

Tabla de contenido

Introducción	17
1. Planteamiento del problema.....	20
2. Justificación del Proyecto	21
3. Objetivos	23
3.1. Objetivo general	23
3.2. Objetivos específicos	23
4. Revisión de la Literatura	24
4.1. La logística inversa desde un enfoque verde y sostenible	24
4.2. Modelos matemáticos aplicados a la logística inversa dentro de la cadena de suministro.....	27
4.3. Aplicaciones de algoritmos y métodos de soluciones en la logística inversa	30
5. Marco Teórico.....	32
5.1. Logística inversa	32
5.2. Logística verde y logística sostenible	33
5.3. Optimización multiobjetivo	34
5.3.1. Métodos de optimización multiobjetivo.	35
5.4. Lógica difusa.....	38
5.4.1. Conjuntos difusos.	38
5.4.2. Número difuso.	39
5.5.1. Strength Pareto Evolutionary Algorithm (SPEA).....	41
5.5.2. Algoritmo Genético	42

DISEÑO DE UNA RED LOGÍSTICA INVERSA SOSTENIBLE	8
5.5.3. NSGA-II (Non-Dominated Sorting Genetic Algorithm).	47
5.6. Valor presente neto.	50
5.7. Eco-Indicador 99.	51
6. Recolección de datos y Metodología	52
6.1. Definición del problema.	52
6.2. Obtención de datos.	56
6.3. Formulación del modelo matemático.	57
7. Algoritmo de solución propuesto.	62
7.1. Codificación de las soluciones	66
7.1.1. Cromosoma.	66
7.1.2. Matriz Population.	71
7.2. Parámetros del algoritmo	72
7.3. Generación de la población inicial	73
7.4. Criterio de selección.	76
7.5. Operador de cruce.	77
7.6. Operador de mutación:	78
7.7. Elitismo (selección de mejores individuos)	79
8. Aplicación del algoritmo para la solución del problema	80
8.1. Supuestos.	80
8.2. Transformación del modelo matemático.	81
8.2.1. Modelo matemático auxiliar resultante.	85
9. Validación del algoritmo.	87
9.1. Diseño experimental para el tamaño de la población.	90

DISEÑO DE UNA RED LOGÍSTICA INVERSA SOSTENIBLE	9
9.2. Diseño experimental para el porcentaje de mutación	97
9.3. Análisis del factor α	103
9.4. Resultados	105
9.5. Contraste con instancia de la literatura.	106
10. Conclusiones	107
11. Recomendaciones.....	109
Referencias Bibliográficas	110

Lista de tablas

Tabla 1. Cumplimiento de los objetivos del proyecto	19
Tabla 2. Representación de la ruta que sigue la demanda utilizada en el cromosoma.	67
Tabla 3. Representación de la ruta de ejemplo.	68
Tabla 4. Cromosoma de ejemplo.	69
Tabla 5. Matriz Population	72
Tabla 6. Parámetros del algoritmo con su respectiva explicación	73
Tabla 7. Niveles de los parámetros utilizados	87
Tabla 8. Niveles de los parámetros utilizados en la experimentación.	90
Tabla 9. Intervalos de confianza para el factor tamaño de la población y la variable de estudio Cantidad de soluciones no dominadas encontradas	91
Tabla 10. Intervalos de confianza para el factor tamaño de la población y la variable de estudio tiempo computacional.....	92
Tabla 11. Intervalos de confianza para el factor tamaño de la población y la variable de estudio Cantidad de soluciones no dominadas encontradas por minuto	93
Tabla 12. Intervalos de confianza para el factor porcentaje de mutación y la variable de estudio Cantidad de soluciones no dominadas encontradas	98
Tabla 13. Intervalos de confianza para el factor porcentaje de mutación y la variable de estudio tiempo computacional.....	98
Tabla 14. Intervalos de confianza para el factor porcentaje de mutación y la variable de estudio soluciones no dominadas encontradas por minuto.....	99

Tabla 15. Resultados experimentales de las mejores soluciones encontradas para las variables objetivo con los diferentes niveles del parámetro Tamaño de la población.	105
Tabla 16. Costo adicional del mejor nivel de protección ambiental encontrado para cada nivel del tamaño de la población	106
Tabla 17. Comparación del rango de centros de recolección abiertos para un nivel $\alpha = 90$.	107

Lista de figuras

Figura 1. Representación de una frontera de Pareto.	35
Figura 2. Cruce de un punto.....	46
Figura 3. Cruce de dos puntos.	46
Figura 4. Cruce uniforme	47
Figura 5. Representación de la red logística del problema.	55
Figura 6. Pseudocódigo del algoritmo genético propuesto	63
Figura 7. Ejemplo de la codificación del cromosoma, para un solo periodo	68
Figura 8. Ejemplo de la codificación del cromosoma, para t periodos.....	68
Figura 9. Traducción de la variable x	69
Figura 10. Ejemplo para el proceso de traducción de las variables de flujo.....	69
Figura 11. Representación de la variable q_{jkt} (flujo de jeringas transportadas desde la zona del cliente j al centro de recolección k en el periodo t)	70
Figura 12. Representación de la variable z_{knt} (flujo de jeringas transportadas desde el centro de recolección k hasta en centro de incineración n en el periodo t)	71
Figura 13. Pseudo código del Ordenamiento No Dominado Rápido (Fast Non-Dominated Sorting).	75
Figura 14. Pseudo código del cálculo de la distancia de apilamiento (Crowding-distance)....	76
Figura 15. Operador del cruce utilizado en el algoritmo propuesto	78
Figura 16. Operador de mutación utilizado en el algoritmo propuesto.	79
Figura 17. Selección de mejores individuos con la ayuda del ordenamiento no dominado rápido..	80

Figura 18. Gráfica de intervalos de los niveles del tamaño de población para la variable respuesta cantidad de soluciones no dominadas encontradas	92
Figura 19. Gráfica de intervalos de los niveles del Tamaño de Población para la variable respuesta Tiempo Computacional.....	93
Figura 20. Gráfica de intervalos de los niveles del Tamaño de Población para la variable respuesta Cantidad de soluciones no dominadas encontradas por minuto	94
Figura 21. Comparación de los conjuntos de mejores soluciones de todas las réplicas del experimento del factor tamaño de la población.	94
Figura 22. Gráfica de intervalos de los niveles del porcentaje de mutación para la variable respuesta cantidad de soluciones no dominadas encontradas.	98
Figura 23. Gráfica de intervalos de los niveles del porcentaje de mutación para la variable respuesta tiempo computacional.....	99
Figura 24. Gráfica de intervalos de los niveles del porcentaje de mutación para la variable respuesta cantidad de soluciones no dominadas encontradas por minuto	100
Figura 25. Comparación de los conjuntos de mejores soluciones de todas las réplicas del experimento del factor porcentaje de mutación.	100
Figura 26. Ilustración del efecto del nivel de factibilidad.	104

Lista de Apéndices

(Los apéndices están adjuntos en el CD y puede visualizarlos en base de datos de la biblioteca UIS)

Apéndice A. Análisis Bibliométrico.

Apéndice B. Código del algoritmo genético en Matlab.

Apéndice C. Artículo de carácter publicable.

Apéndice D. Datos del modelo.

Apéndice E. Datos de resultados experimentales.

RESUMEN

TÍTULO: MODELO DE OPTIMIZACIÓN DIFUSO MULTIOBJETIVO MULTIPERÍODO PARA EL DISEÑO DE UNA RED LOGÍSTICA INVERSA SOSTENIBLE*

AUTORES: JORGE ESTEBAN CABALLERO RODRÍGUEZ
MILENA PATRICIA AGUILAR GIRALDO**

PALABRAS CLAVE: LOGÍSTICA INVERSA, ALGORITMO GENÉTICO, LÓGICA DIFUSA, MULTI-PERÍODO, MULTI-OBJETIVO, MEDIO AMBIENTE, SOSTENIBILIDAD.

DESCRIPCIÓN:

El presente trabajo de investigación tiene como finalidad aplicar un modelo de optimización difuso, multiobjetivo y multiperíodo para el diseño de una red logística inversa sostenible de una empresa iraní fabricante de agujas y jeringas médicas. El modelo aplicado se caracteriza por integrar de manera simultánea el enfoque económico, medioambiental y social. Este problema ha ganado relevancia en los últimos años debido a la frecuente preocupación del impacto ambiental y la responsabilidad social en las empresas, lo cual ha causado interés en el diseño de redes logísticas que no solo tengan miras en la reducción de costos, sino que también busquen una reducción del impacto ambiental y un aumento de la responsabilidad social, lo anterior con la intención de seguir el camino hacia la sostenibilidad.

La función objetivo relacionada con los costos busca minimizar el valor presente de estos (VPN), el impacto ambiental se cuantifica con base en la evaluación del ciclo de vida (LCA) y la responsabilidad social se mide según el número de trabajos utilizados en la red logística en el horizonte de planeación. La solución del modelo se realiza mediante el uso de un algoritmo genético con base en el NSGA-II adaptado al caso de estudio.

La incertidumbre presente en los parámetros se maneja aplicando lógica difusa. Para validar el método de solución propuesto se realizó un diseño experimental de un factor para tres factores con el fin de determinar su influencia en el algoritmo, además de la calidad y cantidad de las soluciones halladas.

* Trabajo de grado

** Facultad de Ingenierías físico-mecánicas. Escuela de estudios industriales y empresariales. Director: Henry Lamos Diaz

ABSTRACT

TÍTULO: A FUZZY MULTI-OBJECTIVE AND MULTI-PERIOD OPTIMIZATION MODEL FOR SUSTAINABLE REVERSE LOGISTIC NETWORK DESIGN*

AUTORES: JORGE ESTEBAN CABALLERO RODRÍGUEZ
MILENA PATRICIA AGUILAR GIRALDO**

PALABRAS CLAVE: REVERSE LOGISTICS, GENETIC ALGORITHM, FUZZY LOGIC, MULTI-PERIOD, MULTI-OBJECTIVE, ENVIROMENT, SUSTAINABILITY.

DESCRIPTION:

In this work, the application of a fuzzy multi-period and multi-objective optimization model is proposed to design a sustainable reverse logistic network for the case study of a medical syringe manufacturer from Iran. The model used is characterized for simultaneously implementing economical, environmental and social aspects. The proposed problem has gain has gain relevancy in recent years due to the organizations' frequent concern for the environmental impact and social responsibility, which has led to interest in the design of logistics networks that not only take in account the reduction of cost but also the reduction of environmental impacts and the increase of social responsibility.

The cost related objective function minimize the Net Present Value of costs (NPV), the environmental-focused objective function minimize the environmental harm done quantified with a Life Cycle Assessment (LCA) based method, and the social responsibility is measured by the amount of works used by the reverse logistic network in the planning horizon. The proposed mathematical model is solved using a NSGA-II-based Genetic Algorithm, adapted to the case of study.

Uncertainties present in the model parameters are assessed by using fuzzy logic. To validate the proposed solution method, one-factor designs are used to measure the effect of three factors in the algorithm, in addition to the quantity and quality of the solutions.

* Bachelor Thesis

** Faculty of Physico-mechanical Engineering. School of Industrial and Business Studies. Director: Henry Lamos Diaz

Introducción

En la sociedad actual, las empresas son constantemente desafiadas por las exigencias de los clientes que cada vez están más orientadas a la sostenibilidad; entre los principales desafíos se encuentra el manejo del flujo de materiales o productos que son devueltos por los clientes a lo que también podría llamarse logística inversa. Históricamente el estudio de la logística inversa como parte de la gestión de la cadena de suministros ha tenido un enfoque centrado en el rendimiento económico y la reducción de costos, mientras que el término de logística verde nacido por los años 80, presenta un enfoque hacia el uso e implementación de tecnologías que contribuyan a la reducción del impacto ambiental causado por las prácticas logísticas.

El diseño de una red logística inversa sostenible y eficiente debe integrar el enfoque eco amigable de la logística verde, puesto que garantizar el menor costo económico generalmente no va de la mano con la reducción del impacto ambiental y ser eficaz en la recolección de productos no implica que se esté siendo amigable con el medio ambiente; de hecho, de acuerdo con estudios realizados por Henriques y Sadowsky en 1996, las prácticas de logística inversa han causado impactos negativos desde el punto de vista ambiental y social.

Actualmente y debido a razones ya mencionadas con anterioridad, se ha empezado a generar interés en el estudio del diseño de redes de logísticas que no solo sean eficientes sino que también sean sostenibles, esto se logra a partir de la integración de consideraciones tanto económicas como sociales y ambientales, buscando así la optimización de los tres aspectos y no solo de uno de ellos; cabe resaltar que solo en escasas ocasiones se ha tenido en cuenta dentro de las investigaciones el componente de impacto social, por lo tanto, esta consideración podría llamarse emergente.

Con aras de hacer frente a los asuntos mencionados con anterioridad y que en parte son debidos a las prácticas de la logística inversa, se han desarrollado técnicas y modelos matemáticos para que los tomadores de decisiones puedan decidir estratégicamente teniendo en cuenta las métricas y los resultados integrales de cómo deben manejar la red de logística inversa para obtener los mejores resultados en los objetivos planteados.

El presente trabajo tiene como finalidad plantear una solución alterna a la problemática existente en un caso práctico en el que Govindan, Paam y Abtahi basados en Pishvaei y Razmi (2012) realizan el diseño de una red logística inversa para una empresa productora de insumos médicos en Irán y , ellos proponen un modelo matemático de optimización con lógica difusa que luego se soluciona utilizando diversas metaheurísticas adaptadas como el algoritmo MOPSO, mientras que en esta investigación se propone un modelo de optimización difuso utilizando como método de solución un algoritmo genético (AG) basado en el NSGA-II y adaptado a las necesidades del problema, los modelos planteados en ambos casos comprenden no solo los aspectos económicos y ambientales, sino que además incluyen el impacto social.

Para el manejo de la incertidumbre presente en el problema debido a la dinámica del flujo inverso, se utilizaron técnicas fuzzy o difusas en el modelado matemático. El diseño de una red logística inversa sostenible es un tema muy reciente y relativamente inexplorado en Colombia, por lo que se desea con investigaciones de este tipo despertar el interés en estos temas a nivel nacional.

El presente documento se encuentra estructurado de la siguiente forma: en el capítulo 4 se encuentra la revisión de literatura relacionada con el objeto de estudio y en el capítulo 5 se presenta el marco teórico; en el capítulo 6 se encuentra la descripción del problema y la metodología utilizada, en el capítulo 7 está el algoritmo de solución propuesto y en capítulo 8 la aplicación de

dicho algoritmo. En el capítulo 9 está la validación del algoritmo propuesto, en los capítulos 10 y 11 se encuentran las conclusiones y recomendaciones respectivamente.

Tabla 1.

Cumplimiento de los objetivos del proyecto

Objetivos	Cumplimiento
Realizar una revisión de literatura sobre el problema de logística inversa, y su solución mediante metodologías de modelos difusos que permitan trabajar con varios periodos y objetivos.	Numeral 4
Definir el modelo de logística inversa difuso, multiperíodo y multiobjetivo a utilizar con la respectiva caracterización de las variables.	Numeral 6
Desarrollar un algoritmo genético (AG) para la solución del problema de logística inversa con un modelo difuso, multiperíodo y multiobjetivo.	Numeral 7 Numeral 8 Apéndice B
Validar la eficiencia del algoritmo genético haciendo uso de instancias de la literatura, y medir el impacto de los parámetros del AG mediante experimentos diseñados.	Numeral 9
Realizar un artículo académico de carácter publicable basado en el trabajo de investigación realizado.	Apéndice C

1. Planteamiento del problema

El proceso logístico de una empresa asociado a transportes y distribuciones de producto trae consigo ciertas consecuencias que no generan bienestar al ser humano, entre estas podemos encontrar el impacto ambiental generado por los gases producidos por los vehículos de transporte, las sustancias tóxicas o residuos generados por la disposición final e incluso el reciclaje. Además de esto, también se deben considerar los costos acarreados por el proceso para la empresa y la responsabilidad social que debe tener esta para evitar un impacto social negativo sobre la población.

La logística inversa, principal objeto de estudio en este documento, en donde los productos van desde los clientes hasta las organizaciones con la finalidad de reutilización, reciclaje o disposición final y la cual generalmente ha tenido un enfoque económico con miras en la reducción de costos, se queda corta frente a la sociedad actual cada vez más exigente y ambientalista, cabe aclarar que las actividades de logística inversa generalmente traen consigo un efecto negativo en el desarrollo social y deterioro medio ambiental, puesto que la devolución de productos implica un mayor consumo de recursos, además de esto en el ambiente existen varios agentes donde las prácticas de logística de cada uno pueden ir en contravía en lo que se refiere al impacto social y medioambiental; por lo anterior se hace necesario que se incluya dentro de la planeación y optimización de operaciones objetivos de tipo ambiental, donde estos tengan en cuenta la incertidumbre presente en los diferentes escenarios de aplicación.

En Colombia el estado actual de las mayas viales complica el proceso de logística inversa pero a pesar de esto se han puesto en marcha varios proyectos que tienen como objetivo sacarle provecho a los productos devueltos por los clientes, particularmente en el caso de productos

alimenticios próximos a vencer, en el país se pueden encontrar diferentes puntos donde se comercializan exclusivamente productos que cuentan con poca vida útil a un precio más bajo que el del mercado, también existen bancos de alimentos dirigidos a poblaciones con bajos recursos abastecidos con productos que deben ser consumidos de forma inmediata. Lo anterior se hace con la finalidad de recuperar una parte de los costos de producción y transporte, al igual que se busca la mejor manera de darle una disposición final a los productos generando un beneficio para la comunidad.

2. Justificación del Proyecto

La sociedad actual se caracteriza por tener una visión de economía y crecimiento sostenible, esto ha llevado a que las organizaciones empiecen a tener interés en mostrarse como eco amigables y responsables socialmente, incluso existen estudios respecto al efecto que genera el interés de las empresas por ser verdes en el desempeño de estas. Diseñar una red de logística inversa sostenible implica que mientras se lleve a cabo el proceso se minimice el impacto ambiental junto con el costo económico, y a su vez se optimice la responsabilidad social, lo anterior trae consigo la necesidad de contar con modelos matemáticos de optimización que tengan en cuenta la incertidumbre presente en los escenarios de aplicación permitiendo así la reducción del impacto ambiental y los costos asociados a la operación, y que de igual forma se incremente el nivel de responsabilidad social.

Muchos investigadores intentan explorar las relaciones que vinculan la sostenibilidad y la cadena de suministro verde en logística inversa (Govindan, 2015; Rostamzadeh et al., 2015). La mayoría de las investigaciones hechas con anterioridad se han centrado en cuestiones sociales y

ecológicas a través de la cadena de suministro de circuito cerrado, pero rara vez se ha considerado la integración de logística inversa en el diseño de cadenas de suministro sostenibles y verdes, en un enfoque cuantitativo, esta integración puede hacerse agregando algunas variables de decisión, funciones objetivo y restricciones en los modelos matemáticos (Govindan et al., 2015b, Soleimani y Govindan, 2015).

En este trabajo se pretende integrar los problemas ecológicos y de sostenibilidad en la logística inversa, donde la minimización de los costos utilizando el concepto de valor presente neto (VPN), la minimización de los impactos ambientales y la maximización de la responsabilidad social, considerando que esta última aumenta las oportunidades de carrera y reduce los daños en el trabajo, son las funciones objetivo con las cuales se considerarían los tres aspectos de la sostenibilidad.

3. Objetivos

3.1. Objetivo general

Aplicar un modelo de optimización difuso, multiobjetivo y multiperíodo para el diseño de una red logística inversa sostenible.

3.2. Objetivos específicos

- Realizar una revisión de literatura sobre el problema de logística inversa, y su solución mediante metodologías de modelos difusos que permitan trabajar con varios periodos y objetivos.
- Definir el modelo de logística inversa difuso, multiperíodo y multiobjetivo a utilizar con la respectiva caracterización de las variables.
- Desarrollar un algoritmo genético (AG) para la solución del problema de logística inversa con un modelo difuso, multiperíodo y multiobjetivo.
- Validar la eficiencia del algoritmo genético haciendo uso de instancias de la literatura, y medir el impacto de los parámetros del AG mediante experimentos diseñados.
- Realizar un artículo académico de carácter publicable basado en el trabajo de investigación realizado.

4. Revisión de la Literatura

En el presente documento la revisión de literatura se realizó para los tres ejes de estudio fundamentales en los que se basa el problema de diseño de una red logística inversa sostenible por medio de un modelo de optimización difuso multiobjetivo multiperíodo para, dichos ejes son los siguientes: La logística inversa desde un enfoque verde y sostenible, modelos matemáticos aplicados a la logística inversa dentro de la cadena de suministro y por último las aplicaciones de algoritmos y métodos de solución en la logística inversa.

4.1. La logística inversa desde un enfoque verde y sostenible

La logística inversa es un tema que ha entrado en auge en los últimos años debido a que la humanidad ha tomado conciencia de la grave situación que presenta el planeta en lo que respecta al medio ambiente; actualmente los líderes y empresarios no sólo buscan rentabilidad en costos, sino que también encuentran posibilidades de negocio y mejora en el diseño de una red de logística inversa que favorezca tanto a las empresas como al medio ambiente y a la comunidad.

Gómez (2010) afirma que el papel de la logística inversa es gestionar adecuadamente los retornos, desechos y devoluciones de la cadena de suministros, buscando una reducción de los impactos ambientales e intentando desarrollar un enfoque de rentabilidad; con lo que se promueve la recuperación de los recursos desde las perspectivas económicas y ecológicas.

A continuación se mencionan artículos donde se presentan estudios que relacionan la logística inversa con la responsabilidad social y el medio ambiente (red sostenible), por ejemplo, Korchi y Millet (2011) obtuvieron 18 estructuras generales diferentes de una red de logística inversa al

cambiar los lugares de los clientes y los centros de agrupación, almacenamiento y producción, al comparar y analizar estas estructuras, lograron ofrecer unas con mayores ganancias económicas y menores daños ambientales. Chaabane et al. (2012) introdujeron un modelo para diseñar una RL sostenible; el modelo tenía dos funciones objetivo para optimizar los costos y los impactos negativos en el medio ambiente (emisiones sólidas, líquidas y gaseosas que emanan de los procesos de producción y los sistemas de transporte). Pishvaei y Razmi (2012) diseñaron un modelo para una cadena de suministro que buscaba optimizar los costos y los impactos ambientales negativos; para medir los impactos ambientales, emplearon un método llamado evaluación del ciclo de vida.

Pishvaei et al. (2012) afirman que las cadenas de suministro desempeñan funciones importantes en el entorno empresarial actual, la cuestión de la responsabilidad social debería ser considerada cuidadosamente al diseñar y planificar las cadenas de suministro para avanzar hacia la sostenibilidad, por lo anterior, abordaron el problema de diseño de la cadena de suministro siendo esta socialmente sostenible y responsable bajo condiciones inciertas. Desarrollaron un modelo de programación con dos objetivos donde se buscaba minimizar el costo total y maximizar la responsabilidad social empresarial dentro de la cadena.

Lopes Silva et al. (2013) estudiaron logística inversa y desarrollaron un modelo de embalaje retornable para optimizar la generación de residuos y los impactos ambientales mientras se minimizan los costos y el consumo de los recursos, utilizaron datos de campo, mapeo de flujo de RL y evaluación ambiental en su investigación. Ramos et al. (2014) consideraron los aspectos económicos, ambientales y sociales en el diseño de los sistemas RL, resolvieron un problema de enrutamiento de vehículos multiobjetivo con rutas entre depósitos por medio del frente de Pareto obtenido para el estudio del caso, finalmente se revisaron las compensaciones entre los objetivos.

Fahimnia et al. (2015) sugirió un modelo de planificación táctica para la cadena de suministro, que se puede usar para examinar las relaciones entre los costos y los impactos ambientales; desarrollaron un procedimiento de solución conocido como entropía cruzada integrada anidada (NICE), lo aplicaron al modelo matemático no lineal propuesto y se centraron en examinar la asociación entre los intentos lean y las consecuencias ecológicas; encontraron que (1) no todas las implicaciones lean en el nivel táctico resultan en ventajas verdes, y (2) una comparación entre cadenas de suministro lean y centralizadas con cadenas de suministro verdes demuestra que la cadena de suministro verde es una opción más flexible y eficiente .

Govindan et al. (2015) utilizaron un método difuso de objetivos múltiples haciendo que el método indique las relaciones causales entre las prácticas en la gestión de las cadenas de suministro verdes y el rendimiento; encontraron las principales prácticas que podrían tener un impacto positivo en el desempeño de la cadena de suministro en términos de aspectos ambientales y económicos; los resultados revelaron que la gestión, la colaboración, la compra ecológica y la certificación ISO 14001 son las prácticas más importantes en la administración de las cadenas de suministro ecológicas.

Mangla et al. (2015) analizaron los riesgos relevantes para la adopción y la implementación efectiva de las prácticas de la cadena de suministro verde (GSC) desde un punto de vista industrial, los resultados arrojaron que los riesgos operacionales son los de mayor importancia en GSC. Coskun et al. (2015) diseñó la red de la cadena de suministro verde basado en las expectativas ecológicas de los consumidores al proponer un modelo de programación de objetivos considerando tres segmentos de consumidores, es decir, consumidores verdes, incoherentes y rojos, Resolvieron un hipotético problema real de ejemplo para mostrar el valor y la aplicabilidad del modelo sugerido y estudiaron un conjunto de escenarios para ofrecer una idea de cómo el nivel de determinación

del consumidor de ser amigable con el medio ambiente afecta la red de suministro verde; los resultados presentaron una forma de medir las relaciones entre las cadenas de suministro verdes y el comportamiento del consumidor.

Ayvaz et al. (2015) encontraron que razones tanto políticas como económicas, imagen ambientalista y la responsabilidad social influyen en que las organizaciones se vean forzadas a desarrollar estrategias para actualizar sus sistemas, con base en eso, realizaron un estudio donde propusieron el diseño de una red genérica de logística inversa bajo el retorno de cantidades, clasificación de productos según calidad y la incertidumbre del costo de transporte. Para esto presentaron un modelo de programación estocástico multi ϵ , multi producto y restricciones de capacidad, esto con la finalidad de tener en cuenta la incertidumbre.

4.2. Modelos matemáticos aplicados a la logística inversa dentro de la cadena de suministro

Uno de los primeros estudios realizados en pro de mejorar los sistemas utilizados para la logística inversa fue propuesto por Fleischmann et al. (1997), ellos dividieron el tema en tres ejes principales: gestión de producción, inventario y distribución, luego de esto midieron las implicaciones que traía cada eje sobre los requisitos para realizar el proceso de reutilización de las partes devueltas. Minner (2001) combinó en sus estudios las devoluciones de productos tanto internos como externos y la mejora que esto implicaba al problema de control de inventarios dentro de la cadena de suministro, esto lo hizo buscando la minimización de los costos teniendo en cuenta la existencia de un inventario excedente y sus implicaciones, Minner optó por resolver un problema de minimización convexo y concluyó que la reutilización de los productos crea el inventario excedente.

Altıparmak et al. (2006) buscaban ofrecer un plataforma óptima para la gestión eficiente y eficaz de la cadena de suministro; para esto consideraron en el problema de diseño redes de cadenas de suministro los siguientes tres objetivos: La minimización del costo total compuesto por costos fijos, costos de plantas y centros de distribución (DC), como también los costos de distribución entrante y saliente, la maximización del servicio de atención al cliente que pueden prestarse en términos de tiempo de entrega aceptable (cobertura), y por último la maximización del equilibrio de utilización de la capacidad para los países en desarrollo (es decir, capital en los índices de utilización).

Lee y Chan (2009) ofrecieron una red de logística inversa basada en la tecnología de identificación por radiofrecuencia (RFID), donde su objetivo era encontrar un modelo para la optimización de la producción, ellos sugirieron que la RFID podría usarse para contar elementos almacenados en puntos o bodegas de almacenamiento y también para enviar señales al centro de retorno, Chen (2012) luego de resolver el problema de la balanza de reciclaje en la logística inversa, se centró en el equilibrio en condiciones en que los precios del mercado y los flujos de reciclaje tienen efectos interactivos y los flujos de materiales reciclados de entrada y salida en la cadena de suministros no están equilibrados.

Das y Chowdhury (2012) ofrecieron un modelo de reciclaje y logística para varios desechos de productos electrónicos a fin de minimizar los costos totales de procesamiento; su modelo consistía en cuatro fases de reciclaje: recolección, separación, reciclaje y reparación, el sitio final incluía un punto de vertido, mercado primario y mercado secundario; descubrieron que los costos de transporte constituyen una parte importante de los costos de reciclaje, por lo tanto, concluyeron que la reducción de los costos de transporte es la mejor manera de reducir los costos generales del sistema.

Nikolaou et al. (2013) presentó un modelo combinatorio para la responsabilidad social de las empresas y la sostenibilidad en el sistema RL sobre la base de un marco operativo completo. Su marco incluía índices matemáticos combinatorios para la evaluación de la responsabilidad social en RL. Su trabajo puede mejorar el desempeño de las empresas en responsabilidad social a través de los procesos de logística inversa.

Ramezani et al. (2013) desarrollaron un modelo para diseñar un sistema logístico asimétrico de avance y retroceso. Sugirieron un proceso para optimizar la calidad, el beneficio y la capacidad de respuesta del cliente como funciones objetivo en su modelo.

Pereira et al. (2014) consideraron dentro de su plan de logística inversa sostenible tres objetivos donde apuntan al balanceo de costos con las preocupaciones ambientales y sociales en un caso de estudio real de un sistema de recolección de residuos reciclables. El objetivo económico consistió en la recolección de los costos variables en función de la distancia recorrida por los vehículos en el enrutamiento, dentro del modelo no se incluyeron los costos fijos, el objetivo ambiental fue relacionado directamente con las emisiones de CO₂ asociados a la actividad de transporte, el objetivo social se orientó a la minimización del máximo número de horas trabajadas por los conductores buscando así una carga equilibrada en el horizonte de planeación. El problema se modela como de enrutamiento periódico de vehículos con múltiples objetivos y múltiples depósitos y se determinan las soluciones por medio del frente de Pareto.

Suyabatmaz et al. (2014) presentaron dos métodos de modelado para el diseño estocástico del sistema de logística inversa para el manejo de las incertidumbres en el problema. Uno abordó el diseño de una red de distribución y otro consideró el desarrollo de un método genérico basado en una nueva metodología de solución propuesta recientemente.

Ghayebloo et al. (2015) presentó un modelo de programación bi-objetivo para una red de logística asimétrica directa / inversa. Se obtienen múltiples soluciones no dominadas para demostrar la compensación entre los objetivos de beneficio y de verdor. Desarrollaron algunas ideas gerenciales a través de varios experimentos computacionales.

Alshamsi y Diabat (2017) publicaron un documento enfocado en desarrollar una programación lineal de enteros mixtos (MILP), con el objetivo de determinar la ubicación óptima y la capacidad de los nodos importantes de la red RL, como los centros de inspección y las instalaciones de remanufactura, las decisiones de transporte tales como si usar vehículos internos o externos, lo anterior con una base en la rentabilidad. El problema está formulado para el caso de un almacén de electrodomésticos ubicado en el Consejo de Cooperación para los Estados Árabes del Golfo (CCG). Se consideran 68 ciudades, lo que lleva a un gran número de variables y restricciones; por lo tanto, un enfoque heurístico, a saber, un algoritmo genético (GA), se elige para resolver el problema. La principal contribución del documento fue desarrollar una AG muy eficiente capaz de resolver un problema a gran escala en poco tiempo.

4.3. Aplicaciones de algoritmos y métodos de soluciones en la logística inversa

Altıparmak et al. (2006) proponen como método de solución al problema de diseño de redes logísticas con varios objetivos algoritmos genéticos que permiten encontrar el conjunto de soluciones Pareto óptimas, el estudio es validado en una empresa productora de productos plásticos en Turquía. Para lidiar con objetivos múltiples y permitir que el que toma las decisiones evalúe un mayor número de soluciones alternativas, se aplican dos enfoques diferentes de ponderación o de peso en el procedimiento de solución propuesto; los resultados experimentales mostraron que el primer enfoque era capaz de generar un mayor número de soluciones óptimas de Pareto.

Min et al. (2006) introdujeron un modelo no lineal para resolver problemas de logística inversa, su principal objetivo era reducir los costos de los procesos de realizar RL, este modelo puede crear un equilibrio en el inventario y reducir los costos de retención y transporte. Lin et al. (2007) abordaron el problema de logística inversa de una manera flexible mediante el diseño de red por múltiples etapas con una estructura no adyacente, es decir, en el problema escalones de la red que no eran vecinos estaban conectados por arcos puesto que en algunos casos prácticos los arcos de conexión no adyacentes hacían que las redes de logística fueran rentables y adaptables a los cambios de los diferentes escenarios que se pudieran presentar. Formularon el problema como de ubicación – asignación y propusieron un algoritmo genético híbrido como metodología de solución.

Du y Evans (2008) propusieron un modelo para optimizar los costos totales y la tardanza total del tiempo del ciclo. Los resultados de su modelo revelaron que minimizar los costos y los tiempos de ciclo hacen que la red sea una estructura más centralizada.

Kim et al. (2009) introdujeron un programa de localización de rutas para vehículos en Corea del Sur transportando desechos de material electrónico que estaban al final de su ciclo de vida; su objetivo era minimizar las rutas de transporte identificando las cuatro áreas principales de los centros de reciclaje, para resolver este problema se empleó el algoritmo de búsqueda tabú.

Pishvae et al. (2010) propusieron un modelo de programación matemática para una red de logística de períodos múltiples con el fin de disminuir los costos de transporte y los costos fijos. La red de RL incluía lugares de clientes, recolección, control, reciclaje y descarga con capacidad limitada, utilizaron un algoritmo de recocido simulado como procedimiento de solución.

Diabat et al. (2013) diseñaron un problema de red RL de varios pasos para determinar la cantidad y las posiciones de los lugares de recolección primaria, los lugares de retorno

centralizados y el tiempo máximo de conservación para pequeños volúmenes de productos devueltos, presentaron un modelo para reducir los costos en el problema del diseño de la red RL, incluidos los costos de mantenimiento, preparación, transporte, etc. Para resolver este problema, utilizaron algoritmos genéticos y del sistema inmunológico artificial.

Feitó-Cespón et al. (2017) presentaron un problema no lineal entero mixto multiobjetivo estocástico (SMOMINLP) para rediseñar una cadena de suministro inversa sostenible para reciclar ciertos productos. El modelo integra objetivos económicos, ambientales y sociales para respaldar decisiones estratégicas como la ubicación de las instalaciones, el flujo de materiales y el diseño y selección de transportes. El objetivo de impacto ambiental se calcula a través del Ciclo de Vida. Se utilizó una programación multi-criterio con un algoritmo de aproximación para gestionar varios objetivos vinculados con la programación estocástica y así, abordar la incertidumbre. Además, para evaluar las soluciones obtenidas y reducir el efecto de incertidumbre en la toma de decisiones, se propuso un indicador de desempeño. Los resultados experimentales mostraron las configuraciones de la cadena de suministro que mejoran el rendimiento de la sostenibilidad.

5. Marco Teórico

5.1. Logística inversa

Alshamsi & Diabat (2017) definen la logística inversa (RL) como una secuencia de operaciones que comienza en el nivel del consumidor y finaliza en el fabricante, lo opuesto al enfoque tradicional de la cadena de suministro. El reciclaje, la reutilización y el reprocesamiento de productos son actividades de las redes de RL, y son cada vez más frecuentes debido a las crecientes

preocupaciones ambientales y socioeconómicas. La investigación ha comenzado a estudiar tales redes en un esfuerzo por maximizar la eficiencia y mejorar las operaciones.

El campo de RL se ocupa de recuperar un producto al final de la vida útil para el cliente y enviarlo al fabricante a través de una serie de puntos tales como centros de recolección e inspección, con el propósito de deshacerse adecuadamente del producto o recuperar valor adicional del mismo (Prahinski y Kocabasoglu, 2006).

Como se dicen Alshamsi y Diabat (2017), se puede hacer una distinción inicial con respecto al tipo de devoluciones de productos, a saber, los rendimientos al final de la vida útil, los rendimientos comerciales y los rendimientos de fallas (Ardeshirilajimi y Azadivar, 2014). La primera categoría se refiere a productos que han alcanzado su máxima utilización y de los cuales el cliente ya no puede usar o extraer más valor. Las redes de RL para tales productos son frecuentes en la literatura (Fleischmann et al., 1997, Guide, 2000, Guide y Wassenhove, 2001). La segunda categoría comprende productos que el consumidor desea devolver a cambio de otro producto (Guide et al., 2006; Azadivar y Ordoobadi, 2012; Masoud y Azadivar, 2013), mientras que el tercero incluye productos con defectos menores o nulos que, no obstante, se devuelven porque los consumidores se han acostumbrado a devolver cualquier producto incluso por la sospecha de un defecto (Lawton, 2008; Ferguson et al., 2006).

5.2. Logística verde y logística sostenible

Govindan et al., (2016) afirman que el diseño una red de RL que pueda minimizar los costos y, al mismo tiempo, considerar problemas ecológicos y sociales se considera como una red de RL verde y sostenible. No solo las entidades de la cadena de suministro deben ser verdes en sus propias operaciones de logística, sino que también deben cooperar entre sí en lo que respecta a

consideraciones ecológicas. Por lo tanto, de una manera sostenible, los factores ambientales también deben considerarse (Zhou et al., 2000; Björklund et al., 2012; Hervani et al., 2005; Zhang et al., 2015). Además, las empresas deben preocuparse por su responsabilidad social, esto implica que las empresas además de obtener ganancias deben pensar en la conformidad con las necesidades legales, los principios éticos y la estima para las personas y las comunidades en todas sus actividades (Pai et al., 2015). Por lo tanto, el término sostenibilidad se usa cuando el medio ambiente y los factores sociales se tienen en cuenta además de los aspectos económicos.

Seuring y Müller, (2008) definieron la sostenibilidad como "el diseño y el empleo de sistemas humanos, así como de sistemas industriales para utilizar los recursos naturales y para asegurar que los ciclos normales no reduzcan la calidad de vida y las oportunidades económicas futuras", de esta forma tampoco tienen ningún impacto negativo en las condiciones sociales, la salud humana y el ecosistema. "Todas las empresas y sus individuos deben desarrollarse en lo que respecta a la salud y la seguridad de todas las criaturas, a un medio ambiente más limpio y ayudar a garantizar una sociedad equilibrada" (Kotzee y Reyers, 2016).

5.3. Optimización multiobjetivo

La optimización multiobjetivo es usada cuando se quiere evaluar varios criterios simultáneamente, es decir, no posee un único criterio medible por el cual pueda declararse que una solución sea completamente satisfactoria; cuando la región factible sea continua.

Es común que los objetivos a evaluar entren en conflicto unos con otros y esto no permita que haya una única solución que simultáneamente satisfaga a todos, por lo anterior, en estos casos se usa el óptimo de Pareto el cual está basado en el criterio de optimalidad paretiana el cual es definido por Pareto en 1896 citado en (Vitoriano, 2009) como sigue: "Una alternativa es eficiente (óptimo

de Pareto) si toda alternativa que proporcione una mejora en un atributo produce un empeoramiento en al menos otro de los atributos.” De lo anterior se deriva la definición de la alternativa dominada o no eficientes, siendo esta una alternativa para la que existe otra alternativa que presenta mejores atributos.

La Frontera De Pareto se define según (Vitoriano, 2007) como el conjunto de soluciones eficientes no dominadas y esto es lo que generalmente se busca con la programación multiobjetivo, puesto que la única solución o solución de mejor compromiso al final es tomada por el ente decisor.

A continuación, se presenta un ejemplo gráfico de una frontera de Pareto para dos objetivos que quieren ser maximizados.

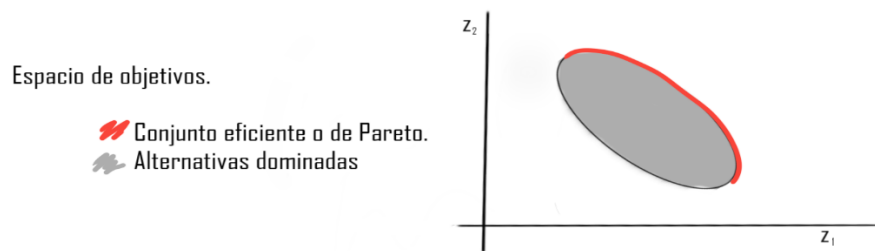


Figura 1. Representación de una frontera de Pareto. Adaptado de Vitoriano, 2007

5.3.1 Métodos de optimización multiobjetivo. Los métodos de optimización multiobjetivo sirven para encontrar el conjunto de soluciones dominantes o frente de Pareto las cuales cumplen unas restricciones y determinan la región factible. Según (Vitoriano, 2007) existen métodos de optimización multiobjetivo para generar la totalidad del conjunto de soluciones eficientes como hay métodos que ofrecen una solución compromiso, dichas técnicas se describen a continuación.

5.3.1.1 Técnicas generadoras del conjunto eficiente. (Vitoriano, 2009) afirma que “son técnicas de carácter mecánico, en las que no se incluyen las preferencias del ente decisor”; Están diseñadas con la finalidad de obtener todo el conjunto eficiente generalmente tras aplicar métodos de programación paramétrica.

5.3.1.1.1 Método de las ponderaciones. Según (Zadeh, 1963) El método consiste que cada objetivo sea multiplicado por un peso o factor no negativo y agregarlos en una única función objetivo. Al variar los pesos se puede obtener todo el conjunto eficiente, resolviendo de esta manera los distintos problemas planteados mediante programación paramétrica. Para aplicar este método se deben haber normalizado con anterioridad los criterios para que de esta manera no influya la diferencia de unidades que tiene cada objetivo.

5.3.1.1.2 Método de las ε -restricciones. Según (Marglin, 1967) este método “consiste en optimizar uno de los objetivos e incorporar el resto como restricciones paramétricas, resolviendo el problema resultante mediante programación paramétrica”. Cabe resaltar que normalmente se optimiza la función objetivo que se considera de mayor importancia.

5.3.1.1.3 Método simplex multiobjetivo. (Yu y Zeleny, 1975) definen este método de la siguiente manera: “Se trata de una extensión del método simplex que evalúa en cada iteración la eficiencia de las soluciones básicas obtenidas (puntos extremos), obteniendo así todos los puntos extremos eficientes”. El conjunto eficiente serán todas las combinaciones lineales convexas de

puntos extremos eficientes que sean adyacentes. Este método sólo es aplicable para objetivos y restricciones lineales.

5.3.1.1.4 Método programación compromiso. La finalidad de la programación compromiso consiste en escoger un punto del frente de Pareto o conjunto eficiente, es de esperarse que para esto se deban tener muy presentes las preferencias e inclinaciones del ente decisor.

Se define el punto o alternativa ideal como un vector que escrito de la manera transpuesta queda de la siguiente manera: $z^*=(z_1, z_2, z_3, \dots, z_n)$ en el cual se recogen los valores óptimos de cada objetivo individualmente tratado (Rangoaga, 2009).

5.3.1.2 Programación por metas: Se basa principalmente en la lógica satisfaciente planteada por Simon en 1955 y citada por Vitoriano en 2009 donde se expresa lo siguiente:

En los contextos actuales de decisión (información incompleta, recursos limitados, conflictos de intereses, ...) el decisor más que optimizar unas funciones objetivo intenta que una serie de metas se aproximen lo más posible a unos niveles de aspiración prefijados (p.117.)

La programación por metas plantea que para cada objetivo el ente decisor tiene una aspiración o meta de manera que al alcanzar dicho nivel de aspiración la solución se considera satisfactoria ([Charnes y Cooper, 1961], [Lee, 1972] e [Ignizio,1976]).

Uno de los principales tipos de programación por metas es el *método lexicográfico*, en este se establecen niveles de prioridad para cada una de las metas planteadas y se resuelve el problema de manera secuencial según el nivel de prioridad dado; de modo de que estas se puedan ordenar según el peso asignado y siempre se priorice la mejor solución para el objetivo catalogado como más

importante, para esto se deben mantener los valores que van siendo obtenidos para metas con mayores niveles de prioridad y fijándolos para la posterior solución de las metas que cuentan con menor prioridad.

5.4. Lógica difusa

(Morales, 2002) dice que “las lógicas difusas procuran crear aproximaciones matemáticas en la resolución de ciertos tipos de problemas. Pretenden producir resultados exactos a partir de datos imprecisos, por lo cual son particularmente útiles en aplicaciones electrónicas o computacionales”. El adjetivo “difuso” aplicado a ellas se debe a que los valores de verdad no-deterministas utilizados en ellas tienen, por lo general, una connotación de incertidumbre.

(Medina, 2006) Afirma que un sistema de lógica difusa o lógica borrosa convierte variables de entrada (cuantitativas y cualitativas) en variables lingüísticas a través de funciones de pertenencia o conjuntos difusos, los cuales son evaluados mediante un conjunto de reglas difusas del tipo si-entonces. Luego las salidas del sistema se convierten en valores nítidos (crisp) mediante un proceso de concreción (de-fuzzyfication), que permiten brindar información para la toma de decisiones. Un sistema de lógica difusa utiliza cualquier tipo de información y la procesa de manera similar que el pensamiento humano; por ello, los sistemas de lógica difusa son adecuados para tratar información cualitativa, inexacta e incierta que permiten tratar con procesos complejos, lo que la hace una alternativa interesante para modelar problemas de toma de decisiones.

5.4.1 Conjuntos difusos. (Morales, 2002) un conjunto difuso es pues una correspondencia (o función) que a cada elemento del universo le asocia su grado de pertenencia, dicho grado se encuentra entre 0 y 1. Enunciada así esta definición parece ser cíclica, mas no lo es: un conjunto difuso es una

función cuyo dominio es el universo y cuyo rango es el intervalo $[0, 1]$. En tanto el grado de pertenencia sea más cercano a 1 tanto más estará el elemento en el conjunto y en tanto el grado de pertenencia sea más cercano a 0 tanto menos estará el elemento en el conjunto.

5.4.2 Número difuso. Un número difuso es definido por Reina y Moscovitz (2008) de la siguiente manera: “un conjunto normalizado y convexo $A \subseteq \mathbb{R}$, cuya función de pertenencia es al menos, continua a trazos y tiene el valor funcional $A(x) = 1$ justo para un elemento” pag.21. Para todo número difuso $\tilde{c}$, el intervalo esperado (IE) y valor esperado (VE) se pueden definir como sigue (Jimenez., et al 2007):

$$IE(\tilde{c}) = [E_1^c, E_2^c] = \left[\int_0^1 f_c^{-1}(x) dx, \int_0^1 g_c^{-1}(x) dx \right]$$

$$VE(\tilde{c}) = \frac{E_1^c + E_2^c}{2}$$

5.4.2.1 Número triangular. Según Ávila (2011) se denomina número difuso triangular, al número borroso real y continuo, tal que la forma de su función de pertenencia determina con el eje horizontal un triángulo. Su función de pertenencia es lineal, a izquierda y derecha, y el α corte para $\alpha=1$ tiene un solo elemento, es decir, la función de pertenencia alcanza el valor uno para un único número real. Un número triangular está determinado de manera única por tres números reales a , b y c , tales que $a \leq b \leq c$, la función de pertenencia y sus α cortes son los siguientes:

$$\mu_{\bar{A}}(x) = \begin{cases} 0 & \text{si } x < a \\ \frac{x-a}{b-a} & \text{si } a \leq x \leq b \\ 1 & \text{si } x = b \\ \frac{c-x}{x-b} & \text{si } b \leq x \leq c \\ 0 & \text{si } x > c \end{cases}$$

$$A_{\alpha} = [a + \alpha * (b - a), c - \alpha * (c - b)]$$

5.4.2.2 Número trapezoidal. Según Ávila (2011) se denomina número difuso trapezoidal al número real continuo, tal que, la forma de su función de pertenencia determina con el eje horizontal un trapecio. Su función de pertenencia es lineal a izquierda y derecha, y el α corte para $\alpha=1$ es un intervalo de números reales. Un número trapezoidal está determinado de manera única por cuatro números reales a, b, c y d , tales que $a \leq b \leq c \leq d$. De igual forma, la función de pertenencia y los α cortes por los cuales se puede representar este número difuso son:

$$\mu_{\bar{A}}(x) = \begin{cases} 0 & \text{si } x < a \\ \frac{x-a}{b-a} & \text{si } a \leq x \leq b \\ 1 & \text{si } b \leq x \leq c \\ \frac{c-x}{x-b} & \text{si } c \leq x \leq d \\ 0 & \text{si } x > d \end{cases}$$

$$A_{\alpha} = [a + \alpha * (b - a), d - \alpha * (d - c)]$$

5.4.3 Sistema de inferencia difuso. Es una forma de representación de datos inexactos y conocimientos de manera similar a como lo hace el pensamiento humano (Jang et al., 1997). Como se cita en (Medina, 2006), los pasos esenciales para el diseño de un sistema difuso según (Jang et al., 1997; Kasabov, 1998; Kosko, 1994) son los siguientes:

1. Identificación del tipo de problema y el tipo de sistema difuso que mejor se ajusta a los datos.
2. Definición de variables de entrada y salida, sus valores difusos y sus funciones de pertenencia (parametrización de variables de entrada y salida).
3. Definición de la base de conocimiento o reglas difusas.
4. Obtención de salidas del sistema mediante la información de las variables de entrada utilizando el sistema de inferencia difuso, el cual utiliza operadores de composición.
5. Traslado de la salida difusa del sistema a un valor nítido o concreto mediante un sistema de conversión.
6. Ajuste del sistema validando los resultados.

5.5. Algoritmos de optimización multiobjetivo evolutivos

Son métodos de optimización basados en la evolución biológica, dónde se mantiene un conjunto de soluciones o “individuos” que se cruzan y compiten entre sí, de tal manera que solo los mejores individuos van quedando a lo largo del tiempo, evolucionando hacia mejores soluciones cada vez. (Chaiyaratana et al., 1997). Los métodos de este tipo tienen una estructura muy similar entre sí, diferenciándose en el manejo del elitismo (la selección de la mejor población resultante de la iteración del algoritmo), la estrategia de evaluación (la manera de evaluar el conjunto de individuos) y el mecanismo de reproducción (la selección de individuos para la generación de las nuevas generaciones) (Von Lücken et al., 2004).

5.5.1 Strength Pareto Evolutionary Algorithm (SPEA). Zitzler y Thiele (1998) proponen el uso del algoritmo multiobjetivo SPEA que integra diferentes algoritmos evolutivos con el objetivo de

combinar lo mejor de ellos. En este algoritmo, se usa un archivo de soluciones no dominadas que se han obtenido anteriormente. En cada generación se copian los individuos no dominados de la población a este archivo, y se calcula un valor de Strength o Fortaleza proporcional al número de soluciones que un individuo domina. Este valor de Fortaleza se utiliza para estimar el nivel de fitness o adaptación del individuo.

El conjunto de soluciones no dominadas se va depurando, dejando solo las soluciones no dominadas respecto al conjunto, es decir, si un individuo es dominado por otro en el conjunto, el individuo dominado se elimina. Sobre este conjunto no dominado se aplican técnicas de clustering que mejoran la calidad de las soluciones, pero disminuyen la eficiencia computacional (Zitzler et al., 1991).

5.5.2 Algoritmo Genético Multiobjetivo. Es un algoritmo de optimización multiobjetivo evolutivo (MOEA) inspirado en procesos de evolución natural y evolución genética. Consiste en la iteración de una serie de transformaciones o generaciones sobre un conjunto de soluciones o población. Estas generaciones residen en la aplicación de operadores de selección, cruce y de mutación sobre la población, haciendo que de generación en generación la población vaya presentando una mejora. (Gupta y Ghafir, 2012). Este algoritmo ha sido usado ampliamente en la solución de problemas multiobjetivo de todo tipo dada su flexibilidad y relativa facilidad de adaptación a las necesidades del problema (Chaiyaratana et al., 1997). La definición de algoritmo genético dada por Koza (1992) se presenta a continuación:

Un algoritmo matemático altamente paralelo que transforma un conjunto de objetos matemáticos individuales con respecto al tiempo usando operaciones modeladas de acuerdo al principio Darwiniano de reproducción y supervivencia del más apto, y tras

haberse presentado de forma natural una serie de operaciones genéticas de entre las que destaca la recombinación sexual. Cada uno de estos objetos matemáticos suelen ser una cadena de caracteres (letras o números) de longitud fija que se ajusta al modelo de las cadenas de cromosomas, y se les asocia con una cierta función matemática que refleja su aptitud (p.819).

La técnica emula la evolución natural para explorar con eficiencia el espacio de búsqueda con el supuesto de que unos individuos con ciertas características son aptos para sobrevivir y transmiten estas características a su descendencia (Sait y Youssef, 1999). Los algoritmos genéticos operan sobre una población o conjunto de soluciones representadas como cadenas que pueden ser binarias o numéricas también llamadas cromosomas. Durante la ejecución, el algoritmo cruza los individuos de mayor aptitud para renovar la población y elimina los de menor aptitud. Al final, el cromosoma de mayor aptitud es la solución al problema (Metaxiotis y Psarras, 2004). Algunos conceptos básicos son:

5.5.2.1 Cromosoma. Se le denomina así a cualquier solución potencial de un problema y esta puede ser representada dándole valores a una serie de parámetros, generalmente se codifica en una cadena binaria o numérica que representa un individuo o solución, donde cada elemento en la cadena se conoce como gen (Pose, 2000).

5.5.2.2 Población. Es un conjunto finito de cromosomas que constituyen la población de cada generación. No es aconsejable tener un tamaño pequeño de población puesto que es probable que no se abarque de la mejor manera el espacio de búsqueda, pero tampoco es recomendable trabajar con poblaciones muy grandes puesto que puede causar un alto costo computacional (Moujahid et

al., 2008).

5.5.2.3 Aptitud. Criterio que evalúa la calidad de un cromosoma. A mayor aptitud, mejor la solución y mayor la probabilidad de que sobreviva y transmita sus características a su descendencia (Davis, 1991).

5.5.2.4 Selección. Esta operación es la encargada de escoger a los individuos que tienen más oportunidades de reproducirse, se les otorga un mayor número de oportunidades de reproducción a aquellos que son más aptos. Sin embargo, no debe descartarse la opción de reproducir a los individuos menos aptos para en futuras generaciones evitar la homogeneidad de la población o, dicho de otra forma, garantizar la diversidad de esta (Pose, 2000). Los principales criterios de selección se explican a continuación.

5.5.2.4.1 Selección por ruleta. Este método también conocido como selección de Montecarlo, consiste en asignar a cada individuo de la población una parte proporcional a su nivel de adaptación, de tal forma que al sumar todos los porcentajes se obtiene la unidad, como es de esperarse, los individuos más aptos recibirán una porción más grande que aquellos que se encuentran dentro de los peores. Para seleccionar un individuo es suficiente generar un número aleatorio ente 0 y 1, luego se devuelve el individuo que estaba ubicado en esa posición de la ruleta; la posición suele obtenerse al hacer un recorrido por los individuos de la población y acumular sus proporciones de ruleta hasta que la suma sea superior al valor obtenido aleatoriamente. La selección por ruleta se caracteriza por su sencillez, pero se vuelve ineficiente en la medida que

aumenta el tamaño de la población, además los peores individuos pueden ser seleccionados más de una vez (Pose, 2000).

5.5.2.4.2 Selección por torneo. Este método busca realizar la selección por medio de comparaciones directas entre los individuos, el método puede ser determinístico o probabilístico. En el primer caso se selecciona al azar un número p de individuos, generalmente $p=2$ y de entre los individuos seleccionados se escoge el más apto para pasarlo a la siguiente generación. La versión probabilística se diferencia de la anterior en que para escoger el individuo ganador se genera un número aleatorio entre 0 y 1, si este es mayor que un parámetro p fijado para todo el proceso evolutivo que y que generalmente toma valores entre 0,5 y 1, se escoge el individuo más apto y en caso contrario el menos apto (Pose, 2000).

5.5.2.5 Cruce. El cruce es una estrategia de reproducción sexual en la que luego de ser seleccionados los individuos que serán padres, estos se recombinan para producir la descendencia de la generación siguiente. Se debe definir la probabilidad de cruce y esta suele estar cerca del 90% (Pose, 2000), la probabilidad indica la frecuencia con la que se producen cruces entre los cromosomas padres, si no existe reproducción, los hijos serán copias exactas de los padres (Arranz y Parra, 2007). Los tipos de cruce más empleados son los siguientes:

5.5.2.5.1 Cruce de un punto. Esta técnica siendo la más sencilla de todas, consiste en que una vez seleccionados los padres se cortan sus cromosomas en un punto seleccionado de forma aleatoria para generar de esta manera dos segmentos diferenciados en cada uno de ellos (la cabeza

y la cola), luego se intercambian las colas entre los dos individuos para generar a los nuevos descendientes (Pose, 2000). El procedimiento se ilustra en la Figura 2:

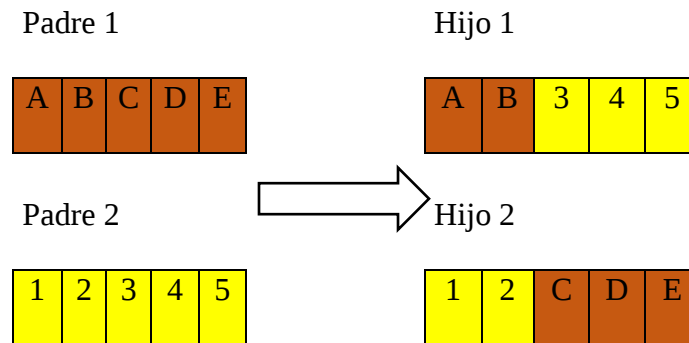


Figura 2. Cruce de un punto. Adaptado de Pose (2000)

5.5.2.5.2 *Cruce de dos puntos.* En este caso el cromosoma de los padres se corta en dos puntos diferentes y ninguno de estos debe coincidir con los extremos del cromosoma para de esta manera garantizar la generación de tres segmentos. El descendiente se forma escogiendo el segmento central del primer padre y los segmentos laterales del segundo padre (Pose, 2000). El procedimiento se ilustra en la Figura 3:

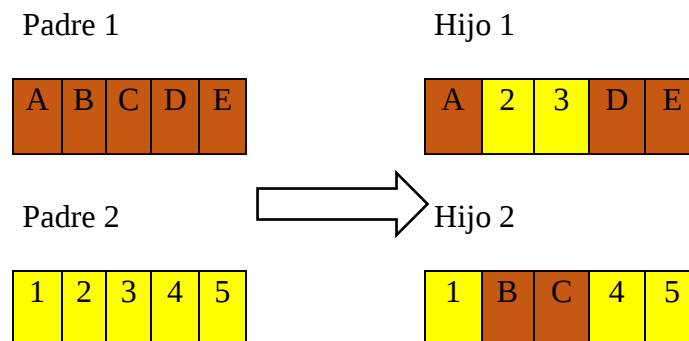


Figura 3. Cruce de dos puntos. Adaptado de Pose (2000)

5.5.2.5.3 Cruce uniforme. Al usar esta técnica, cada gen de la descendencia tiene las mismas posibilidades de pertenecer a uno u otro padre. Implica la generación de una máscara de cruce con valores binarios, si en una de las posiciones de la máscara hay un 1, el gen situado en esa posición en uno de los descendientes se copia del primer padre. Si por el contrario hay un 0 el gen se copia del segundo padre. Para producir el segundo descendiente se intercambian los papeles de los padres, o bien se intercambia la interpretación de los unos y los ceros de la máscara de cruce. El número efectivo de puntos de cruce es fijo, pero será por término medio $L/2$, siendo L la longitud del cromosoma (Pose, 2000). Este procedimiento se ilustra en la Figura 4:

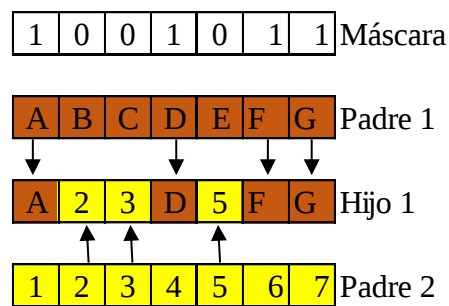


Figura 4. Cruce uniforme. Adaptado de Pose (2000)

5.5.2.6 Mutación. La mutación suele utilizarse de manera conjunta con el operador de cruce, si el cruce es exitoso, al menos uno de los descendientes se muta con cierta probabilidad provocando que alguno de los genes del cromosoma varíe de forma aleatoria. Este procedimiento ocurre con probabilidades muy bajas y se realiza debido a que cuando se genera la descendencia siempre se produce algún tipo de error en el paso de la carga genética de los padres a los hijos (Pose, 2000).

5.5.3 NSGA-II (Non-Dominated Sorting Genetic Algorithm). Es un tipo de algoritmo genético que utiliza el enfoque de la dominancia y la distancia de apilamiento (la densidad de la

zona en que se encuentra una solución) para la selección de la población y la generación de la descendencia (Deb et al., 2002).

Durante cada corrida del algoritmo se aplican operadores genéticos (cruce y mutación). Para la selección de individuos por estos operadores, se realiza un torneo binario y se escogen los individuos con menor rango, es decir, el menos dominante de la pareja o con mayor distancia de apilamiento (en una zona con menos soluciones) (Das et al., 2007). Se unen los individuos de la población antes y después de los operadores genéticos y se escogen los mejores, de acuerdo con la dominancia (Deb et al., 2002). Los individuos escogidos pasarán a la siguiente generación y repetirán el proceso iterativo, el algoritmo genera una serie de corridas y une las mejores soluciones de cada corrida que de hecho pertenecen a la última generación en un conjunto Q. Se definen como soluciones de Pareto aquellas soluciones no dominadas del conjunto Q, es decir, a las mejores soluciones que se encontraron (Das et al., 2007). Algunos conceptos importantes para el NSGA-II son:

5.5.3.1 Dominancia. La dominancia simboliza que una solución es mejor que otra. En los algoritmos de optimización multi objetivo, la dominancia hace referencia a que una solución es estrictamente mejor que otra (Deb et al. 2002). Algunos casos pueden ser que una solución “a” presenta mejores valores en todos los objetivos que una solución “b”, dado que “a” y “b” sean soluciones factibles (Chaiyaratana et al., 1997), o que una solución “a” sea factible mientras que una solución “b” no lo sea (Das et al. 2007). En los casos anteriores, “a” es una solución que domina a “b”. En algoritmos genéticos, la dominancia suele medirse con la medida de adaptación del algoritmo, dado que mayor adaptación mayor dominancia (Deb et al., 2002).

5.5.3.2 Elitismo. El elitismo en los MOEA podría definirse como la acción de seleccionar individuos no dominados o con mejor adaptabilidad para pertenecer a la siguiente generación del algoritmo. El conjunto de soluciones no dominadas se conoce también como un conjunto élite, y va de la mano con la solución de Pareto de los problemas multi objetivos: La frontera de Pareto es, por su definición, el conjunto elite o no dominado de todo el espacio factible (Zitzler y Thiele, 1998). El elitismo es una herramienta de los algoritmos que permite una convergencia más rápida y un mejor desempeño computacional (Das et al., 2007).

5.5.3.3 Ordenamiento no dominado rápido. Es un procedimiento descrito en Deb et al. (2002) para el NSGA-II, en el cual la población se ordena de acuerdo con su dominancia; se ordenan todos los individuos del más dominante al menos dominante, para todo el tamaño de la población de la siguiente manera: En una primera instancia, se separa la población en dos: las soluciones factibles y las no factibles. Para las soluciones factibles, se define la dominancia con la función fitness, mientras que para las soluciones no factibles se define según la magnitud de violación de las restricciones. De acuerdo con lo anterior, se generan subconjuntos llamados rangos en los cuales se encuentran las soluciones con el mismo nivel de dominancia, es decir, soluciones que no se dominan entre sí. Los individuos del rango 1 son las mejores soluciones, mientras que el último rango generado tendrá las peores.

5.5.3.4 Rango. Es uno de los resultados del ordenamiento no dominado y hace referencia al conjunto conformado por las soluciones en el orden de dominancia, es decir las soluciones que se encuentran en el rango 1 dominan a todas las demás soluciones y se les llama soluciones no dominadas, en el rango 2 se encuentran las soluciones dominadas por el rango 1 y que a su vez

dominan a todas las demás y así sucesivamente para los demás rangos (Deb et al., 2002).

5.5.3.5 Distancia de apilamiento. Consiste en obtener una estimación de la densidad de soluciones que rodean una solución particular de la población, para esto se requiere clasificar la población según el valor de cada función objetivo según su magnitud en forma ascendente, a las soluciones con la función más pequeña y más grande se les asigna un valor de distancia infinita y a todas las demás soluciones intermedias se les asigna un valor igual a la distancia normalizada absoluta en función de los valores de dos soluciones adyacentes. El valor total de la distancia de apilamiento se calcula como la suma de los valores de las distancias individuales correspondientes a cada objetivo (Deb et al., 2002).

La fórmula para hallar la distancia de apilamiento del individuo i para M objetivos de aquellas soluciones que no se encuentran en los extremos es la siguiente:

$$Distancia_i = \sum_{m=1}^M \frac{Distancia\ al\ individuo\ siguiente_{mi} - Distancia\ al\ individuo\ anterior_{mi}}{máximo\ del\ objetivo_m - mínimo\ del\ objetivo_m}$$

Para todo i en $[2, I-1]$

5.6. Valor presente neto.

De acuerdo con Varela (2010), el VPN es una herramienta de evaluación de proyectos a largo plazo, permitiendo el cálculo del valor actual del proyecto. Esta herramienta también puede llamarse CPN si se utiliza para el cálculo del valor presente de los costos. Para calcular el VPN, se utiliza la siguiente fórmula:

$$VPN = \sum_{t=1}^n \frac{[Flujo\ de\ dinero_t]}{(1+i)^t}$$

Donde t representa el periodo en el tiempo con una periodicidad dada en el que se encuentra el flujo de dinero, n la cantidad total de periodos e i representa la tasa de interés. El interés debe estar en la misma periodicidad que t , es decir, si t representa periodos separados por 3 años, el interés debe darse para un periodo de tiempo de tres años.

5.7. Eco-Indicador 99

El Eco-Indicador 99 es una metodología basada en la evaluación del ciclo de vida (Life Cycle Assessment, o LCA) utilizada para cuantificar los efectos del impacto ambiental e integrar las políticas de desarrollo sostenible en los modelos matemáticos utilizando unidades amigables para el usuario, permitiendo así tener en cuenta el medio ambiente en toda decisión tomada en el contexto de aplicación (Pishvaei y Razmi, 2012). La metodología de este indicador permite aplicar los resultados de un proceso de diseño LCA sin la complejidad, tiempo y análisis posterior que suele requerir el LCA (Goedkoop et al, 2000).

El primer paso consiste en delimitar el alcance del modelo determinando todos los procesos que forman parte del problema de estudio. Para el segundo paso, se debe definir el ciclo de vida del producto o proceso el cual consiste en las actividades que se realizan durante toda la vida útil de este, por ejemplo, producción, transportes, actividades de uso, actividades de manejo de los desechos del producto, entre otras. Luego, se debe definir y cuantificar el flujo de cada tipo de material durante el ciclo de vida (por ejemplo, metales o plásticos). En el cuarto paso se calcula el eco indicador multiplicando los valores obtenidos por los valores del indicador dados por los creadores del método, se normalizan por proceso y se ajustan a una unidad del producto o proceso,

por ejemplo, para el caso de estudio del presente trabajo, se toman los kilogramos de metal que serían incinerados y se multiplica por el indicador de incineración de metales dado en milipuntos por kilogramo; acto seguido, se consolidan los impactos ambientales de todas las actividades de incineración en el ciclo de vida del proceso o producto y se normalizan de acuerdo de la utilización de metal de una unidad del producto o proceso) (Pishvaei y Razmi, 2012).

Este indicador tiene tres perspectivas: de jerarquía, individualista e igualitaria basados en la teoría cultural. Para el cálculo del eco-indicador en el presente trabajo se utilizan datos que se basaron en la perspectiva de jerarquía. En esta perspectiva, se asume que los impactos a la salud humana, la calidad del ecosistema y el agotamiento de los recursos contribuyen en un 40%, 40% y 20% respectivamente para la puntuación generada por el eco-indicador 99 (Ministry of Housing, Spatial planning and the Environment, 2000). El indicador puede tener valores negativos o positivos, siendo el valor positivo un impacto ambiental dañino y el valor negativo una prevención de un daño ambiental.

6. Recolección de datos y Metodología

6.1. Definición del problema

El problema que se trata en la presente investigación está basado en un caso de la vida real, (Pishvaei y Razmi, 2012). El problema consiste en la red de la cadena de suministro de una empresa iraní que fabrica jeringas y agujas médicas, luego de ser usadas una única vez son desechadas. La empresa satisface cerca del 70% de la demanda doméstica y recientemente también

abastece pedidos importantes a dos países vecinos; se producen en total aproximadamente 600 millones de jeringas al año. Cabe resaltar que, al finalizar el ciclo de vida de las jeringas médicas se genera un impacto ambiental negativo considerable, según la Organización Mundial de la Salud (OMS, 2005) se realizan 16 billones de inyecciones en todo el mundo cada año; se estima que jeringas de uso médico contaminadas no esterilizadas y reutilizadas causan entre 8 y 16 millones de casos de hepatitis B, entre 2,3 y 4,7 millones de caso de hepatitis C y entre 80 y 160 mil casos de VIH. Por lo tanto, se puede afirmar que las jeringas y agujas de uso médico son desechos peligrosos por su alto potencial de propagar infecciones y enfermedades tan solo con una perforación o herida accidental; lo anterior hace que el manejo del fin del ciclo de vida de las jeringas o disposición final de estas sea un asunto de vital importancia y criticidad (Hanson y Hitchcock, 2009).

Con la finalidad de aminorar los peligros originados por las jeringas contaminadas, la OMS proporcionó un documento guía sobre cómo tratar con este tipo de desechos, particularmente los cortopunzantes como agujas y jeringas, en este se sugiere el uso de cajas de seguridad para mantener las agujas que han sido utilizadas y catalogadas como riesgo biológico (OMS, 2003). Para manejar el fin de la vida útil de las jeringas y agujas usadas existen los tres siguientes métodos: Incineración, no incineración y reciclaje.

En la incineración suele usarse el incinerador de cemento y el incinerador de horno rotatorio, este método es uno de los más usados para los fines deseados en este documento debido al bajo costo, uso conveniente y la capacidad de recuperación de energía, pero la incineración tiene un alto impacto ambiental negativo debido a la contaminación del aire y además está clasificado como una de las principales fuentes de emisiones (OMS, 2005), (Hanson y Hitchcock, 2009).

En cuanto a los métodos de no incineración se pueden encontrar los autoclaves de vapor y la desinfección por microondas los cuales buscan la esterilización y desinfección de los residuos médicos, son soluciones ecológicas de las cuales también se puede recuperar energía y su costo es competitivo frente al del método de incineración. Con el método de reciclaje de jeringas y agujas médicas consideradas como peligrosas, se presenta que de manera general se considera como una actividad prohibida (Zhao et al., 2009). Sin embargo, es relevante mencionar que estudios realizados en India bajo la supervisión de la OMS sobre la posibilidad de reciclar las agujas y jeringas médicas muestran que esta opción es viable siempre y cuando se maneje de forma conjunta con procesos de esterilización y desinfección (OMS, 2005), (Pishvaei y Razmi, 2012). El presente trabajo se enfoca en los métodos de incineración y reciclaje teniendo en cuenta las consideraciones dadas en el caso de estudio.

El proceso de logística inversa planteado se ilustra en la Figura 5 donde al cumplir las jeringas y agujas médicas el fin del ciclo de su vida, son llevadas por los clientes a los centros de recolección, estando ahí, las jeringas pasan por el proceso de separación de las partes de metal y plástico, luego de ser separadas, las partes son dirigidas a los centros de reciclaje de plástico y centros de reciclaje de metal según corresponda. Las jeringas que definitivamente no pueden ser recicladas son enviadas los centros de incineración para darle una correcta disposición final.

En cuanto a la cantidad de jeringas devueltas por los clientes porque han cumplido su ciclo de vida, se calcula con un porcentaje predefinido que es multiplicado por la demanda total de jeringas que fueron adquiridas en un principio por los clientes. Se asume que se recogen de cada cliente todas las jeringas que han sido utilizadas y/o desperdiciadas.



Figura 5. Representación de la red logística del problema. Adaptado de Govindan et al. (2016)

Los temas de a tratar en este documento y que hacen referencia a la red de logística inversa planteada en condiciones de incertidumbre, consisten en determinar cuántos centros de recolección abrir y la ubicación de estos, así como también los flujos de jeringas que serán transportados a los diferentes tipos de instalaciones según el método que se defina para la disposición final de estas. Lo anterior se plantea mediante tres funciones objetivo que son: (1) la minimización de el valor presente neto de los costos totales, (2) la minimización del impacto ambiental total y (3) la maximización de la responsabilidad social.

Para la resolución del problema se cuenta con datos históricos, pero nada garantiza que en el futuro el comportamiento de los parámetros cumpla con el patrón dado por los datos históricos (Pishvaei y Razmi, 2012). Para el manejo de la incertidumbre presente en los parámetros obtenidos, estos se presentan como números difusos triangulares según las distribuciones de probabilidad estimadas por expertos de los datos insuficientes y el conocimiento de los tomadores de decisiones (Zadeh, 1978).

Tomando de base el análisis del ciclo de vida (LCA por sus siglas en inglés), el método llamado Eco-Indicador 99 ha sido utilizado para cuantificar y modelar la segunda función objetivo que

consiste en la minimización del impacto ambiental total generado por la red logística inversa sostenible (Goedkoop y Spriensma, 2000). El LCA es un procedimiento que permite a los investigadores cuantificar la carga ambiental y los posibles impactos de todo el ciclo de vida de procesos, actividades y productos (Rebitzera et al., 2004), (Pishvae et al., 2012), sin embargo, este proceso es costoso, complicado y demanda mucho tiempo, dado lo anterior, el Eco-Indicador 99 es un método que basado en el LCA no presenta los problemas mencionados anteriormente (Goedkoop y Spriensma, 2000).

“El Eco-indicador 99 es capaz de agregar los resultados del LCA en unidades fácilmente comprensibles y fáciles de usar llamadas Eco-indicadores. Los eco-indicadores son números que expresan la carga ambiental total de un proceso o producto” (Pishvae et al., 2012 pag.4). Debido a la estandarización y norma de los eco-indicadores, los diseñadores pueden comparar las alternativas de diseño de acuerdo con el impacto ambiental total. Este método incluye tres categorías de daños: (1) salud humana, (2) calidad del ecosistema y (3) recursos (Ministry of Housing, Spatial planning and the Environment, 2000).

Para cuantificar y modelar los factores relacionados con la responsabilidad social que en este documento hace referencia a las oportunidades de trabajo creadas por la red de logística inversa sostenible, se tuvieron en cuenta las consideraciones de (Pishvae et al., 2012).

6.2. Obtención de datos

Los datos utilizados para la resolución del problema se pueden encontrar en el Apéndice D, estos fueron obtenidos de (Pishvae et al., 2012) y provienen del caso de estudio mencionado con anterioridad que consiste en el diseño de la red logística inversa de un fabricante iraní que produce jeringas y agujas médicas.

6.3. Formulación del modelo matemático.

Índices

j	Zonas fijas de los clientes, $j=1 \dots J$
k	Candidatos a centro de recolección, $k=1 \dots K$
l	Zonas fijas de reciclaje de metales, $l=1, \dots, L$
m	Zonas fijas de reciclaje de plástico, $m=1 \dots M$
n	Zonas fijas de incineración, $n=1, \dots, N$
t	Periodo, $t=1 \dots T$

Variables de decisión

q_{jkt} Flujo de jeringas (unidades) transportadas desde la zona del cliente j al centro de recolección k en el periodo t .

v_{klt} Flujo de partes de metal de las jeringas (unidades) transportadas desde el centro de recolección k al centro de reciclaje de metales l en el periodo t .

w_{kmt} Flujo de partes de plástico de las jeringas (unidades) transportadas desde el centro de recolección k al centro de reciclaje de plásticos m en el periodo t .

z_{knt} Flujo de jeringas (unidades) transportadas desde el centro de recolección k al centro de incineración n en el periodo t .

$x_{kt0} \begin{cases} = 1 & \text{si el centro de recolección } k \text{ se abre en el periodo } t \\ = 0 & \text{c. c.} \end{cases}$

Parámetros

Nota: Los parámetros con negrilla y tilde ($\sim$) son números difusos.

- $\tilde{d}_{jt}$ Demanda de jeringas de la zona de clientes j en el periodo t .
- $\tilde{i}$ Tasa de descuento en el periodo de planeación.
- $\tilde{r}_{jt}$ Tasa de devoluciones de jeringas del cliente j en el periodo t .
- $\tilde{g}_{kt}$ Costo fijo de abrir un centro de recolección k en el periodo t .
- a_{jkt} Costo de transportar una jeringa desde la zona de clientes j al centro de recolección k en el periodo t .
- b_{klt} Costo de transportar la parte metálica de una jeringa desde el centro de recolección k al centro de reciclaje de metales l en el periodo t .
- h_{kmt} Costo de transportar la parte plástica de una jeringa desde el centro de recolección k al centro de reciclaje de plástico m en el periodo t .
- o_{knt} Costo de transportar una jeringa desde el centro de recolección k al centro de incineración n en el periodo t .
- $\tilde{\beta}_{lt}$ Costo de procesar una jeringa en el centro de reciclaje de metales l en el periodo t .
- $\tilde{\tau}_{mt}$ Costo de procesar una jeringa en el centro de reciclaje de plásticos m en el periodo t .
- $\tilde{\theta}_{nt}$ Costo de procesar una jeringa en el centro de incineración n en el periodo t .
- $\tilde{c}k_k$ Capacidad máxima del centro de recolección k .
- $\tilde{c}l_l$ Capacidad máxima del centro de reciclaje de metales l .
- $\tilde{c}m_m$ Capacidad máxima del centro de reciclaje de plásticos m .
- $\tilde{c}n_n$ Capacidad máxima del centro de incineración n .
- $\tilde{w}c_{kt}$ Cantidad de trabajos ocupados al abrir un centro de recolección k en el periodo t .
- ei_{jk}^{cc} Impacto ambiental (milipuntos/unidad) de transportar cada jeringa de la zona de clientes j al centro de recolección k .

ei_{kl}^{cs} Impacto ambiental (milipuntos/unidad) de transportar cada jeringa del centro de recolección k al centro de reciclaje de metales l .

ei_{km}^{cp} Impacto ambiental (milipuntos/unidad) de transportar cada jeringa del centro de recolección k al centro de reciclaje de plásticos m .

ei_{kn}^{ci} Impacto ambiental (milipuntos/unidad) de transportar cada jeringa del centro de recolección k al centro de incineración n .

ei_l^{st} Impacto ambiental (milipuntos/unidad) de procesar cada jeringa en el centro de reciclaje de metales l .

ei_m^{pl} Impacto ambiental (milipuntos/unidad) de procesar cada jeringa en el centro de reciclaje de plásticos m .

ei_n^{in} Impacto ambiental (milipuntos/unidad) de procesar cada jeringa en el centro de incineración n .

Funciones objetivo

Objetivo 1: Min(F1)

$$F1 = [\sum_k g_{kt} * x_{kt} + \sum_j \sum_k (a_{jkt}) * q_{jkt} + \sum_k \sum_l (b_{klt} + \tilde{\beta}_{lt}) * v_{klt} + \sum_k \sum_m (h_{kmt} + \tilde{\tau}_{mt}) * w_{kmt} + \sum_k \sum_n (o_{knt} + \tilde{\theta}_{nt}) * z_{knt}] / (1 + \tilde{t})^t$$

La función objetivo consta de los componentes: costos fijos de abrir un centro de recolección y costos acarreados por el procesamiento y transporte de las jeringas en los diferentes centros de reciclaje e incineración que componen la red de logística inversa. Los costos de transporte de las jeringas hacia los diferentes centros son calculados multiplicando el costo de transportar una

unidad entre dos centros por la cantidad de jeringas transportadas entre estos. Cabe resaltar que el cálculo se realizó aplicando la fórmula económica del Valor Presente Neto.

Objetivo 2: Min(F2)

$$F2 = \sum_t \left[\sum_j \sum_k ei_{jk}^{cc} * q_{jkt} + \sum_k \sum_l (ei_l^{st} + ei_{kl}^{cs}) * v_{klt} + \sum_k \sum_m (ei_m^{pl} + ei_{km}^{cp}) * w_{kmt} + \sum_k \sum_n (ei_n^{in} + ei_{kn}^{ci}) * z_{knt} \right]$$

La función objetivo orientada a minimizar el impacto ambiental está basada en el concepto de ciclo de vida (LCT por sus siglas en inglés) usado generalmente para medir el impacto ambiental generado por los diferentes sistemas y el análisis del ciclo de vida (LCA por sus siglas en inglés) es utilizado para cuantificar dichos impactos (Rigamonti et al., 2016). Por medio de este método se puede asignar valores numéricos a los factores que afectan al medio ambiente y de esta forma poder medir los efectos potenciales durante el ciclo de vida del producto. El eco-indicador 99 es una metodología que consolida los resultados obtenidos por el LCA, es decir, los transforma en índices ambientales, estos son valores numéricos que identifican explícitamente los aspectos ambientales relacionados con un producto o proceso. Este método considera tres áreas de daño que consisten en la salud humana, la calidad del ecosistema y los recursos. Estos indicadores se miden en milipuntos por unidad, y permiten comparar los impactos ambientales de las actividades. Un valor positivo del indicador implica un impacto ambiental negativo, mientras que un valor negativo del indicador implica prevención de daños ambientales.

Objetivo 3: Min(-F3)

$$F3 = \sum_t \sum_k (\widetilde{w}c_{kt} * x_{kt})$$

La función objetivo mide el impacto social, está basada en el concepto de la utilización de empleos durante el periodo. La función mide los trabajos que utilizan o generan los centros de recolección al tomarse la decisión de abrirse en un periodo. Por medio de esta función, se busca incentivar el uso de la mayor cantidad posible de trabajos en cada periodo, apoyando al desarrollo social de las zonas en las cuales se decida abrir el centro de recolección y generando así una mejoría en la sociedad (Veleva y Ellenbecker, 2001).

Conjunto de restricciones:**De recolección**

$$\sum_k q_{jkt} \geq \tilde{d}_{jt} * \tilde{r}_{jt} \text{ para todo } j, t \quad (1)$$

Esta restricción garantiza que se recojan todos los elementos que entregan los clientes, en este caso son las jeringas desperdiciadas o utilizadas. El lado derecho de la restricción se llamará a través del documento como la demanda, dado que esta es la magnitud requerida por la red logística inversa.

De balanceo de flujo

$$\sum_j q_{jkt} = \sum_m w_{kmt} + \sum_n z_{knt} \quad \text{para todo } k, t \quad (2)$$

$$\sum_m w_{kmt} = \sum_l v_{klt} \quad \text{para todo } k, t \quad (3)$$

La restricción (2) garantiza que la cantidad de jeringas que entran al centro de recolección sea igual a la cantidad que salen, la restricción (3) se encarga de balancear el flujo de las partes plásticas y metálicas de las jeringas que se reciclan puesto que dichas cantidades deben ser iguales.

De capacidad

$$\sum_j q_{jkt} \leq x_{kt} * \widetilde{c\bar{k}}_k \quad \text{para todo } k, t \quad (4)$$

$$\sum_k v_{klt} \leq \widetilde{c\bar{l}}_l \quad \text{para todo } l, t \quad (5)$$

$$\sum_k w_{kmt} \leq \widetilde{c\bar{m}}_m \quad \text{para todo } m, t \quad (6)$$

$$\sum_k z_{knt} \leq \widetilde{c\bar{n}}_n p \quad \text{para todo } n, t \quad (7)$$

Las restricciones de capacidad garantizan que ningún elemento de la red logística inversa supere su máxima capacidad.

De las variables de decisión

$$w_{kmt}, z_{knt}, q_{jkt}, v_{klt} \geq 0 \quad \text{para todo } j, k, l, m, n, t \quad (8)$$

$$x_{kt} \in \{0,1\} \quad \text{para todo } k, t \quad (9)$$

La restricción (8) garantiza la no negatividad de las variables mientras que la (9) hace referencia a la condición binaria de la variable que define la apertura de un centro de recolección.

A continuación, se describe el algoritmo numérico de resolución del problema planteado

7. Algoritmo de solución propuesto

El algoritmo propuesto es una variación del NSGA-II planteado por (Deb et al., 2002) y desarrollado en MATLAB por (Baskar et al., 2015), adaptándolo a las necesidades del problema

descrito en este trabajo, por ejemplo, la existencia de variables binarias o el uso de un conjunto de elite. El ciclo general de operaciones del algoritmo utilizado se presenta a continuación:

Algoritmo Genético Modificado basado en el NSGA-II	
1.	Inicio
2.	Inicialización de parámetros
3.	Repetir para cada iteración
4.	Generar individuos aleatoriamente.
5.	Traducir y evaluar los individuos en las funciones objetivo
6.	Evaluar la factibilidad de los individuos
7.	Ordenamiento no dominado rápido de los individuos
8.	Repetir para cada generación
	Si la población converge de acuerdo con los criterios definidos (75% o más en el rango 1)
9.	
10.	Reemplazar el 75% de la población aleatoriamente
11.	Fin
12.	Generar una copia de la población
13.	Selección aleatoria de padres a cruzar
14.	Aplicar el operador genético de cruce a la copia de la población
15.	Aplicar el operador de mutación al resultado del punto 14
16.	offspring=Unir cromosoma con los resultados de 14 y 15
17.	Traducir y evaluar offspring en las funciones objetivo
18.	Evaluar la factibilidad de los individuos de offspring
19.	Ordenamiento no dominado rápido (offspring)
20.	Selección de los mejores individuos de offspring.
21.	Se guarda el resultado de la generación (mejores individuos de offspring)
22.	Resultados=Resultados U Mejores individuos de offspring
23.	Se realiza un ordenamiento no dominado de 22.
24.	Se actualiza Resultados solo con los individuos no dominados
25.	Fin
26.	Q_{corrida} =Mejores individuos de la corrida
27.	Fin
28.	Ordenamiento no dominado rápido (Q)
29.	Graficar los valores no dominados de Q
30.	Fin

Figura 6. Pseudocódigo del algoritmo genético propuesto

El pseudo-código presentado en la

Figura 6 conforma a grandes rasgos una corrida del algoritmo propuesto. Este algoritmo sigue de manera general la estructura del NSGA-II clásico propuesto en Deb. et al (2002), con algunas

modificaciones realizadas con el fin de la preservación de la diversidad y la aplicación de elitismo para el problema tratado. Las modificaciones son:

1. La inclusión de un operador de inyección bajo un criterio de convergencia.
2. Utilización de un conjunto élite con las soluciones no dominadas encontradas.
3. Mantener un porcentaje de población entre iteraciones
4. El uso de un conjunto que incluye la población, la población después del operador de cruce y la población después del operador de mutación

El operador de inyección consiste en sustituir parte de la población con población generada aleatoriamente si el algoritmo converge (Gupta et al.,2012). Durante el desarrollo del algoritmo, hubo la necesidad de promover una mayor medida la diversidad de la población, dado que por la estructura de la codificación de la solución el algoritmo tendía a converger a mínimos locales. Por ello, se decidió usar un criterio de convergencia para la aplicación del operador de inyección que consiste en lo siguiente: si un porcentaje de la población son individuos no dominados, se sustituye de manera aleatoria un porcentaje fijo de la población. Este procedimiento resulta útil para solventar el problema de convergencia a mínimos locales, sin embargo, presenta el inconveniente de afectar y reemplazar a individuos de la población que convergen al frente de Pareto. Para evitar este inconveniente, se introduce el uso de un conjunto élite similar al utilizado en el SPEA, en el cual se guardan las soluciones antes de ser reemplazadas.

El conjunto élite es un conjunto que contiene las soluciones no dominadas que ha encontrado el algoritmo. En cada generación, se agregan a este conjunto las soluciones no dominadas encontradas y se extraen las soluciones que sean dominadas por otros integrantes del conjunto. Si una solución es una solución de Pareto, dicha solución no abandonará el conjunto dado que no puede ser dominada por otra. Así, cuando el algoritmo converja, la utilización de este conjunto

permite que se preserven las soluciones ya encontradas y se exploren nuevas soluciones. De esta manera, se aprovecha la desventaja que trae consigo el operador de inyección para explorar la frontera de Pareto, dado que permite la exploración de la frontera como mecanismo de diversidad, es decir, si se converge al frente, las soluciones ya encontradas se guardan en el conjunto elite y se procede a buscar nuevas soluciones.

El conjunto Q presentado en el pseudo código se genera por la unión de los conjuntos elite de cada iteración y no por la población final de la iteración como ocurre en el NSGA-II clásico (Deb et al., 2002). Este ajuste es realizado para tener coherencia en la inclusión del conjunto elite y la inyección en el algoritmo, dado que la última generación de la iteración no siempre contendrá solamente soluciones no dominadas, mientras que el conjunto elite sí.

De manera similar al conjunto elite, entre iteraciones mantiene un porcentaje de la población final de la generación anterior, la población se escoge según la distancia de apilamiento, prefiriendo la mayor distancia, de esta manera se garantiza que se escojan las soluciones extremas, dando un punto de partida a la iteración que incentive la exploración en la zona cercana a las soluciones extremas encontradas. De esta forma, las siguientes iteraciones tendrán un mecanismo que les permite continuar la búsqueda en la cercanía de las anteriores iteraciones, mejorando cada vez la cantidad y calidad de las soluciones.

Adicionalmente, los operadores de cruce y de mutación y sus probabilidades respectivas se escogen de tal manera que se incentive la diversidad de la población, y por esta misma razón se amplía el conjunto utilizado en el NSGA-II (Deb et al, 2002) incluyendo el resultado del cruce. La intención de este cambio es permitir que la mayor cantidad de soluciones encontradas puedan entrar al conjunto elite, incluyendo las poblaciones de los pasos intermedios. Si bien, la población antes y después del operador de mutación es similar entre sí y podría dar lugar a que tras aplicar

el elitismo exista una pérdida en la diversidad al escoger individuos similares, los mecanismos de generación de diversidad anteriormente descritos corrigen este inconveniente.

A continuación, se describe el desarrollo del algoritmo de manera más específica.

7.1. Codificación de las soluciones

Las soluciones se representan en paralelo mediante un cromosoma y una matriz *population*. El cromosoma tiene la codificación de las variables de decisión, mientras que la matriz *population* contiene los datos requeridos por el algoritmo tales como el valor de cada función objetivo de cada cromosoma.

Cabe destacar que el proceso de traducción del cromosoma a las variables de decisión garantiza el cumplimiento de las tres primeras restricciones que son las siguientes: la demanda de cada cliente es atendida por un centro de recolección, el flujo que entra a un centro de recolección es equivalente al que sale y, por último, el flujo que sale hacia los centros de reciclaje de plástico es equivalente al que sale hacia los centros de reciclaje de metal.

A continuación, se explica en mayor detalle el cromosoma, su proceso de traducción y el contenido de la matriz *population*.

7.1.1 Cromosoma. El cromosoma diseñado para el problema codifica la ruta que sigue la demanda de cada cliente a través de la red en cada periodo; para ello, para cada cliente j se determina el centro de recolección que lo atenderá y su destino (a qué centro de reciclaje de plástico m , centro de reciclaje de metal l y/o centro de incineración n irá), además de un porcentaje que simboliza la división del flujo entre los centros de reciclaje e incineración. La estructura se puede ver en la siguiente tabla:

Tabla 2.

Representación de la ruta que sigue la demanda utilizada en el cromosoma.

Cliente j en el periodo t				
Centro de recolección k que recogerá la demanda del cliente j	% de la demanda del cliente j que irá a algún centro de reciclaje	Centro de reciclaje de plástico m que atiende la demanda del cliente j	Centro de reciclaje de metales l que atiende la demanda del cliente j	Centro de incineración n que atiende la demanda del cliente j
k	$\%$	m	l	n

La estructura mostrada en la Tabla 2 codifica explícitamente el flujo de material que cada cliente genera, diciendo que la totalidad de su demanda es atendida por un centro de recolección k dado, y a partir de allí se clasifica y envía un porcentaje del material a centros específicos de reciclaje de metales y de plástico (l y m respectivamente) mientras que el resto del flujo generado por el cliente se envía al centro de incineración n . Por ejemplo, si el cliente j en el periodo t es atendido por el centro de recolección 3, y de allí, la mitad de lo recolectado se envía al centro de incineración 2, mientras que la mitad restante se reparte en cantidades iguales entre el centro de reciclaje de plástico 1 y el centro de reciclaje de metales 3, se representaría en el cromosoma de la siguiente manera:

Tabla 3.

Representación de la ruta de ejemplo.

Cliente j en el periodo t				
k=3	%=0.5	m=3	l=1	n=2

Esta estructura se repite horizontalmente tantas veces como clientes haya y cada una de ellas representa a un cliente, como se muestra en la figura:

Cliente 1					Cliente 2					Cliente 3					...	Cliente j				
k1	%1	m1	l1	n1	k2	%2	m2	l2	n2	k3	%3	m3	l3	n3	...	kj	%j	mj	lj	nj

Figura 7. Ejemplo de la codificación del cromosoma, para un solo periodo

Verticalmente, la estructura se repite para todos los t periodos que se tengan:

	Cliente 1					Cliente 2					Cliente 3					...	Cliente j				
Periodo=1	k11	%11	m11	l11	n11	k21	%21	m21	l21	n21	k31	%31	m31	l31	n31	...	kj1	%j1	mj1	lj1	nj1
Periodo=2	k12	%12	m12	l12	n12	k22	%22	m22	l22	n22	k32	%32	m32	l32	n32	...	kj2	%j2	mj2	lj2	nj2
...																...	...				
Periodo=t	k1t	%1t	m1t	l1t	n1t	k2t	%2t	m2t	l2t	n2t	k3t	%3t	m3t	l3t	n3t	...	kjt	%jt	mjt	ljt	njt

Figura 8. Ejemplo de la codificación del cromosoma, para t periodos.

Para ilustrar la traducción del cromosoma a las variables de resultado se hacen las siguientes suposiciones: Existen tres clientes, hay tres opciones para los centros de recolección, de reciclaje de plástico, de reciclaje de metal e incineración ($k=l=m=n=3$); se consideran dos períodos y la

demanda en todos los periodos para cada cliente es de 1000. El cromosoma descrito en el ejemplo es el siguiente:

Tabla 4.

Cromosoma de ejemplo.

Periodo	Cliente 1						Cliente 2						Cliente 3					
t=1	1	0.5	1	2	3	2	0	2	3	1	3	1	3	2	1			
t=2	1	0.5	1	2	3	2	1	2	3	1	1	1	1	2	3			

Para la traducción a la variable x (decisión de abrir o no un centro de recolección), se revisan las casillas correspondientes al centro de recolección y se asigna el valor de 1 si en el periodo dado, se utilizó ese centro de recolección, se asigna 0 en caso contrario:

1	0.5	1	2	3	2	0	2	3	1	3	1	3	2	1
1	0.5	1	2	3	2	1	2	3	1	1	1	1	2	3

Periodo	Decisión de abrir el centro de recolección k		
	$k=1$	$k=2$	$k=3$
t=1	1	1	1
t=2	1	1	0

Las casillas marcadas son aquellas que deben revisarse para la traducción de x

Figura 9. Traducción de la variable x

Para la traducción de las variables de flujo se recurre a la ruta dada, por ejemplo, para traducir la variable del flujo de cada cliente a cada centro de recolección, se revisa qué centro de recolección atiende a cada cliente en cada periodo. Por ejemplo, si se toma el cliente 1 en el periodo 1, se miraría la casilla marcada de azul en la siguiente figura:

Periodo	Cliente 1						Cliente 2						Cliente 3					
1	1	0.5	1	2	3	2	0	2	3	1	3	1	3	2	1			
2	1	0.5	1	2	3	2	1	2	3	1	1	1	1	2	3			

Figura 10. Ejemplo para el proceso de traducción de las variables de flujo.

Dado que la demanda del cliente 1 en el periodo 1 se atiende desde el centro de recolección 1, se reemplaza 1000 en la casilla correspondiente. A continuación, se presenta el resultado final de la traducción de toda la variable, y se encuentra marcado de azul el valor correspondiente al generado por la casilla señalada de azul en la Figura 10:

t=1				t=2			
Centro de recolección				Centro de recolección			
Cliente	k=1	k=2	k=3	Cliente	k=1	k=2	k=3
j=1	1000	0	0	j=1	1000	0	0
j=2	0	1000	0	j=2	0	1000	0
j=3	0	0	1000	j=3	1000	0	0

Figura 11. Representación de la variable q_{jkt} (flujo de jeringas transportadas desde la zona del cliente j al centro de recolección k en el periodo t)

En la Figura 10, el elemento marcado de azul determina el origen (centro de recolección en el que se genera la salida), mientras que lo marcado en verde permite conocer el destino y realizar el cálculo de la magnitud del flujo. Para calcular el flujo de material que va al centro de incineración se tiene en cuenta el porcentaje del material que va a ser reciclado y la demanda de cada cliente mediante el siguiente procedimiento:

Flujo del Centro de recolección 1 al centro de incineración 3 en el periodo 1 (z_{131})

$$= z_{131} + (1 - \%) * \text{Demanda recolectada en el centro 1 del cliente 1 en el periodo 1}$$

$$= 0 + (1 - 0.5) * 1000 = \mathbf{500}$$

Este proceso se realiza para cada cliente y cada periodo. Para los centros de reciclaje, se toma el porcentaje (y no el complemento, es decir, uno menos el porcentaje como se realizó en el ejemplo anterior). El cálculo del flujo de material que va a los centros de recolección de plástico (w_{kmt}) se realiza de la misma manera puesto que la jeringa se divide en dos partes, la de metal y la de plástico; por lo tanto, se puede decir los flujos z_{knt} y w_{kmt} son iguales para un mismo k . La traducción de la variable z_{knt} basada en la ecuación mostrada anteriormente quedaría de la siguiente manera:

		t=1					t=2		
		Centro incineración					Centro de incineración		
Centro de recolección		$n=1$	$n=2$	$n=3$	Centro de recolección		$n=1$	$n=2$	$n=3$
$k=1$		0	0	500	$k=1$		0	0	500
$k=2$		1000	0	0	$k=2$		0	0	0
$k=3$		0	0	0	$k=3$		0	0	0

Figura 12. Representación de la variable z_{knt} (flujo de jeringas transportadas desde el centro de recolección k hasta en centro de incineración n en el periodo t)

De manera análoga se traducen el resto de las variables relacionadas con el flujo de material que parte de los centros de recolección a los centros de reciclaje de partes plásticas y metálicas.

7.1.2 Matriz Population. La estructura de la matriz population para dos objetivos utilizada se puede observar en la siguiente tabla:

Tabla 5.

Matriz Population

Cromosoma	Función Objetivo 1	Función Objetivo 2	Índice de Violación de Restricciones	Rango	Distancia de Apilamiento
1	f1	g1	0	1	∞
2	f2	g2	2	4	∞
3	f3	g3	0,8	2	∞
4	f4	g4	1,5	3	∞

La primera columna de la matriz representa el cromosoma generador de los valores y seguidamente se encuentra el valor de la función objetivo respectiva. Las últimas tres columnas representan el índice de violación de restricciones que es una medida de factibilidad, el rango que es una medida de dominancia y la distancia de apilamiento que mide la densidad de soluciones de la zona en que se encuentra cada individuo.

7.2. Parámetros del algoritmo

Además de los parámetros delimitados por el modelo matemático y los datos históricos del problema, se deben definir parámetros del algoritmo tales como la probabilidad de mutación o el nivel de confianza de las desigualdades. A continuación, se listan dichos parámetros con su respectiva explicación.

Tabla 6.

Parámetros del algoritmo con su respectiva explicación

Parámetro	Explicación
α	Nivel de confianza de cumplimiento de las desigualdades.
<i>popSize</i>	Tamaño de la población (Número de individuos tomados al final de cada generación).
<i>gen_max</i>	Número total de generaciones por iteración.
<i>max_runs</i>	Número total de iteraciones del algoritmo.
<i>Prob_Mut</i>	Probabilidad de mutación.
<i>Prob_Cruce</i>	Probabilidad de cruce.
<i>%Conver</i>	Porcentaje máximo aceptable de convergencia.
<i>%Sust</i>	Porcentaje de la población que se sustituye si se cumple el criterio de convergencia.
<i>%Legacy</i>	Porcentaje de la población final de una iteración que pasará a ser parte de la población inicial de la siguiente.

7.3. Generación de la población inicial

En la primera iteración del código, la población inicial consiste en dos elementos: el cromosoma y una matriz llamada *population* que contiene la información respectiva a los valores de las funciones objetivo, el índice de violación de restricciones, el frente, la distancia de apilamiento y el cromosoma respectivo que representa. Para la generación del cromosoma, se fijan los valores mínimos y máximos que cada casilla puede tener (LI y LS, respectivamente) y se generan números aleatorios de tal forma que cada casilla tenga un valor dentro del rango posible. Para la generación de *population*, se genera una matriz de tamaño [*popSize*, 6] y a la primera columna, se le asigna el valor del vector 1: *popSize* (es decir, se enumera de 1 hasta el tamaño de la población). Acto seguido, se evalúan todos los cromosomas, y para cada uno en la segunda y tercera columna se ubica el valor respectivo de cada función objetivo.

Luego, se evalúa el índice de violación de restricciones o I.V.R. (Deb et al. 2002), el cual consiste en calcular la magnitud de la violación de cada restricción para cada cromosoma, esto se realiza normalizando los valores obtenidos de acuerdo al máximo de toda la población y sumando estos valores normalizados de cada restricción para cada cromosoma, obteniendo así un índice que representa qué tanto está violando las restricciones cada cromosoma en una escala que permite una comparación justa de la factibilidad; cabe resaltar que para ser factible el índice debe ser igual a 0; y una solución que viole en menor magnitud las restricciones que otra debe tener un menor índice. El valor de este índice oscila entre 0 y el número total de restricciones del problema, tal como se resume en la siguiente ecuación.

$$\text{Si el índice de violación de restricciones (I.V.R.)} \begin{cases} = 0, \text{ la solución es factible} \\ > 0, \text{ la solución es no factible} \end{cases}$$

$$I.V.R. \in [0, \# \text{ Restricciones}]$$

Después, se utiliza la función de ordenamiento no dominado rápido en la cual se generan y ordenan los frentes de acuerdo con la dominancia y la factibilidad, además de calcular la distancia de apilamiento. En el ordenamiento no dominado rápido (Deb et al. 2002), se compara la dominancia de cada individuo sobre el resto de la población. La metodología usada modifica la definición de dominancia levemente para ajustarse al tratamiento de las restricciones. Una solución (a) domina una solución (b) si:

1. La solución a es factible y la solución b no lo es.
2. Ambas soluciones son no factibles, pero a tiene un menor índice de violación de restricciones.
3. Ambas soluciones son factibles y a domina a b respecto las funciones objetivo evaluadas.

El efecto de usar estos principios de dominancia basados en las restricciones es que toda solución factible tendrá un mejor frente asignado que una solución no factible, además de eliminar la necesidad de la utilización de un parámetro de penalización. Las soluciones factibles se ordenan de acuerdo con su dominancia en los diferentes frentes, mientras que las no factibles se ordenan de acuerdo con el índice de violación de restricciones. El pseudo código del ordenamiento no dominado es el siguiente:

Ordenamiento no dominado rápido

Para cada individuo (p) en la población (P)
 $S_p = \emptyset$
 $n_p = 0$
Para cada individuo (j) en la población (P)
Si p domina a j entonces
 $S_p = S_p \cup \{j\}$
Se agrega j al set de soluciones dominadas por p
Caso contrario
 $n_p = n_p + 1$
Se incrementa el contador de dominancia de p
Si $n_p = 0$ entonces
Rango de p = 1
 $F_1 = F_1 \cup \{p\}$
 $i = 1$
Mientras $F_i \neq \emptyset$
 $Q = \emptyset$
Para cada individuo p en F_i
Para cada individuo j en S_p
 $n_q = n_q - 1$
Si $n_q = 0$ entonces
Rango de q = i + 1
 $Q = Q \cup \{q\}$
 $i = i + 1$
 $F_i = Q$

Figura 13. Pseudo código del Ordenamiento No Dominado Rápido (Fast Non-Dominated Sorting). Adaptado de Deb et al., 2002

Después de realizar el ordenamiento no dominado, se calcula la distancia de apilamiento (crowding distance) de cada individuo.

Cálculo de la distancia de apilamiento o crowding-distance

Inicio

Se extrae el tamaño de cada frente f

Para cada individuo i en el frente f , inicializar la distancia en 0

Para cada objetivo m

Se ordena el frente de acuerdo con el objetivo m

Se fija la distancia de los extremos (máximo y mínimo) en infinito.

Para $i=2$ hasta $f-1$

Se consolida la distancia de apilamiento total de cada individuo

$$D(i) = D(i) + (D(i+1).m - D(i-1).m) / (f_{\max}.m - f_{\min}.m)$$

Figura 14. Pseudo código del cálculo de la distancia de apilamiento (Crowding-distance).

Adaptado de Deb et al., 2002

En las siguientes iteraciones del algoritmo, se mantiene el 10% de la población de la generación anterior y se genera aleatoriamente el 90% restante siguiendo todos los demás pasos anteriormente descritos. Adicionalmente, el 10% de la población que se mantiene se escoge de acuerdo con la distancia de apilamiento, prefiriendo las soluciones que presenten una mayor distancia. Esta metodología garantiza la selección de las soluciones extremas y soluciones que se encuentren en las zonas menos pobladas, ayudando al algoritmo a buscar soluciones en zonas menos exploradas.

7.4. Criterio de selección

En cuanto a la selección de los padres, esta se realiza por medio de una selección por torneo binario. Para ello, se escogen parejas de individuos al azar de toda la población, y entre ambos individuos, se elige uno siguiendo las siguientes dos reglas:

1. Se escoge el individuo que se encuentre en el mejor rango, es decir, el menor.

2. Si ambos tienen el mismo rango, se elige el individuo con mayor distancia de apilamiento, es decir, el individuo que se encuentra en una región menos poblada.

El anterior proceso se realiza hasta que el torneo haya escogido tantos individuos como hay en la población. A la población seleccionada se le aplicará el operador de cruce y se guarda el resultado de dicha operación. Después, al resultado del operador de cruce se le aplica el operador de mutación, y se guardan los resultados.

7.5. Operador de cruce.

Para este algoritmo, por la naturaleza del cromosoma, se elige un cruce uniforme donde cada gen o elemento de los padres tiene igual probabilidad de ser aplicado a la descendencia. Para ello, se utiliza una matriz auxiliar de ceros y unos que se genera aleatoriamente, dicha matriz representa los elementos de cada padre que serán pasados a los hijos. Se ilustra el procedimiento a continuación:

Máscara binaria auxiliar

1	1	1	1	0	1	0	0	1	0	1	0	0	1	1
0	1	1	1	0	1	1	0	0	0	0	1	0	0	1

Padre 1

2	0	2	1	1	3	0	3	2	2	3	0.5	1	2	3
2	0	2	1	1	3	0	3	1	3	1	0.5	2	2	1

Padre 2

1	0.5	1	2	3	2	0	2	3	1	1	3	3	2	1
1	0.5	1	2	3	2	2	2	3	1	1	1	2	2	3



Descendiente 1

2	0	2	1	3	3	0	2	2	1	3	3	3	2	3
1	0	2	1	3	3	0	2	3	1	1	0.5	2	2	1

Descendiente 2

1	0.5	1	2	1	2	0	3	3	2	1	0.5	1	2	1
2	0.5	1	2	1	2	2	3	1	3	1	1	2	2	3

Figura 15. Operador del cruce utilizado en el algoritmo propuesto

7.6. Operador de mutación:

Se escoge aleatoriamente un elemento del cromosoma y con una probabilidad dada, se reemplaza dicho elemento por un valor dentro del rango posible. Por ejemplo, el porcentaje puede ser cualquier valor entre 0 y 1, con dos cifras decimales. Estos valores fueron discretizados como

se mencionó anteriormente, resultando en un valor entero entre un mínimo y un máximo (esto es una distribución uniforme discreta). Para mayor claridad, se ilustra el procedimiento a continuación.

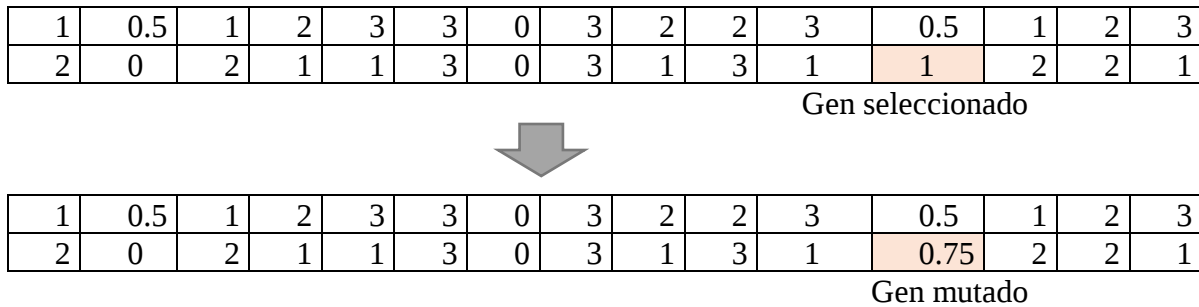


Figura 16. Operador de mutación utilizado en el algoritmo propuesto.

Este operador se aplica 5 veces cuando se realiza la mutación, es decir, se escogen 5 puntos aleatoriamente de cada cromosoma y se mutan con una probabilidad dada.

7.7. Elitismo (selección de mejores individuos)

El procedimiento anterior generó tres conjuntos: La población original, la población a la que se le aplicó el cruce y la población a la cual se le aplicó la mutación luego de haberse realizado el cruce (P_o , P_c y P_m respectivamente) El procedimiento general del NSGA-II consiste en concatenar las tres poblaciones y aplicar el ordenamiento no dominado descrito en la generación de la población inicial. Luego de tener la población ordenada, se procede a extraer los mejores frentes hasta haber seleccionado *popSize* individuos. Si al seleccionar un frente, el conteo de individuos seleccionados excede *popSize*, se escogen individuos de dicho frente hasta completar el tamaño de

población deseado. El resultado es un conjunto con los mejores individuos de las tres poblaciones, es decir, la población descendiente (Pd). El proceso se puede ver ilustrado en la siguiente figura:

Selección mejores individuos: Ordenamiento no dominado rápido

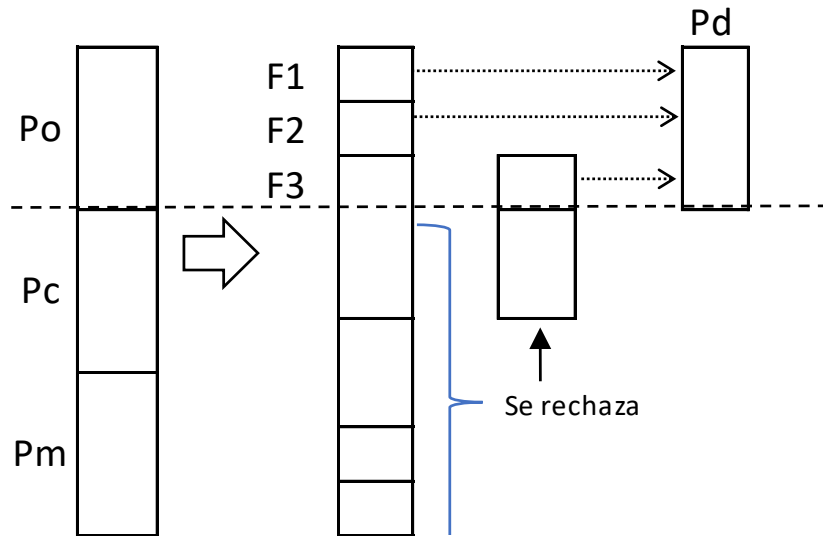


Figura 17. Selección de mejores individuos con la ayuda del ordenamiento no dominado rápido.

8. Aplicación del algoritmo para la solución del problema

8.1. Supuestos.

En el presente trabajo, se toman en cuenta los siguientes supuestos para la solución del problema:

1. La demanda de cada cliente es atendida exclusivamente por un solo centro de recolección.
2. El flujo de material que genera cada cliente es atendido por un único centro de procesamiento de residuos de cada tipo (reciclaje metales, reciclaje plástico o incineración).
3. Se considera que un nivel de confianza aceptable es 90%.

4. El horizonte de planeación hace referencia a tres años (es decir, cada valor de t representa tres años). Todos los parámetros están ajustados a dicho horizonte.

8.2. Transformación del modelo matemático

Para facilitar la resolución del problema se convierte el modelo Fuzzy en un modelo Crisp, para esto, se usa el método del trabajo de Jiménez (Jimenez et al., 2007). En este modelo, se asume que los números difusos poseen forma ya sea trapezoidal o triangular. El modelo se basa en la definición del “valor esperado” y el “intervalo esperado” de un número difuso; transformando las desigualdades usando un criterio alfa que define el grado de factibilidad en que dicha desigualdad se cumple, es decir, el grado de confiabilidad en que un extremo de la desigualdad es mayor al otro, además de modificar la función objetivo teniendo en cuenta que los números difusos que se tienen son de distribución triangular. Este método permite transformar el modelo fuzzy a un modelo auxiliar crisp manteniendo la linealidad del problema y siendo eficiente computacionalmente.

Asumiendo que $\tilde{c}$ es un número difuso triangular, la ecuación (1) puede definirse como la función de pertenencia a $\tilde{c}$:

$$\mu_{\tilde{c}}(X) = \begin{cases} f_c(x) = \frac{x - c^{pes}}{c^{mos} - c^{pes}}, & si \quad c^{pes} \leq x \leq c^{mos} \\ 1, & si \quad x = c^{mos} \\ g_c(x) = \frac{c^{opt} - x}{c^{opt} - c^{mos}}, & si \quad c^{mos} \leq x \leq c^{opt} \\ 0, & si \quad x \leq c^{pes} \text{ o } x \geq c^{opt} \end{cases} \quad (1)$$

en la cual, c^{mos} , c^{pes} y c^{opt} son los tres puntos prominentes de la distribución (el valor más probable, el valor más optimista y el valor más pesimista, respectivamente). Las ecuaciones (2) y (3) definen el intervalo esperado (IE) y el valor esperado (VE) del número triangular $\tilde{c}$:

$$IE(\tilde{c}) = [E_1^c, E_2^c] = \left[\int_0^1 f_c^{-1}(x)dx, \int_0^1 g_c^{-1}(x)dx \right] = \left[\frac{1}{2}(c^{pes} + c^{mos}), \frac{1}{2}(c^{mos} + c^{opt}) \right] \quad (2)$$

$$VE(\tilde{c}) = \frac{E_1^c + E_2^c}{2} = \frac{c^{pes} + 2c^{mos} + c^{opt}}{4} \quad (3)$$

De acuerdo con el método de ranking de Jiménez, (1996), para cada pareja de números difusos $\tilde{a}$ y $\tilde{b}$, el grado en el que $\tilde{a}$ es mayor que $\tilde{b}$ puede definirse de la siguiente manera:

$$\mu_M(\tilde{a}, \tilde{b}) = \begin{cases} 0 & si \quad E_2^a - E_1^b < 0 \\ \frac{E_2^a - E_1^b}{E_2^a - E_1^b - (E_1^a - E_2^b)} & si \quad 0 \in [E_1^a - E_2^b, E_2^a - E_1^b] \\ 1 & si \quad E_1^a - E_2^b > 0 \end{cases} \quad (4)$$

Cuando $\mu_M(\tilde{a}, \tilde{b}) \geq \alpha$, se dice que $\tilde{a}$ es mayor o igual que $\tilde{b}$ en al menos un grado α y se representaría como $\tilde{a} \geq_\alpha \tilde{b}$. Ahora, se considera el siguiente modelo matemático difuso en el cual todos los parámetros se definen cómo números difusos triangulares o trapezoidales:

$$\begin{aligned} \min \quad & z = \tilde{c}^t x \\ \text{sujeto a:} \quad & \\ \tilde{a}_i x \geq \tilde{b}_i, \quad & \text{para } i = 1, \dots, I \\ x \geq 0 \end{aligned} \quad (5)$$

De acuerdo con Jiménez et al. (2007), un vector de decisión $x \in \mathcal{R}^n$ es factible en un grado α si $\min_{i=1, \dots, m} \{\mu_M(\tilde{a}_i x, \tilde{b}_i)\} = \alpha$. De acuerdo con la ecuación (3), la ecuación $\tilde{a}_i x \geq \tilde{b}_i$ es equivalente a:

$$\frac{E_2^{a_i x} - E_1^{b_i}}{E_2^{a_i x} - E_1^{b_i} - (E_1^{a_i x} - E_2^{b_i})} \geq \alpha \quad \text{para } i = 1, \dots, I, \quad (6)$$

La ecuación (6) puede reescribirse como:

$$[(1 - \alpha)E_2^{a_i} + \alpha E_1^{a_i}]x \geq \alpha E_2^{b_i} + (1 - \alpha)E_1^{b_i}, \quad \text{para } i = 1, \dots, I \quad (7)$$

Además, Jiménez et al. (2007) muestra que una solución factible como x^0 es una solución óptima α -aceptable del modelo (4) si y solamente si para todos los demás vectores de decisión factibles x tal que $\tilde{a}_i x \geq \alpha \tilde{b}_i$, para $i=1, \dots, I$, $x \geq 0$, por lo cual la siguiente ecuación se satisface:

$$\tilde{c}^t x \geq \frac{1}{2} \tilde{c}^t x^0 \quad (8)$$

Lo que implica que, con una función objetivo de minimización, x^0 es una mejor solución en un grado al menos de $\frac{1}{2}$ que todos los demás vectores factibles. La ecuación (8) puede ser reescrita de la siguiente manera:

$$\frac{E_1 c^t x + E_2 c^t x}{2} \geq \frac{E_1 c^t x^0 + E_2 c^t x^0}{2} \quad (9)$$

Finalmente, utilizando las definiciones de valor e intervalo esperados de un número difuso definidas en (2) y (3), el modelo crisp α -parametrico equivalente al modelo descrito en (6) se puede escribir de la siguiente manera:

$$\begin{aligned} & \min \quad VE(\tilde{c})x \\ & \text{sujeto a:} \\ & [(1 - \alpha)E_2^{a_i} + \alpha E_1^{a_i}]x \geq \alpha E_2^{b_i} + (1 - \alpha)E_1^{b_i}, \quad \text{para } i = 1, \dots, I \\ & x \geq 0 \end{aligned} \quad (10)$$

Nótese que el procedimiento anterior puede extenderse a un modelo multi objetivo y multi variable, además de permitir cualquier tipo de variable con tal de que cumpla la condición de no negatividad. Igualmente, por las propiedades de la desigualdad, si una restricción tiene la dirección contraria a presentada en las anteriores ecuaciones (es decir, el mayor apuntando hacia el número difuso sin la variable), α y $(1 - \alpha)$ intercambian lugares dado el despeje de (6), pues la desigualdad tiene la dirección opuesta. Por ejemplo, el siguiente modelo:

$$\begin{aligned}
& \min \quad z_1 = \tilde{c}^t x \\
& \min \quad z_2 = \tilde{d}^t y \\
& \text{sujeto a:} \\
& \tilde{a}_i x \geq \tilde{b}_i, \text{ para } i = 1, \dots, I \\
& \tilde{e}_j y \leq \tilde{f}_j, \text{ para } j = 1, \dots, J \\
& x \geq 0 \\
& y \in \{0,1\}
\end{aligned} \tag{11}$$

es un modelo al cuál se le puede aplicar el método descrito anteriormente, y el modelo crisp α -paramétrico auxiliar equivalente al modelo (11) sería:

$$\begin{aligned}
& \min \quad VE(\tilde{c})x \\
& \min \quad VE(\tilde{d})y \\
& \text{sujeto a:} \\
& [(1 - \alpha)E_2^{a_i} + \alpha E_1^{a_i}]x \geq \alpha E_2^{b_i} + (1 - \alpha)E_1^{b_i}, \quad \text{para } i = 1, \dots, I \\
& [(\alpha)E_2^{e_j} + (1 - \alpha)E_1^{e_j}]y \leq (1 - \alpha)E_2^{f_j} + (\alpha)E_1^{f_j}, \quad \text{para } j = 1, \dots, J \\
& x \geq 0 \\
& y \in \{0,1\}
\end{aligned} \tag{12}$$

Adicionalmente, se convierte el objetivo de impacto social en una restricción que garantice el cumplimiento de un porcentaje del máximo valor posible del objetivo. De esta manera, se garantiza que se tenga un nivel de impacto social determinado y moderado que permita el equilibrio entre la mejoría de la sociedad y el costo para la empresa (el costo de apertura de un centro de recolección es el más alto de los costos, y dada la naturaleza de la variable de decisión, el máximo de la función social es la apertura de todos los posibles centros de recolección, es decir, una red completamente descentralizada). Dado el alto costo de apertura de un centro de recolección, la función objetivo económica tiende a buscar una red centralizada, mientras que por otro lado la función ambiental tiende a buscar una red descentralizada que incentive la protección del medio ambiente (Pishvaei y Razmi 2012). Dado lo anterior y la presencia de tres objetivos, se decide que un cumplimiento del 30% del máximo de la función objetivo social es un nivel de compromiso aceptable del objetivo.

8.2.1 Modelo matemático auxiliar resultante. Aplicadas las dos transformaciones anteriores al modelo matemático original, se obtiene el siguiente modelo crisp auxiliar:

Objetivo 1: Min(F1)

$$F1 = \sum_t \{ [\sum_k ((g_{kt}^{pesimista} + 2 * g_{kt}^{probable} + g_{kt}^{optimista})/4) * x_{kt} + \sum_j \sum_k (a_{jkt}) * q_{jkt} + \sum_k \sum_l (b_{klt} + (\beta_{lt}^{pesimista} + 2 * \beta_{lt}^{probable} + \beta_{lt}^{optimista})/4) * v_{klt} + \sum_k \sum_m (h_{kmt} + (\tau_{mt}^{pesimista} + 2 * \tau_{mt}^{probable} + \tau_{mt}^{optimista})/4) * w_{kmt} + \sum_k \sum_n (o_{knt} + (\theta_{nt}^{pesimista} + 2 * \theta_{nt}^{probable} + \theta_{nt}^{optimista})/4) * z_{knt}] / (1 + (i^{pesimista} + 2 * i^{probable} + i^{optimista})/4)^t \}$$

Objetivo 2: Min(F2)

$$F2 = \sum_t [\sum_j \sum_k ei_{jk}^{cc} * q_{jkt} + \sum_k \sum_l (ei_l^{st} + ei_{kl}^{cs}) * v_{klt} + \sum_k \sum_m (ei_m^{pl} + ei_{km}^{cp}) * w_{kmt} + \sum_k \sum_n (ei_n^{in} + ei_{kn}^{ci}) * z_{knt}]$$

Objetivo 3: Max(F3)*

$$F3 = \sum_t \sum_k [\{ (wc_{kt}^{pesimista} + 2 * wc_{kt}^{probable} + wc_{kt}^{optimista})/4 \} * x_{kt}]^*$$

Restricciones

$$\sum_k q_{jkt} \geq \left[\alpha * \left(\frac{d_{jt}^{probable} + d_{jt}^{optimista}}{2} \right) + (1 - \alpha) * \left(\frac{d_{jt}^{probable} + d_{jt}^{pesimista}}{2} \right) \right] * \quad (1)$$

$$\left[\alpha * \left(\frac{r_{jt}^{probable} + r_{jt}^{optimista}}{2} \right) + (1 - \alpha) * \left(\frac{r_{jt}^{probable} + r_{jt}^{pesimista}}{2} \right) \right] \text{ para todo } j, .t$$

$$\sum_j q_{jkt} = \sum_m w_{kmt} + \sum_n z_{knt} \text{ para todo } k, t \quad (2)$$

$$\sum_m w_{kmt} = \sum_l v_{klt} \text{ para todo } k, t \quad (3)$$

$$\sum_j q_{jkt} \leq x_{kt} * \left[(1 - \alpha) * \left(\frac{ck_k^{probable} + ck_k^{optimista}}{2} \right) + (\alpha) * \left(\frac{ck_k^{probable} + ck_k^{pesimista}}{2} \right) \right] \quad (4)$$

para todo k, t

$$\sum_k v_{klt} \leq \left[(1 - \alpha) * \left(\frac{cl_l^{probable} + cl_l^{optimista}}{2} \right) + (\alpha) * \left(\frac{cl_l^{probable} + cl_l^{pesimista}}{2} \right) \right] \quad (5)$$

para todo l, t

$$\sum_k w_{kmt} \leq \left[(1 - \alpha) * \left(\frac{cm_m^{probable} + cm_m^{optimista}}{2} \right) + (\alpha) * \left(\frac{cm_m^{probable} + cm_m^{pesimista}}{2} \right) \right] \quad (6)$$

para todo m, t

$$\sum_k z_{knt} \leq \left[(1 - \alpha) * \left(\frac{cn_n^{probable} + cn_n^{optimista}}{2} \right) + (\alpha) * \left(\frac{cn_n^{probable} + cn_n^{pesimista}}{2} \right) \right] \quad (7)$$

para todo n, t

$$0.3 * \text{Max} \left(\sum_k \sum_t x_{kt} \left[\alpha * \left(\frac{wc_{kt}^{probable} + wc_{kt}^{optimista}}{2} \right) + (1 - \alpha) * \left(\frac{wc_{kt}^{probable} + wc_{kt}^{pesimista}}{2} \right) \right] \right) \leq \sum_k \sum_t x_{kt} \left[\alpha * \left(\frac{wc_{kt}^{probable} + wc_{kt}^{optimista}}{2} \right) + (1 - \alpha) * \right. \quad (8)$$

$$\left. \left(\frac{wc_{kt}^{probable} + wc_{kt}^{pesimista}}{2} \right) \right] \text{ para todo } k, t$$

$$w_{kmt}, z_{knt}, q_{jkt}, v_{klt} \geq 0 \text{ para todo } j, k, l, m, n, \quad (9)$$

$$x_{kt} \in \{0, 1\} \text{ para todo } k, t \quad (10)$$

*El objetivo social se cuantifica dados los resultados de los otros objetivos mediante la fórmula presentada.

9. Validación del algoritmo

El algoritmo se implementó en MATLAB, y se corrió en la versión de MATLAB R2017b, utilizando un computador portátil Lenovo Yoga 520 con 8 Gb de RAM DDR4 y un procesador i5 8250u.

Los parámetros del modelo se estiman a partir del material suministrado por Pishvaei en su sitio web además del publicado en su perfil de ResearchGate y se encuentran recopilados en el Apéndice D. Dentro del material suministrado se encuentran los coeficientes utilizados para la solución del caso de estudio planteado en Pishvaei y Razmi (2012), el cual es el mismo problema estudiado en este trabajo, pero con un modelo matemático y mecánica de solución diferentes.

Los parámetros del algoritmo descritos en el capítulo 7, se fijan en los siguientes valores luego de experimentación empírica:

Tabla 7.

Niveles de los parámetros utilizados

Parámetro	Valor usado
α	90%
<i>popSize</i>	75
<i>gen_max</i>	250
<i>max_runs</i>	10
<i>Prob_Mut</i>	30%
<i>Prob_Cruce</i>	90%
<i>%Conver</i>	75%
<i>%Sust</i>	75%
<i>%Legacy</i>	10%

Los valores anteriores se utilizaron durante el proceso de validación del algoritmo. Durante la experimentación, se varía el parámetro de interés, pero se dejan iguales los demás parámetros.

Para la experimentación y posterior validación de las soluciones encontradas por el algoritmo, se usarán las siguientes métricas:

1. Cantidad de centros de recolección abiertos.
2. Cantidad de soluciones no dominadas encontradas.
3. Velocidad computacional.
4. Soluciones no dominadas por minuto.
5. Dominancia entre las soluciones no dominadas encontradas.

La cantidad de centros de recolección abiertos es un valor reportado en la solución del problema realizada por Pishvaei en Pishvaei y Razmi (2012), y es la medida que se usa para la validación del algoritmo. Para la comparación, dado que el modelo del autor es de un solo horizonte de planeación, se toma el máximo de los valores encontrados para los periodos estudiados. Esta métrica se calcula de la siguiente manera:

$$\text{Centros de recolección abiertos} = \text{Máx}(Z), \quad Z = \left(\sum_k [x_{kt}] \right) \text{ para todo } t$$

Las demás métricas se utilizan para comparar el desempeño y calidad de las soluciones al experimentar con algunos de los parámetros del algoritmo. La cantidad de soluciones no dominadas encontradas consiste en la cantidad de soluciones de Pareto diferentes encontradas por el algoritmo en cada corrida experimental, y se calcula como el tamaño del conjunto Q luego de aplicársele un ordenamiento no dominado rápido (es decir, el conjunto de soluciones no dominadas resultante de la corrida). El tiempo de computación se mide respecto a una corrida experimental, sus unidades son segundos y es un valor calculado por MATLAB.

Adicionalmente, se calculan las soluciones no dominadas encontradas por minuto, dividiendo las soluciones no dominadas encontradas por el tiempo computacional convertido de segundos a

minutos, permitiendo medir en cierta forma el equilibrio entre las soluciones encontradas y el tiempo que toma encontrarlas.

La dominancia entre las soluciones no dominadas encontradas consiste en una comparación de las soluciones halladas con diferentes niveles de un parámetro, básicamente es comprobar si un conjunto de soluciones domina al otro; para ello, se extraen y unen los conjuntos no dominados finales de cada réplica experimental de cada nivel del factor y se realiza el ordenamiento no dominado rápido al conjunto resultante. Este proceso resulta en tantos conjuntos como niveles del factor de estudio se tengan, y para la comparación se grafican todos los conjuntos en una misma figura. La dominancia, dada una gráfica de dos funciones de minimización, se puede evidenciar si alguna figura se encuentra ya sea más a la izquierda y/o más hacia abajo que el resto. Esta métrica se utiliza con base al comportamiento encontrado respecto a las soluciones de los diferentes niveles del porcentaje de mutación y del parámetro alfa, en las cuales un nivel del factor encuentra soluciones mejores que otro, evidenciable al graficar las soluciones.

Con base en los resultados de experimentación empírica realizada al probar el algoritmo en MATLAB y garantizar su correcto funcionamiento, los parámetros que se consideraron de interés y se evaluaron fueron:

1. El tamaño de la población.
2. El valor del porcentaje de mutación.
3. El valor del nivel de confianza alfa.

Las hipótesis a priori que se evaluaron experimentalmente para 1 y 2 fueron:

1. El algoritmo se desempeña mejor con niveles bajos de población, teniendo en cuenta el equilibrio de la velocidad de computación y la cantidad de soluciones no dominadas.

2. Para el algoritmo funciona mejor un porcentaje de mutación más alto del usualmente sugerido.

En cuanto al parámetro 3 a evaluar que consiste en el valor del nivel de confianza de alfa, no se realiza un diseño de experimentos, sino que se revisa analíticamente el efecto del factor sobre el modelo matemático y la razones para la decisión de los investigadores de trabajar con nivel de factibilidad del 90%.

Los niveles de los parámetros utilizados en los diseños de experimentos se presentan en la siguiente tabla:

Tabla 8.

Niveles de los parámetros utilizados en la experimentación.

T. Población	50	75	100	125	150
% Mutación	0.1			0.3	

Se realiza un experimento de un factor con 10 réplicas para los dos factores. El análisis de los resultados se realizó en el software MINITAB 18, para ello, se utilizó la función estadística ANOVA de un factor para el análisis de los efectos de los diferentes niveles de cada factor para cada variable de estudio.

En el Apéndice E se encuentran los valores que se utilizaron para la generación de los intervalos de confianza en MINITAB 18.

9.1. Diseño experimental para el tamaño de la población

Respecto al tamaño de la población se tuvieron las siguientes hipótesis a priori:

1. Existe un tamaño de población a partir del cuál no se presentan mejoras significativas para la cantidad de soluciones no dominadas. De acuerdo a experimentación empírica previa, se cree que a partir del tamaño de población de 75 no existen mejoras significativas.
2. La relación entre el tamaño de la población y el tiempo computacional es directa: mayor tamaño de población implica mayor tiempo computacional.
3. El mejor tamaño de población es 75, dado que genera una mayor cantidad de soluciones en menos tiempo si se cumplen las dos hipótesis anteriores.

A continuación se presentan los resultados experimentales para el tamaño de la población:

Tabla 9.

Intervalos de confianza para el factor tamaño de la población y la variable de estudio Cantidad de soluciones no dominadas encontradas

Factor	N	Media	Desv.Est.	IC de 95%
P50	10	342.9	243.7	(168.5, 517.3)
P75	10	1588	639	(1132, 2045)
P100	10	1019	354	(766, 1273)
P125	10	873	429	(566, 1180)
P150	10	1784	731	(1261, 2307)

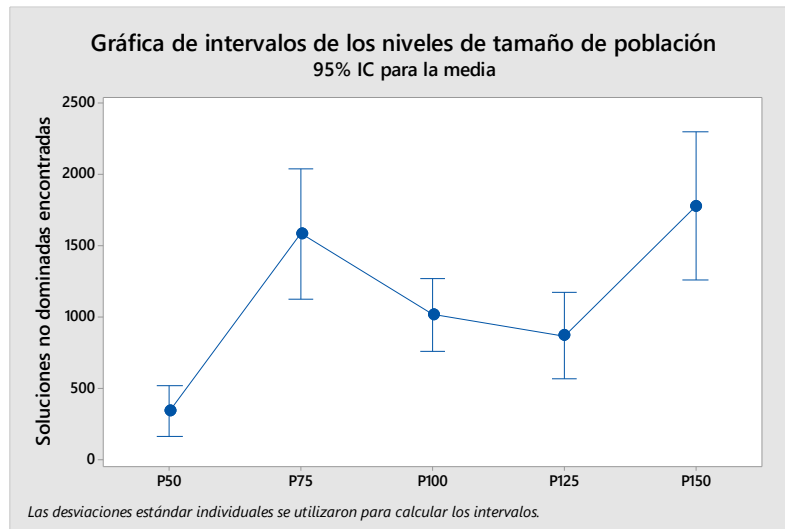


Figura 18. Gráfica de intervalos de los niveles del tamaño de población para la variable respuesta cantidad de soluciones no dominadas encontradas

Tabla 10

Intervalos de confianzas para el factor tamaño de la población y la variable de estudio tiempo computacional

Factor	N	Media	Desv.Est.	IC de 95%
P50	10	2019.0	115.1	(1936.7, 2101.4)
P75	10	2841.2	114.3	(2759.4, 2923.0)
P100	10	4075.8	223.6	(3915.8, 4235.7)
P125	10	5647	537	(5263, 6032)
P150	10	8227	537	(7844, 8611)

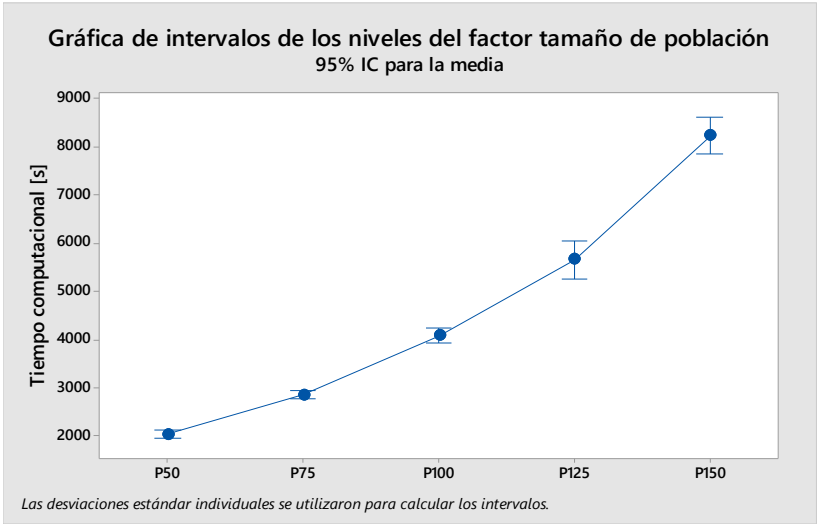


Figura 19. Gráfica de intervalos de los niveles del Tamaño de Población para la variable respuesta
Tiempo Computacional

Tabla 11

*Intervalos de confianza para el factor tamaño de la población y la variable de estudio Cantidad
de soluciones no dominadas encontradas por minuto*

Factor	N	Media	Desv.Est.	IC de 95%
P50	10	10.06	6.71	(5.26, 14.85)
P75	10	33.39	12.85	(24.20, 42.58)
P100	10	15.17	5.74	(11.07, 19.28)
P125	10	9.05	3.83	(6.31, 11.79)
P150	10	13.18	5.96	(8.91, 17.44)

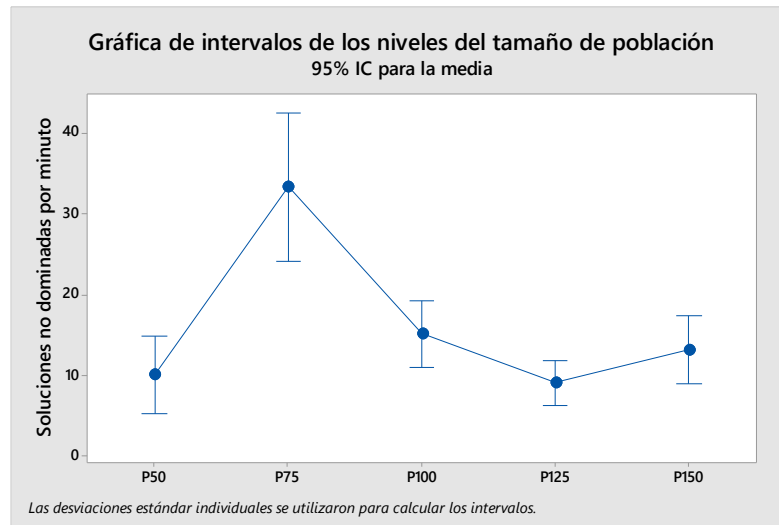


Figura 20. Gráfica de intervalos de los niveles del Tamaño de Población para la variable respuesta Cantidad de soluciones no dominadas encontradas por minuto

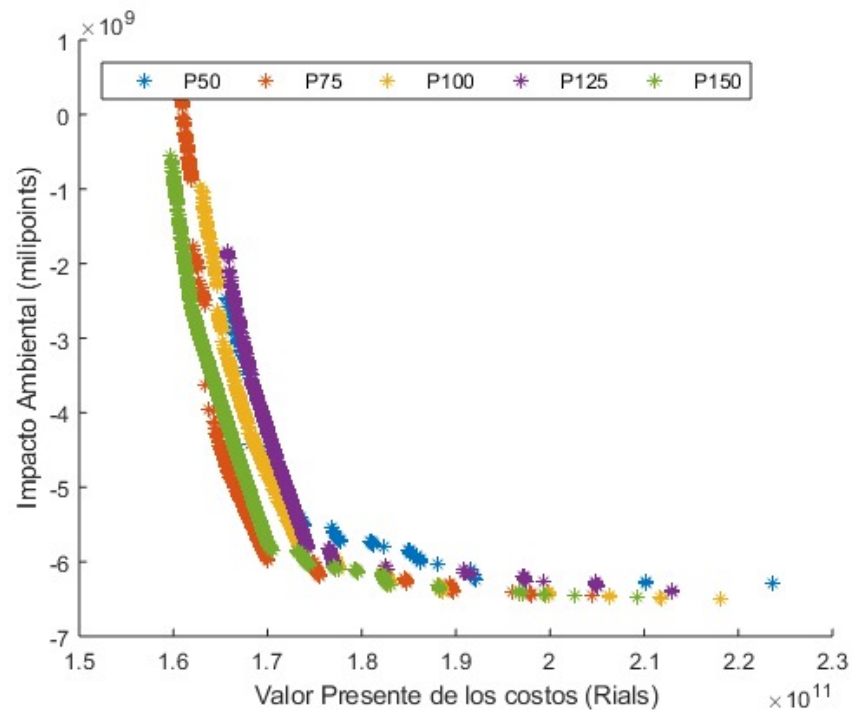


Figura 21. Comparación de los conjuntos de mejores soluciones de todas las réplicas del experimento del factor tamaño de la población.

Se presenta el intervalo de confianza resultante de los niveles del factor tamaño de población para las diferentes variables respuesta en las Tabla 9, Tabla 10 y Tabla 11. El intervalo de confianza se utiliza para comparar los niveles del factor entre sí, y poder determinar el mejor nivel de acuerdo con el valor de la variable respuesta. Se presentan en las Figura 18, Figura 19 y Figura 20 el gráfico de intervalos resultantes de el intervalo de confianza con el fin de visualizar de mejor manera los resultados y poder concluir. En la Figura 21 se presenta la gráfica del conjunto de mejores soluciones encontradas en las corridas experimentales, esta gráfica se utiliza para poder visualizar si existe una dominancia marcada (es decir, que sea estrictamente dominante a todos los demás niveles del factor) entre las mejores soluciones encontradas para un nivel del factor dado.

En el intervalo de confianza resultante presentada en la Tabla 9, se puede evidenciar que la mayor media pertenece al tamaño de población 150 y la sigue muy de cerca la población de 75. Cabe destacar que existe una alta variabilidad en todos los niveles del factor, esto genera intervalos de confianza amplios que permiten concluir que no existe diferencia significativa para las medias de los niveles de tamaño de población 75, 100, 125 y 150, tal como puede verse en la Figura 18 (dado que existe una intersección entre los intervalos de confianza de los niveles mencionados anteriormente, se puede concluir que no existe diferencia significativa entre las medias de cada nivel).

El nivel escogido como el mejor respecto a la cantidad de soluciones no dominadas encontradas es el tamaño de población 75, a pesar de que de aquellas medias que no presentan diferencias significativas no tiene la mayor magnitud tampoco se presente una dominancia marcada entre las mejores soluciones, por tanto, se escoge el nivel 75 pues presenta valores mucho mejores en el tiempo computacional. Dada la igualdad de las medias de los niveles 75, 100, 125 y 150, se puede concluir que no existe evidencia significativa para rechazar la primera hipótesis planteada y que

con un nivel de confianza del 95%, se puede afirmar que existe un tamaño de población (75) a partir del cual no se presenta mejora significativa en la cantidad de soluciones no dominadas encontradas.

El mejor valor en lo que concierne al tiempo computacional es el más bajo, en el intervalo de confianza resultante para esta variable, presentada en la Tabla 10, se puede evidenciar que el menor tiempo computacional pertenece al nivel tamaño de población 50. Los intervalos de confianza muestran que no existe igualdad entre los niveles, dado que no existe intersección entre ninguna pareja de intervalos de confianza. En la gráfica de intervalos, presentada en la Figura 19, se puede ver una relación directa entre el tamaño de la población y el tiempo computacional (a mayor tamaño de población, mayor tiempo computacional).

El mejor nivel respecto al tiempo computacional es el tamaño de población 50, dado que es estrictamente mejor respecto al tiempo computacional al poseer la menor media que es significativamente diferente a las demás. Dado que todas las medias son diferentes entre sí y son mayores para mayores niveles del factor, no existe evidencia significativa para rechazar segunda hipótesis planteada y se puede afirmar con un nivel de confianza del 95% que existe una relación directa entre el tamaño de la población y el tiempo computacional.

Dado que se busca encontrar la mayor cantidad de soluciones no dominadas en el menor tiempo posible, el mejor valor para la variable respuesta es el más alto, en el intervalo de confianza resultante para esta variable, presentada en la Tabla 11, se puede evidenciar que la mejor media la tiene el nivel tamaño de población 75, con un valor superior al doble de la segunda mejor media. Observando los intervalos de confianza, se evidencia que no existe ninguna media que sea significativamente igual a la del tamaño de población 75 (no existe intersección con el intervalo). Este hecho puede verse de manera gráfica en la Figura 20.

El mejor nivel para el factor respecto a la variable de soluciones no dominadas encontradas por minuto es el tamaño de población 75, dado que su media es la más alta y no es significativamente igual a ninguna media de otro nivel del factor. De acuerdo con lo anterior, no existe evidencia significativa para rechazar la tercera hipótesis planteada, y se puede afirmar con un nivel de confianza del 95% que el tamaño de población 75 presenta el mejor resultado de todos los tamaños de población estudiados respecto a la cantidad de soluciones no dominadas encontradas por minuto.

Para el factor tamaño de población, no existe una dominancia marcada entre las soluciones dadas por los diferentes tamaños de población, tal y como puede verse en la Figura 21.

Teniendo en cuenta lo anterior, se puede concluir que *el mejor nivel del factor es el tamaño de población 75*, dado que es el que presenta mejor desempeño teniendo en cuenta todas las variables de estudio.

9.2. Diseño experimental para el porcentaje de mutación

Respecto al porcentaje de mutación, se tuvieron las siguientes hipótesis a priori:

1. Un porcentaje de mutación (30%) más alto al usualmente utilizado (10%) puede generar mejores resultados para el algoritmo propuesto.
2. El porcentaje de mutación no tiene efecto significativo sobre el tiempo computacional.

A continuación se presentan los resultados experimentales para el porcentaje de mutación:

Tabla 12.

Intervalos de confianza para el factor porcentaje de mutación y la variable de estudio Cantidad de soluciones no dominadas encontradas

Factor	N	Media	Desv.Est.	IC de 95%
M10	10	521	410	(228, 814)
M30	10	986	532	(606, 1367)

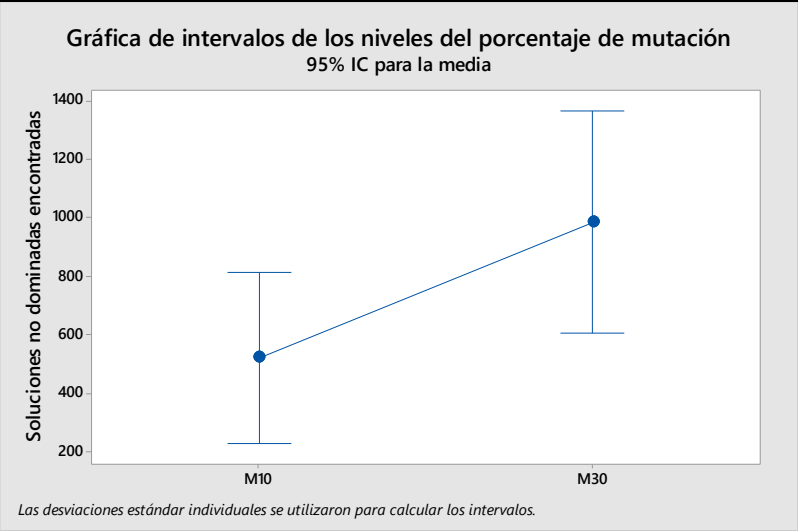


Figura 22. Gráfica de intervalos de los niveles del porcentaje de mutación para la variable respuesta cantidad de soluciones no dominadas encontradas.

Tabla 13.

Intervalos de confianza para el factor porcentaje de mutación y la variable de estudio tiempo computacional

Factor	N	Media	Desv.Est.	IC de 95%
M10	10	2780.2	257.9	(2595.7, 2964.7)
M30	10	2892.2	285.7	(2687.8, 3096.6)

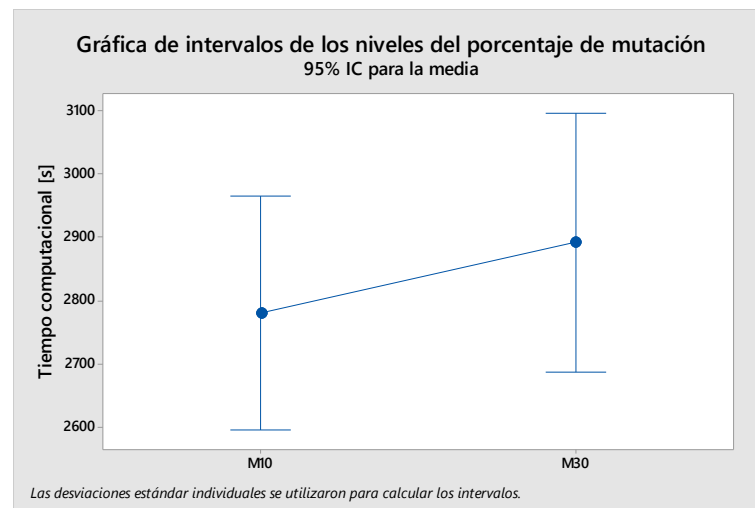


Figura 23. Gráfica de intervalos de los niveles del porcentaje de mutación para la variable respuesta tiempo computacional

Tabla 14.

Intervalos de confianza para el factor porcentaje de mutación y la variable de estudio soluciones no dominadas encontradas por minuto.

Factor	N	Media	Desv.Est.	IC de 95%
M10	10	11.08	8.05	(5.32, 16.84)
M30	10	20.81	12.11	(12.14, 29.47)

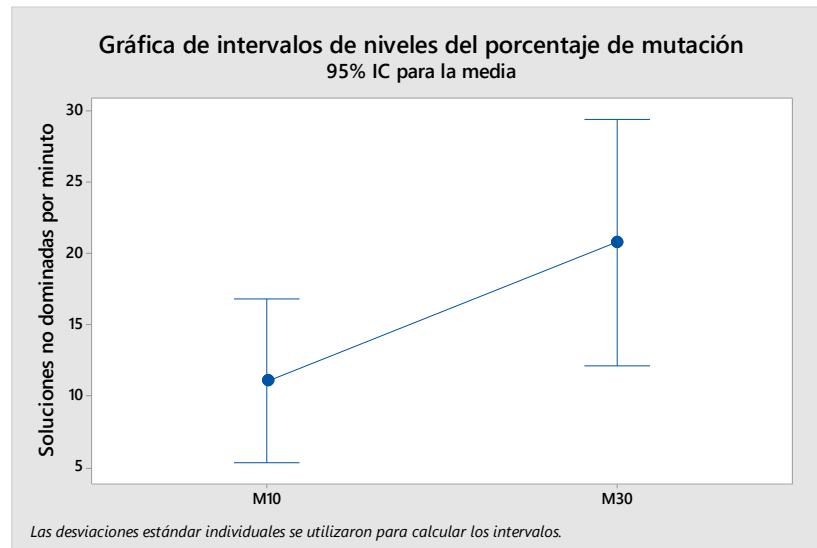


Figura 24. Gráfica de intervalos de los niveles del porcentaje de mutación para la variable respuesta cantidad de soluciones no dominadas encontradas por minuto

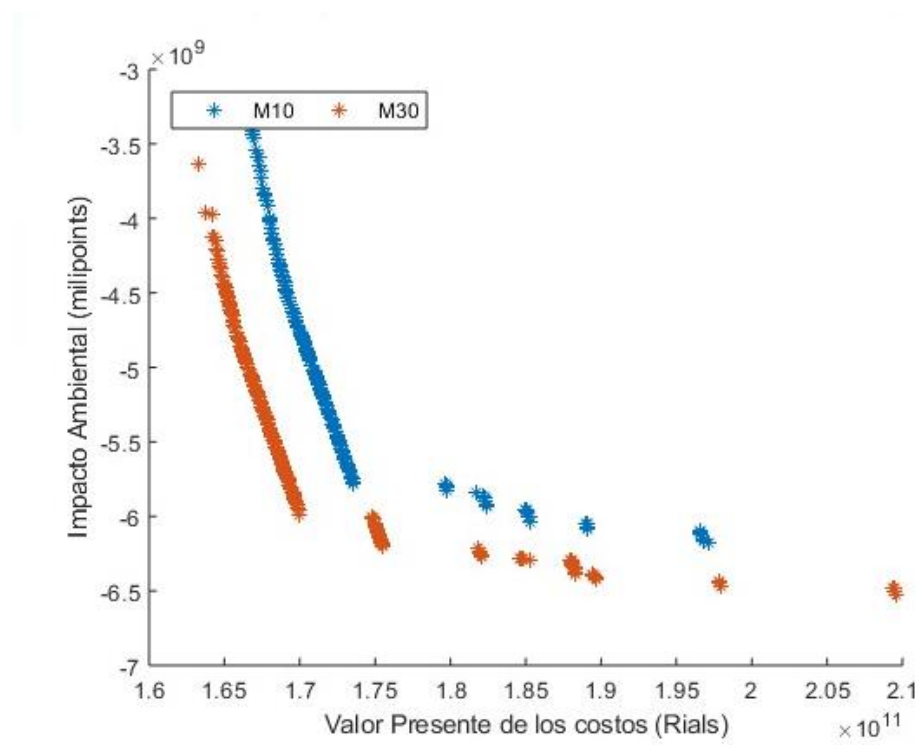


Figura 25. Comparación de los conjuntos de mejores soluciones de todas las réplicas del experimento del factor porcentaje de mutación

Se presenta el intervalo de confianza resultante de los niveles del factor porcentaje de mutación para las diferentes variables respuesta en las Tabla 12, Tabla 13 y Tabla 14. El intervalo de confianza se utiliza para comparar los niveles del factor entre sí, y poder determinar el mejor nivel de acuerdo con el valor de la variable respuesta. Se presentan en las Figura 22, Figura 23 y Figura 24 el gráfico de intervalos resultantes de el intervalo de confianza con el fin de visualizar de mejor manera los resultados y poder concluir. En la Figura 25 se presenta la gráfica del conjunto de mejores soluciones encontradas en las corridas experimentales. Esta gráfica se utiliza para poder visualizar si existe una dominancia marcada (es decir, que sea estrictamente dominante a los todos los demás niveles del factor) entre las mejores soluciones encontradas para un nivel del factor dado.

El mejor valor para las soluciones no dominadas es el mayor, el mejor valor para el tiempo computacional es el menor y el mejor valor para las soluciones no dominadas encontradas por minuto es el menor. En los intervalos de confianza resultantes para los diferentes factores de estudio, presentadas en las Tabla 12, Tabla 13 y Tabla 14, se evidencia que no existe diferencia significativa entre el porcentaje de mutación 30 y el porcentaje de mutación 10, para ninguna de las variables de estudio (cantidad de soluciones no dominadas encontradas, tiempo computacional en segundos y cantidad de soluciones no dominadas encontradas por minuto). No obstante, la media del nivel de mutación 30 es mayor para todas las variables de estudio, tal cómo se puede observar en gráficas Figura 22, Figura 23 y Figura 24.

Dado que no existen diferencias significativas entre las medias de los niveles del factor porcentaje de mutación 10 y 30 para la variable de estudio tiempo computacional, no existe evidencia significativa para rechazar la segunda hipótesis planteada y por tanto se puede afirmar con un nivel de confianza del 95% que los diferentes niveles del factor porcentaje de mutación no

tienen un efecto significativo sobre la variable respuesta tiempo computacional. Dado que las medias para las otras variables de estudio tampoco presentan diferencias significativas, se puede concluir con un nivel de confianza del 95% que el factor porcentaje de mutación no ejerce un efecto significativo en ninguna variable respuesta.

Las mejores soluciones encontradas con el nivel de porcentaje de mutación 30 dominan a las encontradas con el nivel de porcentaje de mutación 10, tal como se evidencia en la Figura 25. Estas mejores soluciones tienen un nivel de confianza del 90% por el proceso de transformación realizado, por lo que no existe evidencia significativa para rechazar la primera hipótesis planteada y se puede afirmar con un nivel de confianza del 90% que las soluciones del porcentaje de mutación 30% son mejores que las del porcentaje de mutación 10%.

Dado lo anterior, se demuestra que una mayor tasa de mutación es benéfica para el algoritmo planteado en el presente documento. Si bien, puede ser contra intuitivo el uso de una alta tasa de mutación para los algoritmos genéticos, se ha encontrado que modelos de mutación más agresivos, ya sea con multi mutación (Hassanat et al., 2005) o tasas de mutación superiores al 10% (Greenwell et al. 1995) tienen el potencial de generar mejoras en los algoritmos genéticos. Se ha propuesto incluso una máxima tasa recomendada del 50% (Senaratna, 2005) o la aplicación de tasas variables decrecientes dependiendo de la generación o el nivel fitness de la población (Matthias et al. 2013). No obstante, en estos estudios se hace la aclaración de que los resultados aplican para los problemas y algoritmos propuestos, y que se debería comprobar si una tasa más alta mejora la calidad de respuestas del algoritmo.

9.3. Análisis del factor α

En la sección 8.2.1, se transforma el modelo en un modelo α -paramétrico. Para el presente trabajo de investigación, se trabajó siempre con un nivel α del 90% por decisión de los investigadores, dado que se consideró que un menor grado de factibilidad no era pertinente para el problema. El grado de factibilidad se puede interpretar como el nivel de confianza que una solución óptima del modelo crisp auxiliar sea también una solución óptima del modelo difuso original. Tomando como ejemplo una de las restricciones del modelo matemático auxiliar, se puede evidenciar el efecto del parámetro en las soluciones:

$$\text{Si } \alpha = 0, \text{ entonces } \sum_k z_{knt} \leq \left[\left(\frac{cn_n^{probable} + cn_n^{optimista}}{2} \right) \right] \quad (1)$$

$$\text{Si } \alpha = 1, \text{ entonces } \sum_k z_{knt} \leq \left[\left(\frac{cn_n^{probable} + cn_n^{pesimista}}{2} \right) \right] \quad (2)$$

para todo n, t

El lado derecho de la desigualdad en (1) es mayor que en (2), es decir, que el valor optimista del factor es mayor que el pesimista. Eso implica que, a mayor nivel de factibilidad, se asume una menor capacidad de procesamiento en los centros de incineración, lo cual lleva a la posibilidad de necesitar usar mayor cantidad de centros de incineración en (2) que en (1). Aplicando el mismo análisis a las demás restricciones permite ver que a un mayor nivel α , la demanda es mayor pero las capacidades son menores, pudiendo generar así la utilización de más recursos para satisfacer la restricción. De manera ilustrativa, se realizó una réplica del algoritmo para $\alpha=50$ y $\alpha=100$, para poder presentar un ejemplo visual del efecto del factor en los resultados:

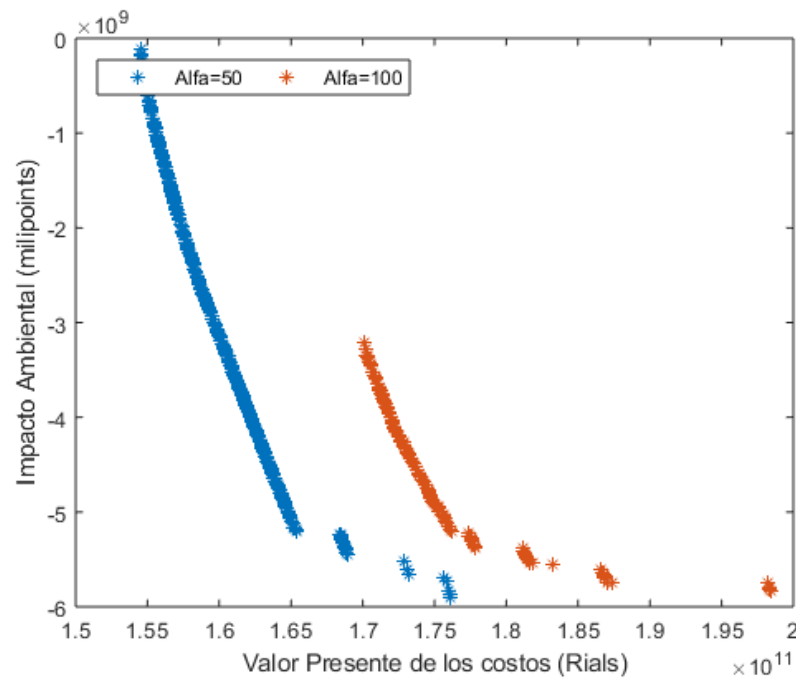


Figura 26. Ilustración del efecto del nivel de factibilidad.

En la anterior figura, se evidencia que las soluciones con el mayor α presentan un mayor uso de recursos y por tanto mayores valores para las funciones objetivo, además de un menor número de soluciones encontradas al hacer menos probable que una solución encontrada aleatoriamente sea factible. Este análisis permite explicar los resultados experimentales de Pishvaei y Razmi (2012), en el cual encuentran que un mayor nivel de factibilidad implica mayores niveles en las funciones objetivo.

Por todo lo anterior, se decidió manejar un nivel de factibilidad del 90%.

9.4.Resultados

Tabla 15.

Resultados experimentales de las mejores soluciones encontradas para las variables objetivo con los diferentes niveles del parámetro Tamaño de la población.

Tamaño de población	Mejor Solución respecto al objetivo	Económico (Miles de millones de Rials*)	Ambiental (Miles de puntos)	Social (Cantidad de trabajos)
50	Económico	166.8953067	-139.636709	1106
	Ambiental	196.0908608	-5175.56087	1106
75	Económico	161.359246	-119.572235	1000
	Ambiental	176.3193546	-6017.92998	1000
100	Económico	163.1914712	-1178.0531	1008
	Ambiental	199.7801608	-6426.54035	1305
125	Económico	174.9262332	137.7550273	1015
	Ambiental	210.7771503	-6051.26085	1015
150	Económico	164.3495973	-1363.95937	1010
	Ambiental	194.323727	-6354.87659	1010

Nota. *El Rial es la moneda de Irán.

En la Tabla 15 se presentan las mejores soluciones encontradas por objetivo en la sección 9.1, utilizando los parámetros descritos en la Tabla 7. A partir de los resultados experimentales, se selecciona la solución que tenga el valor mínimo para el objetivo ambiental y se calculan los valores para los demás objetivos (este mismo procedimiento se realiza respecto al objetivo económico). Se supone por criterio de selección de la mejor solución encontrada la minimización del costo adicional para alcanzar el mejor nivel de protección ambiental, de acuerdo con el nivel de la población. Este costo es la diferencia entre el costo de la mejor solución respecto al objetivo ambiental y el de la mejor solución respecto al objetivo económico, para cada tamaño de población.

La mejor solución de acuerdo con este criterio le alcanza con el nivel de población 75, confirmando ser el mejor nivel del factor.

Tabla 16.

Costo adicional del mejor nivel de protección ambiental encontrado para cada nivel del tamaño de la población

Nivel del factor tamaño de población	Mejor Solución respecto al objetivo	Económico (Miles de millones de Rials)	Costo de la mejor protección ambiental (Miles de millones de Rials)
P50	Económico	166.8953067	29.19555406
	Ambiental	196.0908608	
P75	Económico	161.359246	14.96010861
	Ambiental	176.3193546	
P100	Económico	163.1914712	36.58868963
	Ambiental	199.7801608	
P125	Económico	174.9262332	35.85091707
	Ambiental	210.7771503	
P150	Económico	164.3495973	29.97412962
	Ambiental	194.323727	

9.5. Contraste con instancia de la literatura.

El problema fue resuelto en Pishvae y Razmi (2012) con un modelo matemático diferente, incluyendo en el alcance de su solución los centros de producción y limitándose un solo periodo para el horizonte de planeación. No se planteó el problema con el mismo alcance que los autores dado que si bien fue suministrada información por parte de los autores, en ella no se incluían valores para todos los parámetros que utilizaron. Para contrastar los resultados con los de Pishvae y Razmi, dadas las diferencias en la formulación matemática, se contrasta el rango de centros de recolecciones abiertos en las soluciones no dominadas encontradas para un mismo nivel α (90%),

manejando los mismos valores de los parámetros comunes en los modelos matemáticos. La decisión de abrir el centro de recolección es una decisión con efectos largo plazo, y la más intensiva en costo en la formulación del problema. La comparación se presenta a continuación:

Tabla 17.

Comparación del rango de centros de recolección abiertos para un nivel $\alpha = 90$

Soluciones encontradas en...	Rango de centros de recolección abiertos	Nivel α
Instancia de la literatura*	3 a 7	0.9
Presente trabajo	3 a 6	

Nota: *Pishvaei y Razmi (2012)

10. Conclusiones

El interés por el diseño de redes logísticas inversas verdes que se ha originado en los últimos años viene de la reciente preocupación de las organizaciones por el desarrollo sostenible. En la literatura hay pocas instancias de diseños de redes logísticas inversas sostenibles que tengan consideraciones sociales, ambientales y económica, manejando la incertidumbre a través de parámetros difusos. Un gran número de estas instancias han sido realizadas por un número limitado de autores principalmente ubicados en Irán (Pishvaei) y Dinamarca (Govindan), por lo que este proyecto puede ser considerado como una referencia para futuras investigaciones sobre el problema.

Dado el reciente y limitado interés, no existen instancias conformadas para la solución del problema tratado. Sin embargo, con la información disponible se encuentra un resultado similar a los obtenidos por otros investigadores respecto a la decisión más importante del caso de estudio.

El problema tratado en el presente proyecto de investigación está fuertemente ligado a una decisión importante del sector de salud respecto al reciclaje de ciertos desechos médicos, que al menos en Colombia es obligatoria su incineración (Artículo 5 de la Resolución 482 del 2009). La organización mundial de la salud ha estudiado la viabilidad técnica del proceso de reciclaje de estos desechos, pero para su aplicación es necesario el diseño de una red logística conforme a las normativas del país. La solución propuesta en el presente proyecto ilustra el uso de metodologías, tales como el Eco-Indicador99 utilizado en el modelo matemático o el tratamiento de los parámetros difusos mediante la metodología de Jiménez, que permiten cuantificar la prevención del impacto ambiental del cambio de estas prácticas, además de ayudar en el proceso de evaluación de la viabilidad económica del mismo por parte de los entes interesados.

De acuerdo con los resultados experimentales, se encontró que el algoritmo genético propuesto funciona mejor para tamaños de población pequeños para el problema planteado, siendo el mejor tamaño de población entre los evaluados (50, 75, 100, 125 y 150) el de 75. También se encontró evidencia de que altos porcentajes de mutación pueden tener efectos benéficos para los algoritmos genéticos, dado que el algoritmo utilizado se desempeña mejor con un 30% respecto al 10%.

El algoritmo propuesto no es eficiente respecto al tiempo, no obstante, dado que las decisiones asociadas al problema propuesto son de carácter estratégico (horizonte a largo plazo), son decisiones para las que se justifica la espera. La inexperiencia de los investigadores puede ser una causa de la falta de eficiencia respecto al tiempo computacional, esperable dentro del alcance de un proyecto de pregrado.

11. Recomendaciones

Para futuras investigaciones, existen muchas oportunidades para extender este trabajo en algunos aspectos. Primero, se puede extender el alcance de la red logística para modelar, incluyendo por ejemplo la producción del producto o el uso de diferentes tecnologías de transporte o procesamiento de los productos. Segundo, se pueden utilizar otras metodologías para la estimación de los parámetros sociales y ambientales, por ejemplo, la emisión de gases producto de la incineración o las horas de trabajo perdidas por enfermedades generadas por la contaminación ambiental. Tercero, se pueden considerar otros algoritmos de solución para el caso de estudio, como por ejemplo el SPEA.

Existe la posibilidad de utilizar heurísticas y metaheurísticas para la generación de la población inicial, utilizar una codificación diferente para las soluciones o experimentar con otros mecanismos de generación de diversidad y el uso del elitismo. También es posible adaptar la estructura del algoritmo para otros problemas de estructura similar.

Fomentar en los estudiantes de ingeniería industrial el interés por el estudio de la lógica difusa como herramienta de solución de problemas de investigación de operaciones con datos poco fiables, incentivando así el interés en el desarrollo de este tipo de proyectos.

Referencias Bibliográficas

- Alshamsi, A., & Diabat, A. (10 de 5 de 2017). A Genetic Algorithm for Reverse Logistics network design: A case study from the GCC. *Journal of Cleaner Production*, 151, 652-669.
- Altıparmak, F., Lin, L., & Paksoy, T. (1 de 9 de 2006). A genetic algorithm approach for multi-objective optimization of supply chain networks. *Computers & Industrial Engineering*, 51(1), 196-215.
- Ardeshirilajimi, A., & Azadivar, F. (19 de 4 de 2015). Reverse supply chain plan for remanufacturing commercial returns. *The International Journal of Advanced Manufacturing Technology*, 77(9-12), 1767-1779.
- Arranz, J., & Parra, A. (2007). *Algoritmos Genéticos*. Universidad Carlos III.
- Ávila, S. (2011). *FORMULACIÓN DE UN MODELO DE PROGRAMACIÓN MULTI-OBJETIVO FUZZY PARA LA SELECCIÓN DE PROVEEDORES: CASO DE ESTUDIO*. Universidad del Valle, Santiago de Calí.
- Ayvaz, B., Bolat, B., & Aydın, N. (1 de 11 de 2015). Stochastic reverse logistics network design for waste of electrical and electronic equipment. *Resources, Conservation and Recycling*, 104, 391-404.
- Byron, J., & Iba, W. (s.f.). Population Diversity as a Selection Factor: Improving Fitness by Increasing Diversity.
- Chaabane, A., Ramudhin, A., & Paquet, M. (1 de 1 de 2012). Design of sustainable supply chains under the emission trading scheme. *International Journal of Production Economics*, 135(1), 37-49.

- Chaiyaratana, N. (1997). Recent developments in evolutionary and genetic algorithms: theory and applications. *Second International Conference on Genetic Algorithms in Engineering Systems*. 1997, págs. 270-277. IEEE.
- Charnes, A., & Cooper, W. (1 de 10 de 1957). Management Models and Industrial Applications of Linear Programming. *Management Science*, 4(1), 38-91.
- Chen, P.-y. (1 de 9 de 2012). The investment strategies for a dynamic supply chain under stochastic demands. *International Journal of Production Economics*, 139(1), 80-89.
- Coskun, S., Ozgur, L., Polat, O., & Gungor, A. (1 de 1 de 2016). A model proposal for green supply chain network design based on consumer segmentation. *Journal of Cleaner Production*, 110, 149-157.
- Das, K., & Chowdhury, A. (1 de 1 de 2012). Designing a reverse logistics network for optimal collection, recovery and quality-based product-mix planning. *International Journal of Production Economics*, 135(1), 209-221.
- Das, R., Mallick, S., & Samal, C. (2007). *STUDY OF MULTI-OBJECTIVE OPTIMIZATION AND ITS IMPLEMENTATION USING NSGA-II*. National Institute of Technology.
- Davis, L. (1991). Handbook of genetic algorithms. *Van Nostrand Reinhold, New York*.
- Deb, K., Pratap, A., Agarwal, S., & Meyarivan, T. (2002). *A Fast and Elitist Multiobjective Genetic Algorithm: NSGA-II*.
- Diabat, A., Kannan, D., Kaliyan, M., & Svetinovic, D. (1 de 5 de 2013). An optimization model for product returns using genetic algorithms and artificial immune system. *Resources, Conservation and Recycling*, 74, 156-169.
- Du, F., & Evans, G. (1 de 8 de 2008). A bi-objective reverse logistics network analysis for post-sale service. *Computers & Operations Research*, 35(8), 2617-2634.

- El korch, A., & Millet, D. (1 de 4 de 2011). Designing a sustainable reverse logistics channel: the 18 generic structures framework. *Journal of Cleaner Production*, 19(6-7), 588-597.
- Ergazakis, K., Metaxiotis, K., & Psarras, J. (12 de 10 de 2004). Towards knowledge cities: conceptual analysis and success stories. *Journal of Knowledge Management*, 8(5), 5-15.
- Fahimnia, B., Sarkis, J., & Eshragh, A. (1 de 7 de 2015). A tradeoff model for green supply chain planning: A leanness-versus-greenness analysis. *Omega*, 54, 173-190.
- Feitó-Cespón, M., Sarache, W., Piedra-Jimenez, F., & Cespón-Castro, R. (10 de 5 de 2017). Redesign of a sustainable reverse supply chain under uncertainty: A case study. *Journal of Cleaner Production*, 151, 206-217.
- Fleischmann, M., Bloemhof-Ruwaard, J., Dekker, R., van der Laan, E., van Nunen, J., & Van Wassenhove, L. (16 de 11 de 1997). Quantitative models for reverse logistics: A review. *European Journal of Operational Research*, 103(1), 1-17.
- Ghayebloo, S., Tarokh, M., Venkatadri, U., & Diallo, C. (1 de 7 de 2015). Developing a bi-objective model of the closed-loop supply chain network with green supplier selection and disassembly of products: The impact of parts reliability and product greenness on the recovery network. *Journal of Manufacturing Systems*, 36, 76-86.
- Goedkoop, M., & Spriensma, R. (2001). *The Eco-indicator 99 A damage oriented method for Life Cycle Impact Assessment Methodology Annex*. PRé Consultants B.V.
- Gómez Montoya, R. (2010). Inverse logistics a process with environmental and productivity impacts. *Producción + Limpia*, 5(2), 63-76.
- Govindan, K., Jafarian, A., & Nourbakhsh, V. (1 de 10 de 2015). Bi-objective integrating sustainable order allocation and sustainable supply chain network strategic design with

- stochastic demand using a novel robust hybrid multi-objective metaheuristic. *Computers & Operations Research*, 62, 112-130.
- Govindan, K., Paam, P., & Abtahi, A.-R. (1 de 8 de 2016). A fuzzy multi-objective optimization model for sustainable reverse logistics network design. *Ecological Indicators*, 67, 753-768.
- Greenwell, R., Angus, J., & Finck, M. (1 de 4 de 1995). Optimal mutation probability for genetic algorithms. *Mathematical and Computer Modelling*, 21(8), 1-11.
- Gupta, D., & Ghafir, S. (2012). *International Journal of Emerging Technology and Advanced Engineering An Overview of methods maintaining Diversity in Genetic Algorithms*.
- Hanson, J., & Hitchcock, R. (2009). Towards sustainable design for single-use medical devices. *2009 Annual International Conference of the IEEE Engineering in Medicine and Biology Society*. 2009, págs. 5602-5605. IEEE.
- Hassanat, A., Al-Nawaiseh, N., Abbadi, M., Alkasassbeh, M., & Alhasanat, M. (s.f.). *Enhancing Genetic Algorithms using Multi Mutations: Experimental Results on the Travelling Salesman Problem*.
- Henriques, I., Sadorsky, P., Henriques, I., & Sadorsky, P. (1996). The Determinants of an Environmentally Responsive Firm: An Empirical Approach. *Journal of Environmental Economics and Management*, 30(3), 381-395.
- Ignizio, J. (1976). *Goal programming and extensions*. Lexington Books.
- Indicators, T. E. (s.f.). *Trading Economics: Iran Indicators*. Obtenido de https://tradingeconomics.com/iran/indicators,%20http://www.iberglobal.com/files/2018/iran_iec.pdf,
- Jang, J.-S., Sun, C.-T., & Mizutani, E. (2010). *Neuro-fuzzy and soft computing : a computational approach to learning and machine intelligence*. PHI Learning.

- Jiménez, M., Arenas, M., Bilbao, A., & Rodríguez, M. (2007). Linear programming with fuzzy parameters: An interactive method resolution. *European Journal of Operational Research*, 177(3), 1599-1609.
- Jiménez, M. (8 de 1996). Ranking fuzzy numbers through the Comparison of its Expected Intervals. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 04(04), 379-388.
- Kannan, D., Govindan, K., & Shankar, M. (17 de 2 de 2016). India: Formalize recycling of electronic waste. *Nature*, 530(7590), 281-281.
- Kasabov, N., & Book, A. (1996). *Foundations of Neural Networks, Fuzzy Systems, and Knowledge Engineering*.
- Kim, H., Yang, J., & Lee, K.-D. (1 de 7 de 2009). Vehicle routing in reverse logistics for recycling end-of-life consumer electronic goods in South Korea. *Transportation Research Part D: Transport and Environment*, 14(5), 291-299.
- Kosko, B. (1994). Fuzzy systems as universal approximators. *IEEE Transactions on Computers*, 43(11), 1329-1333.
- Kotzee, I., & Reyers, B. (1 de 1 de 2016). Piloting a social-ecological index for measuring flood resilience: A composite index approach. *Ecological Indicators*, 60, 45-53.
- Koza, J. (1992). *Genetic programming : on the programming of computers by means of natural selection*. MIT Press.
- Kühn, M., Severin, T., & Salzwedel, H. (2013). *Variable Mutation Rate at Genetic Algorithms: Introduction of Chromosome Fitness in Connection with Multi-Chromosome Representation*.

- Lee, C., & Chan, T. (1 de 7 de 2009). Development of RFID-based Reverse Logistics System. *Expert Systems with Applications*, 36(5), 9299-9307.
- Lee, S. (1972). *Goal programming for decision analysis*. Auerbach Publishers.
- Lin, L., Gen, M., & Wang, X. (2007). A Hybrid Genetic Algorithm for Logistics Network Design with Flexible Multistage Model. *International Journal of Information Systems for Logistics and Management*, 3(1), 1-12.
- Medina Hurtado, S. (2006). *Cuadernos de administración*. (Vol. 19). Departamento de Administración, Facultad de Ciencias Económicas y Administrativas, Pontificia Universidad Javeriana.
- Min, H., Jeung Ko, H., & Seong Ko, C. (1 de 1 de 2006). A genetic algorithm approach to developing the multi-echelon reverse logistics network for product returns. *Omega*, 34(1), 56-69.
- Ministry of Housing Spatial Planning and the Environment Communications. (2008). *Eco-indicator 99 Manual for Designers A damage oriented method for Life Cycle Impact Assessment*. Ministry of Housing, Spatial Planning and the Environment Communications.
- Minner, S. (6 de 5 de 2001). Strategic safety stocks in reverse logistics supply chains. *International Journal of Production Economics*, 71(1-3), 417-428.
- Morales-Luna, G. (2002). *Introducción a la lógica difusa*. Centro de Investigación y Estudios Avanzados del IPN.
- Moujahid, A., Inza, I., & Larrañaga, P. (2008). *Algoritmos Genéticos*. Universidad del País Vasco–Euskal Herriko Unibertsitatea.
- Oficina de Información Diplomática. (2019). *Ficha País de la República Islámica de Irán*. Oficina de Información de Irán.

- Organización Mundial de la Salud (OMS). (2005). *Safe Management of Bio-medical Sharps Waste in India A Report on Alternative Treatment and Non-Burn Disposal Practices*.
- Organización Mundial de la Salud. (2003). *Procuring single-use injection equipment and safety boxes*. OMS.
- Osman, I., & Kelly, J. (1996). Meta-Heuristics: An Overview. En I. Osman, & J. Kelly, *Meta-Heuristics* (págs. 1-21). Boston, MA: Springer US.
- Pai, P.-F., Chen, L.-C., & Lin, K.-P. (1 de 4 de 2016). A hybrid data mining model in analyzing corporate social responsibility. *Neural Computing and Applications*, 27(3), 749-760.
- Pishvae, M. (s.f.). *Página web de Pishvae*. Obtenido de <http://www.pishvae.com/>
- Pishvae, M. (s.f.). *Perfil de ResearchGate de Pishvae*. Obtenido de https://www.researchgate.net/profile/Mir_Pishvae/publications
- Pishvae, M., & Razmi, J. (1 de 8 de 2012). Environmental supply chain network design using multi-objective fuzzy mathematical programming. *Applied Mathematical Modelling*, 36(8), 3433-3446.
- Pishvae, M., Farahani, R., & Dullaert, W. (1 de 6 de 2010). A memetic algorithm for bi-objective integrated forward/reverse logistics network design. *Computers & Operations Research*, 37(6), 1100-1112.
- Pishvae, M., Kianfar, K., & Karimi, B. (17 de 3 de 2010). Reverse logistics network design using simulated annealing. *The International Journal of Advanced Manufacturing Technology*, 47(1-4), 269-281.
- Pishvae, M., Razmi, J., & Torabi, S. (1 de 11 de 2012). Robust possibilistic programming for socially responsible supply chain network design: A new approach. *Fuzzy Sets and Systems*, 206, 1-20.

- Prahinski, C., & Kocabasoglu, C. (1 de 12 de 2006). Empirical research opportunities in reverse supply chains. *Omega*, 34(6), 519-532.
- Quiceno, A., & Ramos, M. (2014). *PROPUESTA DE DESARROLLO SOSTENIBLE EN EL MANEJO DE LOS RESIDUOS SÓLIDOS ORGÁNICOS QUE PRODUCE LA PLAZA DE MERCADO, DEL MUNICIPIO DE BUESACO - DPTO DE NARIÑO*. Universidad San Buenaventura.
- Ramos, T., Gomes, M., & Barbosa-Póvoa, A. (1 de 10 de 2014). Planning a sustainable reverse logistics system: Balancing costs with environmental and social concerns. *Omega*, 48, 60-74.
- Rangoaga, M. (2009). *A decision support system for multi-objective programming problems*. University of South Africa.
- Reina, D. (2008). *Fundamentos de Matemática Difusa*. Funcación Universitaria Kornad Lorenz.
- Rigamonti, L., Sterpi, I., & Grosso, M. (2016). Integrated municipal waste management systems: An indicator to assess their environmental and economic sustainability. *Ecological Indicators*, 60, 1-7.
- Rodríguez, N., & Rodríguez, E. (2017). *PROPUESTA DE MEJORA PARA EL AREA DE LOGISTICA INVERSA EN LA PLANTA DE PRODUCCION DE LA INDUSTRIA COLOMBIANA DE LACTEOS (INCOLACTEOS) UBICADA EN SIMIJACA (CUNDINAMARCA)*. Universidad de la Salle.
- Roghanian, E., & Pazhoheshfar, P. (1 de 7 de 2014). An optimization model for reverse logistics network under stochastic environment by using genetic algorithm. *Journal of Manufacturing Systems*, 33(3), 348-356.

- Rusli, K., Abd Rahman, A., & Ho, J. (2012). Green supply chain management in developing countries : a study of factors and practices in Malaysia.
- Sait, S., & Youssef, H. (1999). *Iterative computer algorithms with applications in engineering : solving combinatorial optimization problems*. IEEE Computer Society.
- Sanabria, D., Yulayth, R., & Prada, K. (2015). *UN MODELO MULTINIVEL PARA LA EVALUACIÓN DE UN CORREDOR LOGÍSTICO TIPO GREEN*. Universidad Industrial de Santander.
- Schweickardt, G. (2009). Metaheurística FPSO-X multiobjetivo. Una aplicación para la planificación de la expansión de mediano/largo plazo de un sistema de distribución eléctrica. *Energía*(32), 77-88.
- Senaratna, N. (2005). *Genetic Algorithms: The Crossover-Mutation Debate*.
- Seuring, S., & Müller, M. (1 de 10 de 2008). From a literature review to a conceptual framework for sustainable supply chain management. *Journal of Cleaner Production*, 16(15), 1699-1710.
- Shemshadi, A., Shirazi, H., Toreihi, M., & Tarokh, M. (15 de 9 de 2011). A fuzzy VIKOR method for supplier selection based on entropy measure for objective weighting. *Expert Systems with Applications*, 38(10), 12160-12167.
- Silva, D., Santos Renó, G., Sevegnani, G., Sevegnani, T., & Serra Truzzi, O. (1 de 5 de 2013). Comparison of disposable and returnable packaging: a case study of reverse logistics in Brazil. *Journal of Cleaner Production*, 47, 377-387.
- Simon, H. (1 de 2 de 1955). A Behavioral Model of Rational Choice. *The Quarterly Journal of Economics*, 69(1), 99.

- Soleimani, H., & Kannan, G. (15 de 7 de 2015). A hybrid particle swarm optimization and genetic algorithm for closed-loop supply chain network design in large-scale networks. *Applied Mathematical Modelling*, 39(14), 3990-4012.
- Suyabatmaz, A., Altekin, F., & Şahin, G. (1 de 4 de 2014). Hybrid simulation-analytical modeling approaches for the reverse logistics network design of a third-party logistics provider. *Computers & Industrial Engineering*, 70, 74-89.
- Sverdlin, A., & Ness, A. (1997). *Steel Heat Treatment: Fundamental Concepts in Steel Heat Treatment* (Taylor&Francis ed.). (George E. Totten, Ed.) New York: CRC Press and Taylor&Francis.
- Tamilselvi, S., Baskar, S., & Varshini, P. (2015). *MATLAB code for Constrained NSGA II*.
Obtenido de MATLAB File Exchange:
<https://la.mathworks.com/matlabcentral/fileexchange/49806-matlab-code-for-constrained-nsga-ii-dr-s-baskar-s-tamilselvi-and-p-r-varshini>
- TRADING ECONOMICS. (2019). *Iran - Economic Indicators*. Obtenido de Iran Economic Indicators: <https://tradingeconomics.com/iran/indicators>
- Tseng, M., & Chiu, A. (9 de 10 de 2012). Grey-Entropy Analytical Network Process for Green Innovation Practices. *Procedia - Social and Behavioral Sciences*, 57, 10-21.
- Varela Villegas., R. (2010). *Evaluación económica de proyectos de inversión*. McGraw Hill Educación.
- Veleva, V., Ellenbecker, M., & Ellenbecker, M. (2001). Indicators of sustainable production: Framework and methodology. *Available from: Michael J. Ellenbecker Retrieved on, 9, 519-549.*

- Vélez, M., & Montoya, J. (4 de 10 de 2013). METAHEURÍSTICOS: UNA ALTERNATIVA PARA LA SOLUCIÓN DE PROBLEMAS COMBINATORIOS EN ADMINISTRACIÓN DE OPERACIONES. *Revista EIA*, 4(8), 99-115.
- Vitoriano, B. (2007). *Decisión con Incertidumbre, Decisión Multicriterio y Teoría de Juegos*.
- Vitoriano, B. (2009). *Modelos operativos de Gestión*. Universidad Plutense de Madrid.
- Von Lücken, C., Hermosilla, A., & Barán, B. (2004). Algoritmos evolutivos para optimización multiobjetivo: un estudio comparativo en un ambiente paralelo asíncrono.
- Yang, M., Hong, P., & Modi, S. (1 de 2 de 2011). Impact of lean manufacturing and environmental management on business performance: An empirical study of manufacturing firms. *International Journal of Production Economics*, 129(2), 251-261.
- Yu, P., & Zeleny, M. (1 de 2 de 1975). The set of all nondominated solutions in linear cases and a multicriteria simplex method. *Journal of Mathematical Analysis and Applications*, 49(2), 430-468.
- Zadeh, L. (1 de 1963). Optimality and non-scalar-valued performance criteria. *IEEE Transactions on Automatic Control*, 8(1), 59-60.
- Zadeh, L. (1 de 1 de 1978). Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets and Systems*, 1(1), 3-28.
- Zadeh, L. (1 de 1 de 1999). Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets and Systems*, 100, 9-34.
- Zailani, S., Eltayeb, T., Hsu, C.-C., & Choon Tan, K. (18 de 5 de 2012). The impact of external institutional drivers and internal strategy on environmental performance. *International Journal of Operations & Production Management*, 32(6), 721-745.

- Zhao, W., van der Voet, E., Huppes, G., & Zhang, Y. (16 de 3 de 2009). Comparative life cycle assessments of incineration and non-incineration treatments for medical waste. *The International Journal of Life Cycle Assessment*, 14(2), 114-121.
- Zitzler, E., & Thiele, L. (1998). *An Evolutionary Algorithm for Multiobjective Optimization: The Strength Pareto Approach*.
- Zitzler, E., Deb, K., & Thiele, L. (1991). *Comparison of Multiobjective Evolutionary Algorithms: Empirical Results*. Hajela and Lin.