



UNIVERSIDAD DE O'HIGGINS
DIRECCIÓN DE POSTGRADO
ESCUELA DE INGENIERÍA
MAGÍSTER EN CIENCIAS DE LA INGENIERÍA

**MEDIDAS DE AVERSIÓN AL RIESGO Y APRENDIZAJE CENTRADO EN
DECISIONES**

KATHERINE CONSTANZA GUTHRIE TABILO

**TESIS PARA OPTAR AL GRADO DE MAGÍSTER EN CIENCIAS DE LA INGENIERÍA,
MENCIÓN EN GESTIÓN DE OPERACIONES**

MEMORIA PARA OPTAR AL TÍTULO DE INGENIERA CIVIL INDUSTRIAL

PROFESOR GUÍA:
VÍCTOR BUCAREY LÓPEZ

PROFESOR CO-GUÍA:
EMILIO VILCHES GUTIÉRREZ

MIEMBROS DE LA COMISIÓN:
RENAUD CHICOISNE
GONZALO MUÑOZ

RANCAGUA, CHILE
JUNIO, 2024

Este trabajo ha sido parcialmente financiado por la Subdirección de proyectos de investigación de ANID a través de los proyectos Fondecyt Regular N° 1200283 y Fondecyt de Iniciación N° 11220864.

Resumen

La incertidumbre en parámetros de entrada es omnipresente en problemas de toma de decisiones del mundo real, y su manejo enfrenta desafíos específicos debido a la naturaleza del problema y la disponibilidad de datos. En la presente tesis, se explorarán dos enfoques para abordar la incertidumbre en problemas de optimización. El primero, el aprendizaje centrado en decisiones o Decision-Focused Learning, sigue la premisa "predecir, luego optimizar". El segundo se basa en modelos de aversión al riesgo, una herramienta matemática que evalúa la incertidumbre asociada a eventos o decisiones, expresándola en términos de la probabilidad de resultados no deseados.

El objetivo de esta investigación es determinar qué modelos, contruidos bajo estos dos enfoques, gestionan mejor la incertidumbre asociada con la falta de conocimiento sobre los coeficientes del vector de costos en funciones objetivo, enfocándose en los problemas del camino más corto y la mochila. Para lograrlo, se estableció un caso base para cada problema, y se variaron sus parámetros para observar cómo estos cambios influían en los resultados. Esto permitirá identificar qué modelo es más adecuado y en qué situaciones sería adecuado su uso, facilitando la toma de decisiones en contextos reales con incertidumbre.

Los resultados demostraron que un modelo basado en el paradigma de aprendizaje centrado en decisiones (DFL) fue más efectivo para enfrentar la incertidumbre en el problema del camino más corto, mostrando un desempeño superior en diversas configuraciones con un menor nivel de arrepentimiento. En contraste, los modelos aversos al riesgo presentaron resultados superiores en el problema de la mochila, gestionando mejor la incertidumbre en condiciones variadas y manteniendo un rendimiento consistente. No obstante, estos modelos fueron altamente sensibles a los parámetros de la función objetivo, lo cual afectó significativamente los tiempos de ejecución.

Palabras claves: Optimización bajo incertidumbre, Predict-then-Optimize, Programación estocástica, Medidas de riesgo, Decision-Focused Learning.

*Dedicado a mi Zoe, familia
y amigos más cercanos.*

Agradecimientos

Deseo expresar mi más sincero agradecimiento a todas las personas y seres que han jugado un papel esencial en mi trayectoria, proporcionándome un apoyo inestimable a lo largo de estos años. En particular, quiero destacar a mi gata Zoe, quien ha sido una fuente constante de motivación para mí, acompañándome a lo largo de muchas largas jornadas de estudio y trabajo. Su sola presencia ha sido suficiente para reanimar mi espíritu en los momentos de cansancio. Además, agradezco profundamente a mi madre, quien constantemente me ha recordado la importancia de la educación en la vida.

Asimismo, debo reconocer el respaldo incondicional de mis amigos más cercanos, Daniel y Orlando, quienes han demostrado su confianza y fe en mí de manera inquebrantable. Este logro se debe en gran medida a su constante apoyo, valiosos consejos y afecto sincero.

No puedo dejar de mencionar a mis estimados profesores de la Universidad, quienes compartieron su sabiduría y dedicación conmigo. En especial, agradezco a mi tutor, Víctor Bucarey, por su paciencia infinita, y a mi co-tutor, Emilio Vilches, por sumarse a este proyecto con entusiasmo. A ambos les estoy profundamente agradecida por sus enseñanzas, diálogos enriquecedores y momentos de alegría compartidos.

Además, quiero expresar mi gratitud al profesor Gonzalo Muñoz, quien generosamente me brindó acceso a su cluster para llevar a cabo los experimentos necesarios para esta investigación.

Por último, un reconocimiento especial va hacia mi hermano, Rodrigo, quien ha sido mi compañero de estudios y trabajo durante cinco años, así como a mis compañeros de travesía y, fuente constante de conversaciones provechosas y divertidas. En esta lista se incluyen Consuelo, Enzo, Carolina, Tomás, Lía, Víctor, Diego, Adolfo, Sebastián, Sayli y Benjamín. Su apoyo y camaradería han sido invaluable para mí en este viaje.

Tabla de Contenidos

1. Introducción	1
1.1. Problema de investigación	6
1.2. Revisión bibliográfica	7
1.2.1. Aprendizaje centrado en la decisión	7
1.2.2. Optimización aversa el riesgo	8
1.3. Hipótesis y/o supuestos	11
1.4. Objetivos de investigación	11
2. Marco Teórico	12
2.1. Problemas de optimización	12
2.1.1. Camino más corto	13
2.1.2. Mochila 0 - 1	16
2.2. Medida de aversión al riesgo	17
2.3. Métodos de aprendizaje	19
2.3.1. Regresión Lineal	19
2.3.2. K-medias	19
2.3.3. Análisis de componentes principales	22
2.4. Aprendizaje centrado en decisiones	22
2.5. Criterios de evaluación	25
3. Metodología	29

3.1. Generación de datos	31
3.1.1. Problema del camino más corto	32
3.1.2. Problema de la mochila 0-1	33
3.2. Regresiones lineales y su implementación	34
3.3. Formulación lineal de la pérdida Smart “Predict, then Optimize”	35
3.3.1. Problema del camino más corto	37
3.3.2. Problema de la mochila 0 - 1	38
3.4. Optimización aversa al riesgo	39
3.4.1. Valor en riesgo condicional	40
3.4.2. Pérdida media + varianza	41
3.5. K-medias	44
4. Comparativa Empírica: Optimización Aversa al Riesgo vs. Predecir y Optimizar	46
4.1. Resultados computacionales sobre el problema del camino más corto	48
4.1.1. Medidas de riesgos + contexto	57
4.1.2. Conclusiones para SPP	66
4.2. Resultados computacionales sobre el problema de la mochila	67
4.2.1. Medidas de riesgos + contexto	76
4.2.2. Conclusiones para KP01	83
5. Conclusiones generales	84
5.1. Trabajo futuro	85
Bibliografía	87
A. Resultados estadísticos de SPP	94
B. Resultados estadísticos de KP01	100

Índice de Figuras

1.1. Proceso secuencial de enfoques utilizados	3
2.1. Grafo dirigido vs. Grafo no dirigido	14
2.2. Proceso de extracción de datos	20
3.1. Proceso Metodológico	30
4.1. Distribución del arrepentimiento en función del caso base modificando el tamaño de la red para SPP.	49
4.2. Distribución del arrepentimiento en función del caso base modificando la cantidad de escenarios para SPP.	51
4.3. Distribución del arrepentimiento en función del caso base modificando el grado polinómico para SPP.	53
4.4. Distribución del arrepentimiento en función del caso base modificando en el número de atributos para SPP.	55
4.5. Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información global para SPP	57
4.6. Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información local para SPP	58
4.7. Distribución de arrepentimiento para una red 3 · 3 con información global en SPP.	59
4.8. Distribución de arrepentimiento para una red 3 · 3 con información local en SPP.	60

4.9. Distribución de arrepentimiento para una red $5 \cdot 5$ con información global en SPP.	61
4.10. Distribución de arrepentimiento para una red $5 \cdot 5$ con información local en SPP.	61
4.11. Distribución de arrepentimiento para una red $7 \cdot 7$ con información global en SPP.	62
4.12. Distribución de arrepentimiento para una red $7 \cdot 7$ con información local en SPP.	63
4.13. Distribución del arrepentimiento en función del caso base modificando la capacidad de la mochila para KP01.	68
4.14. Distribución del arrepentimiento en función del caso base modificando la cantidad de escenarios para KP01.	70
4.15. Distribución del arrepentimiento en función del caso base modificando el grado polinómico para KP01.	72
4.16. Distribución del arrepentimiento en función del caso base modificando en el número de atributos para KP01.	73
4.17. Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información global para KP01	76
4.18. Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información local para KP01	77
4.19. Distribución de arrepentimiento para una capacidad 60 con información global en KP01.	78
4.20. Distribución de arrepentimiento para una capacidad 60 con información local en KP01.	78
4.21. Distribución de arrepentimiento para una capacidad 120 con información global en KP01.	79
4.22. Distribución de arrepentimiento para una capacidad 120 con información local en KP01.	79

4.23. Distribución de arrepentimiento para una capacidad 180 con información global en KP01.	80
4.24. Distribución de arrepentimiento para una capacidad 180 con información local en KP01.	80

Índice de Tablas

4.1. Tiempos computacionales [s] de RL RA y SPO+ RA según el tamaño de la red para SPP.	48
4.2. Tiempos computacionales [s] de CVaR ε Y Mean Var α según el tamaño de la red para SPP.	48
4.3. Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de escenarios para SPP.	50
4.4. Tiempos computacionales [s] de CVaR ε Y Mean Var α según la cantidad de escenarios para SPP.	50
4.5. Tiempos computacionales [s] de RL RA y SPO+ RA según el grado polinómico para SPP.	52
4.6. Tiempos computacionales [s] de CVaR ε Y Mean Var α según el grado polinómico para SPP.	52
4.7. Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de atributos para SPP.	54
4.8. Puntuaciones según evaluación de gráficos según parámetro variable para SPP.	56
4.9. Tiempos computacionales [s] de RL RA vs. RL según tamaño de red para SPP.	63
4.10. Tiempos computacionales [s] de SPO+ RA vs. SPO+ según tamaño de red para SPP.	64
4.11. Tiempos computacionales [s] de RL RA y SPO+ RA según la capacidad para KP01.	67

4.12. Tiempos computacionales [s] de CVaR Y Mean Var según la capacidad para KP01.	68
4.13. Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de escenarios para KP01.	69
4.14. Tiempos computacionales [s] de CVaR Y Mean Var según la cantidad de escenarios para KP01.	69
4.15. Tiempos computacionales [s] de RL RA y SPO+ RA según el grado polinómico para KP01.	71
4.16. Tiempos computacionales [s] de CVaR Y Mean Var según el grado polinómico para KP01.	71
4.17. Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de atributos para KP01.	73
4.18. Puntuaciones según evaluación de gráficos según parámetro variable para KP01.	75
4.19. Tiempos computacionales [s] de RL RA vs. RL según capacidad para KP01.	81
4.20. Tiempos computacionales [s] de SPO+ RA vs. SPO+ según capacidad para KP01.	81
A.1. Datos estadísticos de Caso Base (Grid: $5 \cdot 5$, Esc: 1000, k: 10, Deg: 8)	94
A.2. Datos estadísticos con grilla $3 \cdot 3$	95
A.3. Datos estadísticos con grilla $7 \cdot 7$	95
A.4. Datos estadísticos con 200 escenarios	96
A.5. Datos estadísticos con 5000 escenarios	96
A.6. Datos estadísticos con grado polinómico 1	97
A.7. Datos estadísticos con grado polinómico 2	97
A.8. Datos estadísticos con grado polinómico 4	98
A.9. Datos estadísticos con grado polinómico 6	98
A.10. Datos estadísticos con 5 atributos	99
A.11. Datos estadísticos con 20 atributos	99

B.1. Datos estadísticos de Caso Base (Cap: 120, Esc: 750, k: 5, Deg: 8)	100
B.2. Datos estadísticos con capacidad 60	100
B.3. Datos estadísticos con capacidad 180	101
B.4. Datos estadísticos con 200 escenarios	101
B.5. Datos estadísticos con 1500 escenarios	102
B.6. Datos estadísticos con grado polinómico 1	102
B.7. Datos estadísticos con grado polinómico 2	102
B.8. Datos estadísticos con grado polinómico 4	103
B.9. Datos estadísticos con grado polinómico 6	103
B.10. Datos estadísticos con 3 atributos	104
B.11. Datos estadísticos con 7 atributos	104

Capítulo 1

Introducción

El presente trabajo de tesis se enfoca en la comparación de dos metodologías dedicadas a la optimización bajo incertidumbre. Los parámetros desconocidos se presentan de forma lineal en la función objetivo, es decir, el vector de costes de cualquier problema de optimización lineal, convexa o entera. Sin embargo, se cuentan con datos que pueden ayudar a predecir dicho objetivo desconocido. Tres ejemplos clásicos donde surge esta situación son:

- El **problema de tamaño de lote económico**, el cual es un problema de inventario determinista en el que se busca satisfacer una secuencia de demandas para minimizar los costos de orden y mantenimiento. Sin embargo, las demandas deben predecirse a partir de datos tales como la estacionalidad, el historial de búsqueda, las ventas de productos similares, etc. (Levi et al., [2004](#); Wagner & Whitin, [1958](#))
- El **problema del camino más corto**, donde se quiere ir de un origen a un destino, en el menor tiempo posible. Sin embargo, esto requiere predecir el tiempo de viaje de cada arista mediante los límites de velocidad, patrones de tráfico, la hora del día, etc. (Ahuja et al., [1995](#))
- El **problema de optimización de cartera**, es un problema donde se quiere minimizar alguna combinación ponderada de rendimientos de los activos, sin embargo, se necesita predecir el éxito de cada inversión, lo que se podría hacer a partir de datos

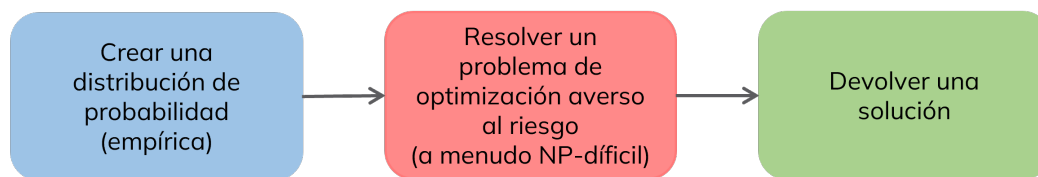
en redes sociales, noticias, indicadores económicos, etc.

Por tanto, todos estos problemas de toma de decisiones puede definirse matemáticamente como $\min_{v \in V} c^\top v$ donde $c \in \mathbb{R}^d$ es el vector de costes del problema de optimización y $v \in \mathbb{R}^d$ es el vector de variables de decisión del conjunto de restricciones que, se espera, minimice el costo, por lo cual se escribe en términos de $c^\top v$, donde c es incierto y v la solución factible.

Un enfoque típico ha sido “Predict-then-Optimize” que separa las predicciones y la optimización en dos etapas, donde primero se entrenan modelos de aprendizaje automático para predecir parámetros relevantes y luego se utilizan algoritmos de optimización especializados para resolver el problema de decisión para estos parámetros predichos. Sin embargo, el entrenamiento del modelo se centra en la precisión de las predicciones de los parámetros que preceden al modelo de decisión, lo que pudiera no conducir a las mejores decisiones (Mandi et al., [2023](#)).

Otra manera común de enfrentar la incertidumbre es mediante la programación estocástica y un enfoque es mediante modelos aversos al riesgo. Una medida de riesgo es una herramienta matemática que se utiliza para evaluar la incertidumbre asociada a un evento o decisión, y generalmente se expresa en términos de la probabilidad de que ocurra un resultado no deseado. Se define como la esperanza de pérdida en exceso de un nivel de umbral, donde la pérdida es medida por una función de utilidad. Existen muchas medidas de riesgo, ejemplo de ellas son *Value-at-Risk* (VaR), *Conditional Value-at-Risk* (CVaR), *Entropic Value-at-Risk* (EVaR), *Upper Semi-Deviation* (USD), etc.

Enfoque de Optimización con Aversión al Riesgo



Enfoque Predict-then-Optimize

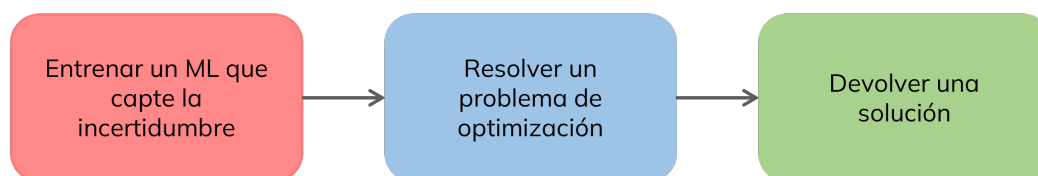


Figura 1.1: Proceso secuencial de enfoques utilizados

Una pregunta muy común es: ¿Cuáles son las diferencias e implicancias de un enfoque predecir y luego optimizar, y uno bajo optimización estocástica?

Los modelos basados en un enfoque de optimización estocástica son inherentemente complejos. Resolver un problema de optimización adverso al riesgo es más complejo que abordar un problema de manera determinista, ya que implica lidiar con múltiples restricciones y variables de decisión enteras. Otra desventaja es que estos modelos no tienen en cuenta el contexto, ya que la función de distribución se construye únicamente a partir de los vectores de costos, lo que pudiera generar un sesgo en la toma de decisiones. (Gowharji & Whitefoot, [2021](#))

Por otro lado, existen los modelos basados en un enfoque de “Predict-then-Optimize” que comienzan con la predicción antes de la optimización. En estos modelos, el énfasis principal recae en el entrenamiento de un modelo que captura la incertidumbre al tomar en cuenta el contexto y predecir los parámetros desconocidos. Esto implica una inversión

significativa de tiempo y esfuerzo. Posteriormente, estas predicciones se insertan en un modelo de optimización, lo que resulta en implementaciones rápidas y, finalmente, en la obtención de una solución.

En ambos enfoques, se destaca en rojo (Ver Figura 1.1) la etapa donde se concentra el mayor esfuerzo necesario para obtener una solución. En la optimización adversa al riesgo, se enfrenta el desafío de resolver un problema complejo con todos los datos disponibles y sin considerar el contexto. En contraste, en el enfoque “predecir y luego optimiza”, se invierte mucho esfuerzo en el entrenamiento para realizar estimaciones precisas para minimizar un error de predicción.

Para abordar estas limitaciones y equiparar ambos enfoques en términos de consideración del contexto, se propone contextualizar las medidas de riesgo mediante el uso de la técnica *KMeans*. A partir del vector correlacionado de características x , se pueden formar grupos o *clusters*, cada uno con su propia función de distribución. Cuando se recibe una nueva observación, esta se clasifica en uno de los *clusters*, y posteriormente se resuelve el problema de optimización, pero utilizando solo los datos del *cluster* al que pertenece. Esta estrategia permite contextualizar el modelo de aversión al riesgo.

Recientemente, el enfoque de predecir y luego optimizar de extremo a extremo se ha convertido en una alternativa atractiva. Este último también conocido como aprendizaje centrado en decisiones (DFL). Este es un paradigma emergente en el aprendizaje automático que entrena directamente modelos para optimizar decisiones, integrando la predicción y la optimización en sistemas de un extremo a otro. A diferencia del enfoque habitual de “dos pasos”, DFL busca mejorar la calidad de las decisiones entrenando modelos de aprendizaje automático especializados en realizar predicciones para optimizar decisiones.

Con base en lo anterior, es que se buscará descubrir que enfoque modelo enfrenta mejor

la incertidumbre dada la naturaleza del problema y la disponibilidad de datos.

De forma paralela, se buscarán respuestas a preguntas subyacentes, tales como:

- ¿Cuál es el efecto de aumentar o disminuir el conjunto de datos, la dimensión del vector de características y/o el tamaño de la red o capacidad (este último dependiendo del problema que se esté tratando)?
- ¿Cuál es el impacto de aumentar o disminuir el ruido en la generación de los datos?
- ¿Cuál es el resultado de contextualizar las medidas de riesgo?
- ¿Cuánto se gana al considerar toda la información contextual disponible?

1.1. Problema de investigación

En esta tesis se estudiarán dos metodologías para tratar con la incertidumbre en parámetros de entrada en problemas optimización. El primer enfoque se basa en el paradigma del aprendizaje centrado en decisiones, donde el modelo de aprendizaje automático se entrena para optimizar los criterios de evaluación que miden la calidad de las decisiones obtenidas. El segundo enfoque se basa en un modelo de aversión al riesgo, que es una herramienta matemática utilizada para evaluar la incertidumbre asociada a un evento o decisión, y se expresa en términos de la probabilidad de que ocurra un resultado no deseado. Por tanto, la pregunta de investigación es ¿Qué enfoque será más efectivo para abordar la incertidumbre en problemas complejos de optimización?

1.2. Revisión bibliográfica

1.2.1. Aprendizaje centrado en la decisión

Los problemas de optimización que forman parte del aprendizaje centrado en decisiones tienen parámetros de optimización parcialmente definidos. Uno de los avances clave en esta área es la capa de optimización diferenciable (Amos & Kolter, 2017), que calcula los gradientes de un programa cuadrático de forma analítica al diferenciar los requisitos de optimización de KKT. Sin embargo, este cálculo de gradientes no es aplicable a los problemas de programación lineal (PL) debido a las singularidades de la función objetivo. Para superar esto, Wilder et al. (2019) añadieron un regularizador cuadrático a la función objetivo, mientras que Mandi y Guns (2020) introdujeron un término de regularización de barrera logarítmica y calcularon los gradientes a partir de la incrustación autodual homogénea del PL. Siguiendo la arquitectura de Wilder et al. (2019), Ferber et al. (2020) estudiaron los PL enteros mixtos y los redujeron a PL agregando un plano de corte al nodo raíz del PL. Todas estas técnicas son exclusivas de la estructura empleada del problema de optimización restringida.

Se pueden emplear otros métodos independientemente del problema de optimización restringida y el solucionador que se aplique. En esta línea, Elmachoub y Grigas (2022) propusieron un enfoque inteligente de predicción y optimización que estudia los problemas de optimización con objetivos lineales y cuándo es necesario predecir el vector de costos. Además, propusieron un subgradiente accesible e introdujeron un límite superior sustituto convexo de la pérdida. De manera similar, Pogančić et al. (2020) también analizaron problemas con objetivos lineales y ajustaron el vector de costo anticipado para calcular el gradiente. Este trabajo fue ampliado por Niepert et al. (2021) mediante la agregación de perturbaciones de ruido para la diferenciación implícita basada en perturbaciones. Al respecto, resulta particularmente recomendable el texto de Kotary et al. (2021) para un resumen del aprendizaje centrado en decisiones.

En otra línea de investigación mostrada por van Staden et al. (2022), los autores implementan un enfoque híbrido novedoso para evaluar y probar información sintetizada y del mundo real a través de una combinación de inteligencia artificial estadística y simbólica.

Sin embargo, todos esos métodos requieren resolver el problema de optimización para cada instancia durante el entrenamiento, lo que puede ser costoso computacionalmente. Para abordar este problema, se estudiará la conveniencia de resolver bajo un enfoque de aprendizaje centrado en la decisión o bajo un modelo averso al riesgo.

1.2.2. Optimización aversa al riesgo

Ser neutral al riesgo, o minimizar el valor esperado, es un objetivo típico cuando se trata con incertidumbre en un problema de optimización. Sin embargo, las realizaciones extremas son más pertinentes en algunas circunstancias. Por ejemplo, cuando se elige la ruta para un servicio de emergencia, es preferible elegir una ruta que siempre llegue en un tiempo razonable que una que tenga un tiempo de viaje proyectado bajo, pero que ocasionalmente pueda tardar una cantidad de tiempo inaceptable.

La primera definición matemática del riesgo se realizó en los primeros trabajos de Von Neumann y Morgenstern (1947), donde se demostró que, bajo los supuestos de racionalidad, el riesgo que percibe un agente racional está directamente relacionado con la concavidad de una función de desutilidad convexa $u(\cdot)$ del resultado aleatorio X y puede modelizarse como el valor esperado $\mathbb{E}_x[u(X)]$. Este método de definir el riesgo también se conoce como equivalente de certeza bajo una función de utilidad (Yaari, 1969, 1987). Formalmente, una medida de riesgo es un mapeo de valor real de una colección de variables aleatorias utilizada para expresar el nivel de riesgo que conlleva una variable aleatoria. El valor condicional en riesgo (CVaR) (Acerbi & Tasche, 2002; Rockafellar, Uryasev et al., 2000) y las medidas de riesgo entrópico (EVaR) (Cominetti, 2015; Föllmer & Schied, 2002)

son métricas de riesgo significativas en la literatura. El modelo de riesgo medio se creó a raíz de los primeros trabajos de Markowitz (1952), que utilizaba la varianza para reflejar la dispersión de los datos y ofrecía una técnica de solución para los problemas aversos al riesgo. Minimizar un objetivo CVaR es un problema de optimización convexo, como demostraron Rockafellar, Uryasev et al. (2000). Sin embargo, en el caso de varias variables aleatorias correlacionadas, esto requiere la minimización de una expectativa, que puede ser difícil de calcular.

Con distribuciones no correlacionadas y normales en las duraciones de los viajes de borde, Nikolova et al. (2006) y Lim et al. (2013) utilizaron este último modelo para abordar problemas de camino más corto de reducción de CVaR y VaR utilizando un método de aproximación personalizado de tiempo completamente polinomial. Ambos trabajos se basan en gran medida en la ausencia de correlaciones y en el hecho de que el CVaR y el VaR se reducen a formas cuadráticas cuando la distribución subyacente del tiempo de viaje es normal multivariante. Los autores Averbakh y Lebedev (2004) también tienen en cuenta los tiempos de viaje inciertos, pero presentan metodologías para resolver una versión robusta del problema del camino más corto que solo tiene en cuenta el soporte de los tiempos de viaje aleatorios. Kwon (2011) proporciona una solución potente para resolver el problema del camino más corto con el objetivo CVaR en redes a gran escala. Sin embargo, en lugar de centrarse en los tiempos de viaje, este enfoque busca minimizar el CVaR de la probabilidad de que ocurra un accidente catastrófico, lo que crea problemas con estructuras claras. En Zhang y Homem-de-Mello (2017), los autores se centran en resolver problemas de caminos más cortos neutrales al riesgo con restricciones de dominancia estocástica en los que los tiempos de viaje no están correlacionados. Para ello, combinan técnicas heurísticas y enfoques de ramificación y corte.

Song y Shen (2016) se concentraron en un problema de interdicción del camino más corto algo distinto. En este problema, para que los caminos más cortos tengan una alta

probabilidad de superar al menos un valor umbral especificado, es necesario bloquear algunas aristas de la red en este problema. Hasta donde se sabe, toda la bibliografía sobre el cálculo de los caminos más cortos con aversión al riesgo parte del supuesto de que los tiempos (o costes) de viaje no están correlacionados y son desconocidos. La correlación en los tiempos de viaje desconocidos está presente en las redes de transporte, como demuestran Yang et al. (2013). Además, descartar cualquier vínculo potencial en los tiempos de viaje aleatorios podría, en general, no producir los mejores resultados (Agrawal et al., 2010, 2012).

Finalmente, se revisó el trabajo de Chicoisne et al. (2018) donde desarrollaron, considerando la medida de EVaR y CVaR en presencia de correlaciones entre variables aleatorias, algoritmos para resolver el problema del camino más corto aversos al riesgo mediante aproximación media muestral, técnicas de linealización y métodos de descomposición.

1.3. Hipótesis y/o supuestos

Se espera que una formulación bajo un enfoque de aprendizaje centrado en decisiones sea más efectiva para abordar la incertidumbre en problemas de optimización en comparación con una formulación que incluye una medida de riesgo.

Esta hipótesis se basa en el hecho de que el enfoque de DFL considera que los vectores predichos van a ser utilizados en el problema de optimización posterior y también considera el vector de características para predecir los vectores de costos. Esto último permite identificar patrones y relaciones entre las características y los costos, lo que puede ser útil para mejorar la eficiencia y la precisión de las predicciones. Por otro lado, la formulación que incluye una medida de riesgo no tiene en cuenta el contexto de los atributos y, por lo tanto, podría tener una menor capacidad para abordar la incertidumbre de manera efectiva en problemas de optimización.

1.4. Objetivos de investigación

El objetivo general de este proyecto de tesis es comparar dos enfoques para la toma de decisiones bajo incertidumbre: Optimización Aversa al Riesgo y Aprendizaje centrado en decisiones, con el propósito de minimizar el impacto de la incertidumbre en un problema de decisión.

Los objetivos específicos de esta investigación son los siguientes:

1. Formular modelos utilizando ambos enfoques para diferentes problemas de decisión.
2. Adaptar los modelos de optimización con medidas de riesgo para considerar el contexto y lograr una evaluación más justa.
3. Evaluar y comparar el impacto de la calidad de las soluciones generadas por ambas metodologías.

Capítulo 2

Marco Teórico

En esta sección se describirán detalladamente los problemas de optimización a tratar, así como las técnicas y medidas de evaluación que se emplearán en la investigación.

2.1. Problemas de optimización

Un problema de programación lineal, sin pérdida de generalidad, se puede escribir de la siguiente forma:

$$\begin{aligned} \min_{v} \quad & c^T v \\ \text{s.a.} \quad & Av \leq b \\ & v \geq 0, \end{aligned} \tag{2.1}$$

donde las variables de decisión son $v \in \mathbb{R}_+^d$, los coeficientes de coste son $c \in \mathbb{R}^d$, los coeficientes de las restricciones son $A \in \mathbb{R}^{k \times d}$, y el lado derecho de las restricciones es $b \in \mathbb{R}^k$. En un programa entero (mixto) algunas de las variables v_i están restringidas a tomar valores enteros:

$$\begin{aligned}
& \underset{v}{\text{mín}} && c^T v \\
& \text{s.a.} && Av \leq b \\
& && v \geq 0 \quad v_i \in \mathbb{Z} \quad \forall i \in D', \quad D' \subseteq \{1, 2, \dots, d\}.
\end{aligned} \tag{2.2}$$

Sea V la región factible en la programación lineal y entera, $z^*(c)$ el valor objetivo óptimo con respecto al vector de coste c , y $v^*(c) \in V^*(c)$ una solución óptima particular que se deriva de algún solucionador. Se define el conjunto de soluciones óptimas $V^*(c)$ debido a que puede haber múltiples óptimos.

En el presente trabajo se utiliza como problema de optimización el camino más corto y el problema de la mochila, que se describen a continuación.

2.1.1. Camino más corto

El problema del camino más corto o Shortest Path (SPP, por sus siglas en inglés) es un problema de optimización en teoría de grafos. Un grafo está formado por nodos (vértices) y aristas (conexiones entre nodos) (Madkour et al., 2017). Un grafo dirigido G es un par (N, A) , donde N es un conjunto de vértices, y A es un conjunto de aristas - pares ordenados de estos vértices $A \subseteq \{(s, t) | s, t \in N\}$. Un grafo dirigido ponderado por aristas es un grafo dirigido $G(N, A)$ emparejado con una función de peso w . Las funciones de peso asignan un peso (valor numérico) a cada arista $(s, t) \in A$. Un camino $s - t$ en un grafo G es una secuencia finita alternada de vértices y aristas de G que comienza con el vértice s , termina con el vértice t y donde cada arista de la secuencia conecta el vértice que le precede en la secuencia con el vértice que le sigue en la secuencia ¹.

¹Un grafo dirigido o digrafo se refiere a un tipo de grafo donde las aristas tienen una dirección asociada, es decir, cada arista va de un vértice a otro en una dirección específica. En un grafo no dirigido, todas las aristas son implícitamente bidireccionales.

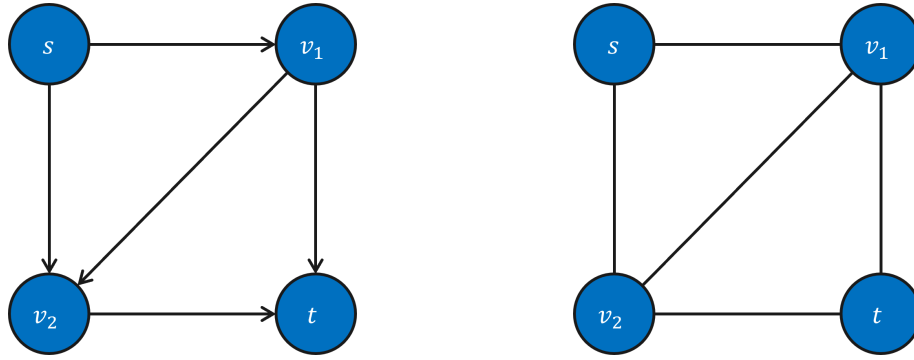


Figura 2.1: Grafo dirigido vs. Grafo no dirigido

Un camino $s - t$ en un dígrafo G es un camino $s - t$ con todas las aristas dirigidas en la dirección de la secuencia. Si G es un grafo de aristas ponderadas, el coste del camino c puede definirse como $c(s - t) = \sum w(a); a \in A \wedge a \in \text{caminos} - t$. El camino $s - t$, p , se denomina el camino más corto $s - t$ si y sólo si $c(p) \leq c(c - t)$.

Matemáticamente, el problema del camino más corto puede plantearse de la siguiente manera:

$$\text{mín} \sum_{(i,j) \in A} c_{ij} v_{ij} \quad (2.3)$$

$$\text{s.a.} \quad \sum_{i:(i,j) \in A} v_{ij} - \sum_{j:(i,j) \in A} v_{ji} = \begin{cases} 1 & \text{si } i = s \\ 0 & \text{si } i \neq s, t \\ -1 & \text{si } i = t \end{cases} \quad i \in N \quad (2.4)$$

$$v_{ij} \geq 0, \quad \forall (i,j) \in A \quad (2.5)$$

El “costo” o “distancia” de una arista puede representar diferentes cosas dependiendo del contexto: podría ser la distancia física entre dos ubicaciones, el tiempo necesario para viajar entre dos ciudades, el costo monetario de una transacción, la cantidad de recursos

necesarios para completar una tarea, entre otros.

El objetivo principal del problema del camino más corto es determinar la ruta con el costo total más bajo posible desde un nodo de origen dado hasta un nodo de destino específico. Existen varios algoritmos para resolver este problema, siendo los más conocidos el algoritmo de Dijkstra y el algoritmo de Bellman-Ford, siendo este último el cual se utilizará. Este problema tiene aplicaciones en el diseño de rutas de navegación, planificación logística, optimización de redes y comunicaciones, entre otros campos (Gao et al., 2014; Yu, 1998). Este problema puede resolverse más eficientemente para grafos sin aristas de ponderación negativa que para grafos con aristas de ponderación negativa.

Grafos sin aristas con ponderación negativa El algoritmo Bellman-Ford, publicado por varias personas, incluidos Richard Bellman (Bellman, 1958) y Lester Ford, entre 1955 y 1959, fue la primera solución a este problema. El algoritmo de Bellman-Ford funciona en $O(|N| \cdot |A|)$ y se puede usar en grafos con aristas con pesos negativos. El algoritmo de Dijkstra (Dijkstra, 1959) es una solución alternativa que todavía se utiliza. El algoritmo de Dijkstra tenía una complejidad temporal inicial de $O(|N|^2)$, pero ahora, con algunas mejoras, se ejecuta en $O(|A| + |N| \log |N|)$ mediante métodos avanzados.

Grafos con aristas de ponderación negativa Puede haber un ciclo negativo en grafos con aristas de ponderación negativa. En un grafo, si hay un ciclo negativo y el ciclo es alcanzable desde el origen s , el costo del camino más corto $s - t$ para cada vértice t , $t \in N$ alcanzable desde el ciclo es $-\infty$. El algoritmo de Bellman-Ford fue la primera solución a este problema. Dado que las predicciones de los costos en los arcos pudiesen llegar a tener valor negativo, este algoritmo es el que se considera en este trabajo $O(|N| \cdot |A|)$.

2.1.1.1. Algoritmo de Bellman-Ford

La programación dinámica es la base del algoritmo de Bellman-Ford. Este algoritmo se encarga de determinar el camino más corto a cada vértice en el grafo, comenzando por

considerar solo una arista, luego dos aristas, y así sucesivamente. Cada etapa del algoritmo utiliza los resultados de las $a - 1$ aristas previas para calcular el resultado para a aristas. Estos pasos se conocen como iteraciones y son fundamentales en el proceso.

En el caso de que el grafo no contenga ciclos negativos, el camino más corto usando como máximo $|N|$ aristas es equivalente al camino más corto sin ninguna restricción. Por lo tanto, para obtener el resultado deseado, es necesario relajar el grafo $|N| - 1$ veces, asegurando así que se encuentre el camino más corto desde el vértice fuente hasta cualquier otro vértice en el grafo.

2.1.2. Mochila 0 - 1

El problema de la mochila 0-1 o Knapsack 0-1 (KP01, por sus siglas en inglés) ha sido objeto de extensos estudios desde los años 60s (Horowitz & Sahni, 1974; Ingargiola & Korsh, 1973; Kolesar, 1967; Nemhauser & Ullmann, 1969). Se trata de un problema especial de programación entera con una restricción. En este contexto, se nos proporcionan dos conjuntos de I enteros positivos: $U = \{u_1, u_2, \dots, u_i\}$ y $W = \{w_1, w_2, \dots, w_i\}$, junto con M también positivo.

Los valores w_i representan los tamaños, pesos o volúmenes de los objetos numerados como $1, 2, \dots, i$, mientras que M se refiere al tamaño o capacidad de la mochila en cuestión. Los elementos u_i en el conjunto U representan los beneficios asociados a cada objeto, lo que significa que al incluir el objeto i en la mochila, se obtiene un beneficio de u_i . Cabe destacar que cada objeto i ocupa un espacio de tamaño w_i cuando se coloca en la mochila. Matemáticamente, el problema de la mochila 0-1 puede plantearse de la siguiente manera:

$$\max \sum_{i=1}^I u_i v_i \quad (2.6)$$

$$\text{s.a.} \sum_{i=1}^I w_i v_i \leq M \quad (2.7)$$

$$v_i \in 0, 1 \quad (2.8)$$

El Problema de la mochila es un ejemplo común y ampliamente estudiado de un problema de optimización combinatoria NP-difíciles (Cao et al., 2018; Merkle & Hellman, 1978), en el que encontramos la solución óptima del problema dado de forma que satisfaga la restricción dada. Los KP01 están presentes en una amplia gama de procesos de toma de decisiones en el mundo real (Feng et al., 2017). Algunos ejemplos de aplicaciones de los KP01 en el mundo real incluyen el problema de asignación de presupuesto de capital (Bas, 2011), el problema de mantenimiento de propiedades inmobiliarias (Taillandier et al., 2017), el problema de carga de mercancías (Pisinger, 2002), el problema de asignación de recursos (Reniers & Sörensen, 2013), el problema de selección de proyectos (Chang & Lee, 2012), las subastas combinatorias (Pfeiffer & Rothlauf, 2007), el problema de disponibilidad de promesas (LAU & LIM, 2004), el problema de corte de existencias (Muter & Sezer, 2018) y la toma de decisiones de inversión (Peeta et al., 2010).

Por tanto, el problema KP01 se enfoca en un conjunto de elementos que tienen asignados un peso y un valor. El objetivo principal es determinar la cantidad de cada artículo para incluirlo en una colección, de tal manera que el peso total de la colección no exceda un límite máximo preestablecido, y al mismo tiempo, se maximice el valor total obtenido de los artículos seleccionados (Feng et al., 2018; Li et al., 2012; H. Wu et al., 2018).

2.2. Medida de aversión al riesgo

Dentro de las metodologías que se utilizarán en la presente investigación se tomarán decisiones en base a la aversión al riesgo y para cuantificar dicha aversión

Una medida de riesgo es una función ρ que asigna un valor real a una variable aleatoria X para cuantificar su grado de riesgo. Para definir el concepto de medida de riesgo, se debe considerar un espacio de probabilidad (Ω, A, P) tal que Ω es un conjunto de todos los eventos simples, A es un σ -álgebra de subconjuntos de Ω , y P es una medida de probabilidad en A . Además, sea L_0 el conjunto de todas las funciones medibles de Borel (variables aleatorias) $X : \Omega \rightarrow \mathbb{R}$, y sea $\mathbb{X} \subseteq L_0$ el espacio modelo, que es un subespacio que incluye todos los reales (funciones constantes). Por tanto, una medida de riesgo se define como $\rho : \mathbb{X} \rightarrow \overline{\mathbb{R}}$, donde $\overline{\mathbb{R}} = \mathbb{R} \cup \{-\infty, +\infty\}$ es la línea real extendida.

Una medida de riesgo se llama coherente si satisface las cuatro propiedades: invarianza traslacional, subaditividad, monotonicidad y homogeneidad positiva (Ahmadi-Javid, 2011; Artzner et al., 1999).

Considerando el problema que se busca resolver (Ver Ecuación 2.1) donde c representa un vector de costo, es incierto, pero sigue una distribución discreta y v la variable de decisión. Luego, c es una variable aleatoria, y por ende, $c^\top v$ también lo es. Con lo anterior, se busca resolver un problema del tipo

$$\begin{aligned} \min \quad & \rho(c^\top v) \\ & v \in V \end{aligned} \tag{2.9}$$

donde ρ es una medida de riesgo cualquiera. Ejemplo de ellas son: Valor del Riesgo (VaR), Valor en riesgo condicional (CVaR), Valor en riesgo entrópico (EVaR), Pérdida media más varianza (Mean Var), Semidesviación superior (USD), etc.

2.3. Métodos de aprendizaje

2.3.1. Regresión Lineal

La regresión lineal es una técnica de análisis de datos que se utiliza para hacer predicciones y se utiliza como uno de los métodos más simples en aprendizaje automático. En el caso de la regresión lineal simple, se crea un modelo con dos variables, una variable de respuesta c , el cual para este caso representará el costo predicho, y una variable explicativa (x_1), con el propósito de predecir la variable de respuesta. Por otro lado, en la regresión lineal múltiple, se amplía el modelo para incluir más de una variable explicativa ($x_1, x_2, \dots, x_n$), lo que da lugar a un modelo con múltiples variables (Tranmer & Elliot, 2008). El modelo de análisis de regresión multivariable se formula del siguiente modo

$$c = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n + \varepsilon$$

donde,

c = variable dependiente

x_i = variable independiente

β_i = parámetro

ε = error

Los supuestos del análisis de regresión multivariable son distribución normal, linealidad, ausencia de valores extremos y no tener vínculos múltiples entre variables independientes (Dagnino et al., 2014).

2.3.2. K-medias

El algoritmo K-medias o K-means es una técnica de aprendizaje no supervisado utilizada para agrupar datos en diferentes conjuntos o clústeres. Su objetivo es dividir un

conjunto de datos en Q clústeres, donde Q es un valor predefinido. Cada clúster representa un grupo de datos con características similares, y el objetivo es minimizar la distancia entre los datos dentro de cada clúster y maximizar la distancia entre los clústeres (Cambronero & Moreno, 2006). Las etapas del algoritmo K-Means son las que se muestran en la Figura 2.2

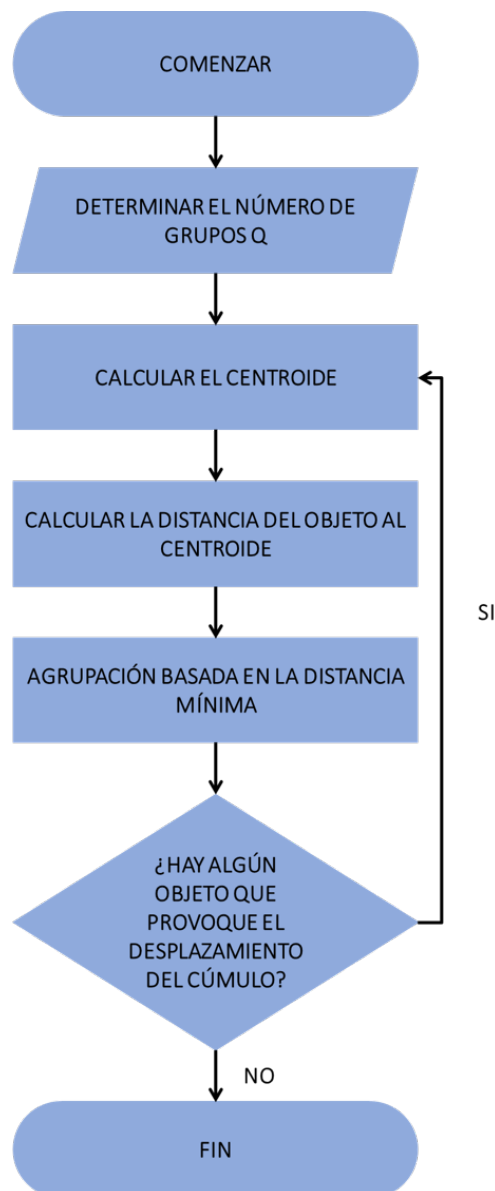


Figura 2.2: Proceso de extracción de datos

El algoritmo de K-medias funciona de la siguiente manera:

1. Determina el valor Q como el número de grupos a generar;
2. Determina de forma aleatoria el centro del clúster o centroide.
3. Agrupa cada dato en el clúster más cercano. Se calcula la distancia de cada dato al centro del conglomerado mediante la ecuación de la distancia euclidiana.

$$distancia(q, r) = |q - r| = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (2.10)$$

Notación:

d = Distancia entre q e r

q = Centro de datos del cluster

r = Datos en el atributo

i = Cada dato

n = Cantidad de datos

q_i = Datos en el centro del cluster i

r_i = Datos en cada dato i

4. Recalcular cada centroide calculando la media de todos los datos del centroide con los miembros actuales del clúster.
5. Volver a agrupar (volver al paso 3). Se repiten en iteraciones sucesivas hasta que se cumpla un criterio de convergencia. El criterio más común es que no haya cambios significativos en la asignación de datos a los clústeres o en la posición de los centroides entre iteraciones consecutivas.
6. Cuando el algoritmo ha convergido, se obtiene una partición del conjunto de datos en Q clústeres, y cada dato se asocia a un clúster específico.

Es importante mencionar que el resultado final del algoritmo puede depender de la inicialización de los centroides, ya que diferentes configuraciones iniciales pueden llevar a soluciones diferentes. Para mejorar la calidad de los resultados, a menudo se realizan múltiples ejecuciones del algoritmo con diferentes configuraciones iniciales y se selecciona la mejor partición basada en algún criterio, como la suma de cuadrados intra-cluster (intra-cluster sum of squares), el índice de silueta, el índice Davies-Bouldin, u otros.

2.3.3. Análisis de componentes principales

El Análisis de Componentes Principales o Principal component analysis (PCA, por sus siglas en inglés) es una técnica estadística que utiliza un principio matemático sofisticado para transformar un conjunto de variables posiblemente correlacionadas en un conjunto reducido de variables no correlacionadas llamadas componentes principales. Esta técnica es ampliamente utilizada en el análisis de conjuntos de datos de gran tamaño, ya que permite una representación más simple y comprensible de los datos originales.

En términos generales, el PCA utiliza una transformación de los vectores para reducir la dimensionalidad de los conjuntos de datos. Mediante una proyección matemática, el conjunto de datos original, que puede incluir muchas variables, a menudo puede interpretarse en sólo unas pocas variables (los componentes principales). Un examen de los datos de dimensión reducida permitirá detectar tendencias, pautas y valores atípicos en los datos con mayor facilidad que si no se hubiera realizado el análisis de los componentes principales (Richardson, [2009](#)).

2.4. Aprendizaje centrado en decisiones

El paradigma del Aprendizaje Centrado en Decisiones o Decision-Focused Learning (DFL, por sus siglas en inglés) entrena modelos de aprendizaje automático (ML) para mejorar las decisiones al estimar parámetros desconocidos y obtener soluciones de alta

calidad. DFL combina los componentes de predicción y optimización en un sistema, optimizando una función de pérdida basada en las decisiones resultantes, ya que la función de pérdida del modelo ML depende de la solución del problema de optimización con restricciones (CO) (Mandi et al., 2023).

En DFL, la función de pérdida depende de la solución de un modelo de optimización, por lo que el solucionador de optimización forma parte del modelo ML. Por tanto, el modelo ML se entrena para optimizar los criterios de evaluación que miden la calidad de las decisiones resultantes. Como las decisiones se toman después de la etapa de optimización, esto implica la combinación de los componentes de predicción y optimización, en un modelo compuesto que genere decisiones completas. Desde esta perspectiva, la generación de los parámetros de predicción $\hat{c}$ es un paso intermedio dentro del enfoque integrado, y el objetivo principal del entrenamiento no es la precisión de $\hat{c}$, sino el error cometido después de la optimización (Mandi et al., 2022; Mulamba et al., 2020).

Una medida del error con respecto a las decisiones prescriptivas del modelo integrado, cuando se usa como función de pérdida para el entrenamiento. La diferencia esencial con la mencionada pérdida de predicción es que mide el error en $v^*(\hat{c})$, en vez de en $\hat{c}$. El valor objetivo obtenido utilizando la predicción $v^*(\hat{c})$ suele ser subóptimo en comparación con los verdaderos parámetros objetivos c . A menudo, el objetivo final es generar predicciones $\hat{c}$ con una solución óptima $v^*(\hat{c})$ cuyo valor objetivo en la práctica (es decir, $f(v^*(\hat{c}), c)$) se acerque al valor óptimo con información completa $f(v^*(c), c)$.

En tales casos, una noción importante de pérdida de tarea es el concepto de arrepentimiento, definido como la diferencia entre el valor objetivo óptimo con información completa y el valor objetivo logrado por la decisión prescriptiva. En otras palabras, es la medida de cuánto se desvía la decisión $v^*(\hat{c})$ de la solución óptima $v^*(c)$ bajo los parámetros verdaderos c :

$$\begin{aligned} \text{Regret}(\hat{c}, c) &= f_c(v^*(\hat{c}), c) - f_c(v^*(c), c) \\ v^* &\in \arg \min f_c(v) \end{aligned} \tag{2.11}$$

En los problemas de minimización, el arrepentimiento siempre es un valor positivo. Sin embargo, puede ser igual a cero cuando la optimización de los valores estimados conduce a la verdadera solución óptima o a una solución equivalente. El objetivo de la predicción y la optimización es aprender los parámetros del modelo β que determinen $\hat{c} = m(\beta, x)$ de manera que se minimice el arrepentimiento de las predicciones resultantes, es decir, $\arg \min_{\beta} \mathbb{E}[\text{Regret}(m(\beta, x), c)]$. Cuando se utiliza la retropropagación como mecanismo de aprendizaje, no es posible emplear directamente el arrepentimiento como función de pérdida debido a que no es continua y requiere diferenciar sobre el argmin $v^*(c)$. Por lo tanto, el desafío principal en la predicción y optimización radica en encontrar una función de pérdida diferenciable y computacionalmente eficiente que considere la estructura de f y $v^*(.)$ de manera más general. El aprendizaje en un conjunto de S instancias de formación puede formularse en el marco empírico de minimización del riesgo como

$$\arg \min_{\beta} \mathbb{E}[\text{Regret}(m(\beta, x), c)] \quad (2.12)$$

$$\approx \arg \min_{\beta} \frac{1}{S} \sum_{i=1}^N f_{c^i}(v^*(\hat{c}^i)) - f_{c^i}(v^*(c^i)) \quad (2.13)$$

El objetivo principal del proceso de aprendizaje es optimizar las decisiones predichas tomadas con base en una distribución de variables características $x \sim X$, considerando un criterio de evaluación aplicado a esas decisiones. Mientras que el modelo de aprendizaje automático m_{β} se entrena para predecir $\hat{c}$, su rendimiento se evalúa en función de las soluciones óptimas correspondientes $v^*(\hat{c})$ (Elmachtoub & Grigas, 2022). En el presente trabajo se emplea el término “*Predecir, luego Optimizar*” *inteligentemente* (SPO, por sus siglas en inglés) para referirse al problema de hacer predicciones sobre $\hat{c}$ con el objetivo de mejorar la evaluación de $v^*(\hat{c})$.

Sin embargo, lo anterior es difícil de resolver, tanto en teoría como en práctica, por tanto, en definitiva se ocupó una función sustituta convexa llamada SPO+ que es básicamente una derivación de $\ell_{\text{SPO}}(\cdot, \cdot)$ que permitirá producir soluciones aproximadas y más fáciles de resolver (Elmachtoub & Grigas, 2022).

$$\ell_{\text{SPO}+}(\hat{c}, c) := \max_{v \in S} \left\{ c^T v - 2\hat{c}^T v \right\} + 2\hat{c}^T v^*(c) - z^*(c) \quad (2.14)$$

2.5. Criterios de evaluación

Para llevar a cabo la evaluación y comparación del impacto de la calidad de las soluciones generadas por ambas metodologías, se utilizarán dos medidas fundamentales: los tiempos computacionales requeridos para encontrar una solución y el *regret*.

- **Tiempos computacionales:** Se refiere al tiempo promedio en unidad de segundos que cada modelo necesita para encontrar una solución al problema tratado.
- **Arrepentimiento o Regret:** Se empleará la moción de arrepentimiento (también conocida como pérdida SPO (Elmachtoub & Grigas, 2022)) y se define como la diferencia de valor objetivo entre una solución óptima (que utiliza el vector de coste real, pero desconocido) y una solución obtenida utilizando el vector de coste previsto (Ver Ecuación 2.11). Por consiguiente, la calidad de las decisiones se evalúa midiendo el error de decisión entre la predicción y la solución real.

Estos resultados serán evaluados mediante una inspección visual y de análisis de estadísticas descriptivas. La evaluación de las estadísticas descriptivas se realizará de la siguiente manera.

1. **Mediana (línea en la caja):** Su evaluación se realizará con respecto al promedio general de los modelos y será evaluado cada $1/3$ de valor de dicho promedio.
 - **Muy buena:** 4 puntos. Si la mediana se encuentra dentro de un rango de $2/3$ inferior al valor del promedio general (Esto, considerando la definición de *Regret*).
 - **Buena:** 3 puntos. Si la mediana se encuentra dentro de un rango de $1/3$ inferior al valor del promedio general.
 - **Promedio:** 2 puntos. Si la mediana es igual o tiene una desviación máxima de 0.1 al promedio.
 - **Mala:** 1 puntos. Si la mediana se encuentra dentro de un rango de $1/3$ superior al valor del promedio general
 - **Muy mala:** 0 puntos. Si la mediana se encuentra dentro de un rango de $2/3$ superior al valor del promedio general o está considerablemente alejada del promedio.

2. **Caja (rango intercuartílico):** La evaluación se realizará con respecto al promedio general de los modelos y se dividirá en tercios de su valor. La evaluación según tercios se presenta de la siguiente manera:
 - **IQR muy bajo:** 4 puntos. Si el IQR se encuentra dentro de un rango de $2/3$ inferior al valor del promedio general.
 - **IQR bajo:** 3 puntos. Si el IQR se encuentra dentro de un rango de $1/3$ inferior al valor del promedio general.
 - **IQR promedio:** 2 puntos. Si el IQR es igual al promedio general.
 - **IQR alto:** 1 punto. Si el IQR se encuentra dentro de un rango de $1/3$ superior al valor del promedio general.
 - **IQR muy alto:** 0 puntos. Si el IQR se encuentra dentro de un rango de $2/3$ superior al valor del promedio general.

Un IQR más pequeño indica una mayor concentración de datos.

3. **Bigotes (líneas que se extienden fuera de la caja):** Se denominarán *Bigotes* a la suma del bigote superior e inferior para evaluar la variabilidad:

$$\text{Bigote superior} = \text{Valor máximo} - Q_3$$

$$\text{Bigote inferior} = Q_1 - \text{Valor mínimo}$$

En este caso, un menor valor de la suma indicará que los bigotes están menos dispersos, lo que sería considerado mejor.

La puntuación se asigna en orden decreciente de los resultados:

- **1er puesto:** 5 puntos.
- **2do puesto:** 4 puntos.
- **3er puesto:** 3 puntos.
- **4to puesto:** 2 puntos.
- **5to puesto:** 1 punto.
- **6to puesto:** 0 punto.

4. **Valores atípicos (outliers):** Para su evaluación, se considerará la lejanía del último dato atípico con respecto a cero. La puntuación se asigna en orden decreciente de los resultados, donde el primer puesto lo ocupará el modelo cuyo último dato atípico esté más cercano a cero, y el quinto puesto se asignará al modelo que esté más alejado de cero.

- **1er puesto:** 5 puntos.
- **2do puesto:** 4 puntos.
- **3er puesto:** 3 puntos.
- **4to puesto:** 2 puntos.

- **5to puesto:** 1 punto.
- **6to puesto:** 0 punto.

Capítulo 3

Metodología

En el presente capítulo, para poder dar respuesta a la pregunta de investigación, la metodología será la siguiente: PRIMERO, se generan datos sintéticos de acuerdo a cada problema que se estará considerando. Luego se resolverá dicho problema bajo dos ramas que por simplicidad bautizaremos: “Predict-then-Optimize” y programación estocástica. Bajo estas dos ramas resolveremos diversos problemas, los cuales se explican a continuación:

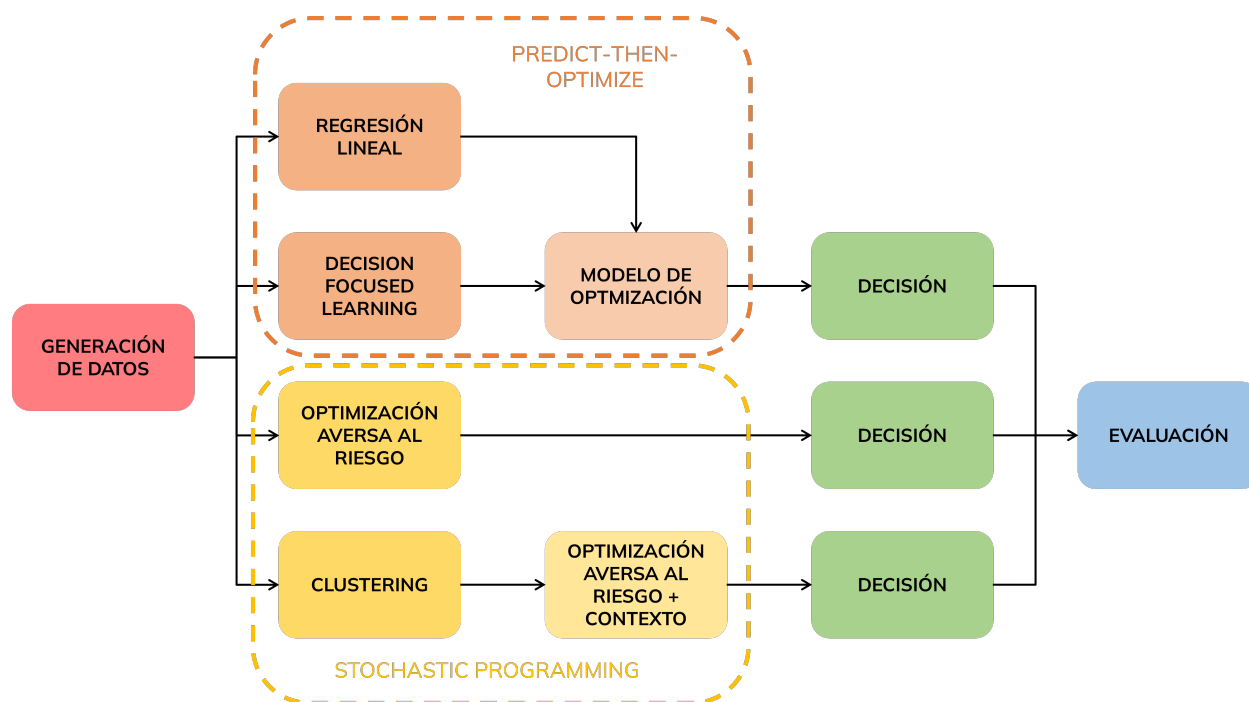


Figura 3.1: Proceso Metodológico

■ PREDICT-THEN-OPTIMIZE

• Regresión lineal

Este modelo se utiliza bajo el enfoque "Predict-then-Optimize", donde se minimiza el error de predicción. Primero, se realizan estimaciones puntuales utilizando un modelo de regresión lineal. Luego, estas predicciones se utilizan como insumos para resolver algoritmos de optimización de rápida ejecución y así obtener la solución óptima.

• Decision-Focused Learning

En este modelo, se emplea el marco SPO+ (Smart "Predict then Optimize"), que minimiza el error de decisión. A diferencia del enfoque predictivo tradicional, este modelo integra el problema de optimización nominal desde el principio. Aunque también se obtienen estimaciones puntuales, estas si consideran el problema de optimización nominal que se busca resolver. Finalmente, para obtener la solución óptima, se resuelven algoritmos de optimización típicos o

modelos de programación lineal.

■ PROGRAMACIÓN ESTOCÁSTICA

• **Optimización con aversión al riesgo**

Este enfoque considera el uso de medidas de riesgo. La forma de tratar este problema es primero estimar las distribuciones de probabilidad a partir de los datos de costos, luego se resuelve un modelo averso al riesgo del que se obtiene un valor esperado y finalmente una solución.

• **Optimización con aversión al riesgo + contexto**

Este problema considera tanto los costos como las características. Se considera que se tienen distribuciones distintas de costos y, por tanto, es posible agrupar dentro del conjunto de características, del cual se tendrá una distribución de probabilidad distinta asociada a cada grupo. Y con lo anterior, resolver un modelo averso al riesgo y obtener una solución por grupo.

Finalmente, los modelos se evalúan considerando los criterios de evaluación mencionados anteriormente.

3.1. Generación de datos

En el contexto de esta investigación, se utilizan datos sintéticos generados artificialmente para realizar experimentos y para poder evaluar el rendimiento del trabajo propuesto. Los conjuntos de datos sintéticos incluyen características $\mathbf{x}$ y coeficientes de coste $\mathbf{c}$. El vector de características $\mathbf{x}^i \in \mathbb{R}^k$ sigue una distribución gaussiana multivariante estándar $\mathcal{N}(0, \mathbf{I})$, y el coste correspondiente $\mathbf{c}^i \in \mathbb{R}^d$ procede de una función polinómica $f(\mathbf{x}^i)$ multiplicada por un ruido aleatorio $\epsilon^i \sim U(1 - \bar{\epsilon}, 1 + \bar{\epsilon})$.

En general, hay varios parámetros que se pueden controlar:

- **Escenarios (e):** Tamaño de los datos.

- **Atributos (k):** Dimensión de las características de los costes $\mathbf{c}$.
- **Deg (deg):** Grado polinómico de la función $f(\mathbf{x}^i)$. Este indica el grado de polinomio de datos en el modelo; cuanto mayor sea el valor de Deg , más se desvía de un modelo lineal y más errores habrá. Se experimentó con grados de 1, 2, 4, 6 y 8.
- **Noise ($\bar{\epsilon}$):** Semiancho de ruido de ϵ . Se muestreará $U(0.5, 1.5)$.

Posteriormente, la herramienta a utilizar para dividir la data en conjuntos de entrenamiento y testeo es la función *train_test_split()* de la biblioteca SCIKIT-LEARN. Scikit-learn es una librería de software gratuito de aprendizaje automático para el lenguaje de programación Python (Hao & Ho, 2019).

La división de datos se realiza con el propósito de entrenar un modelo utilizando el conjunto de entrenamiento (generalmente el 80 % de los datos) y luego evaluar su rendimiento en un conjunto independiente, llamado conjunto de testeo (normalmente el 20 % restante) (Gholamy et al., 2018). Esta separación ayudará a estimar el rendimiento del modelo en datos no vistos previamente y evaluar si el modelo generaliza bien para nuevas instancias (Gholamy et al., 2018).

Es una práctica común en la mayoría de los algoritmos de aprendizaje automático para evaluar y comparar modelos, y prevenir problemas de sobreajuste (overfitting) (Gholamy et al., 2018). Al dividir los datos en conjuntos de entrenamiento y testeo, se evita que el modelo memorice los datos de entrenamiento y pueda realizar predicciones más precisas en datos nuevos.

3.1.1. Problema del camino más corto

Se adaptó el problema de Elmachtoub y Grigas, 2022. Se trata de un problema de un camino más corto en una cuadrícula de 3×3 , 5×5 y 7×7 , con el objetivo de ir desde la esquina suroeste de la cuadrícula hasta la esquina este, donde los bordes pueden ir hacia

el noreste. Se utiliza una distribución gaussiana multivariada para obtener vectores de características x^i .

Obsérvese que la dimensión del vector de costes d , corresponde al número total de aristas de la red cuadriculada que se esté empleando y que k es el número determinado de características.

En primer lugar, se generó una matriz aleatoria $\mathcal{B} \in \mathbb{R}^{d \times k}$ que sigue la distribución de Bernoulli con probabilidad 0.5, codifica los parámetros del modelo verdadero. Luego se generan los vectores de costos según la siguiente fórmula:

$$c^{ij} = \left[\left[\frac{1}{3.5^{deg}} \left(\frac{1}{\sqrt{k}} (\mathcal{B}\mathbf{x}^i)_j + 3 \right)^{deg} + 1 \right] \cdot \epsilon_{ij} \right] \quad (3.1)$$

donde c^{ij} es el componente j -ésimo del vector de costos c^i y $(\mathcal{B}\mathbf{x}^i)_j$ denota el j -ésimo componente de $\mathcal{B}\mathbf{x}^i$. El modelo subyacente utiliza una red neuronal sin capas ocultas, es decir, un modelo de regresión.

Se considerarán conjuntos de datos de tamaño 200, 1000 y 5000, los cuales se dividirán en instancias de entrenamiento y testeo.

3.1.2. Problema de la mochila 0-1

Dado que los coeficientes inciertos supuestamente solo existen en la función objetivo, los pesos de los elementos permanecen fijos en todos los datos. Estos pesos de la mochila fueron recuperados de un archivo generado en la investigación de Mulamba et al., 2020, donde se generaron sintéticamente asignando aleatoriamente un peso de 3, 5 o 7 a cada uno de los 48 objetos considerados.

$$c^{ij} = \left[\left[\frac{5}{3.5^{deg}} \left(\frac{1}{\sqrt{k}} (\mathcal{B}\mathbf{x}^i)_j + 3 \right)^{deg} + 1 \right] \cdot \epsilon_{ij} \right] \quad (3.2)$$

donde, al igual que antes, $\mathcal{B}$ es una matriz aleatoria.

Se considerarán conjuntos de datos de tamaño 200, 750 y 1500, los cuales se dividirán en instancias de entrenamiento y testeo. Se estudiarán tres casos de este problema de mochila con capacidades de 60, 120 y 180.

3.2. Regresiones lineales y su implementación

Como era de esperar, se recurrirá a los conceptos explicados en el Capítulo 2. Se utilizará la Regresión Lineal en su forma multivariable para estimar los valores de los coeficientes β^i , de modo que se minimice la suma de los errores al cuadrado (Ver Ecuación 3.4).

$$\min_{f \in \mathcal{H}} \frac{1}{E} \sum_{i=1}^E \ell(f(x^i, c^i)) \quad (3.3)$$

La Ecuación 3.3 es un modelo de predicción f^* que se entrena (en el caso del presente estudio, una regresión lineal múltiple) usando el principio de minimización de riesgos empíricos (ERM). Por tanto, el correspondiente problema empírico de minimización de riesgos para encontrar el mejor el modelo lineal $\mathcal{B}^*$ se convierte en

$$\approx \min_{\mathcal{B}} \frac{1}{E} \sum_{i=1}^E \|\mathcal{B}x^i - c^i\|^2 \quad (3.4)$$

Sin embargo, dicha función minimizará la pérdida de entrenamiento empírico, utilizando como función de pérdida la del error cuadrático medio en la predicción.

Para el modelo anterior del paradigma “Predict, then Optimize”, lo que se hace es entrenar un modelo de regresión lineal multivariable empleando una matriz de atributos

asociada a su respectiva arista u objeto (dependiendo del problema que se esté tratando) y a un vector de costos de todos los escenarios. Esto se logra mediante la implementación del método `LINEARREGRESSION` de la librería de `SCIKIT-LEARN` en Python, en particular de la función `LinearRegression().fit()`. De este modelo se obtiene un intercepto y tantos coeficientes como atributos por instancia se tengan.

Por otro lado, se tiene el modelo bajo el enfoque Decision-Focused learning, que también entra en el paradigma de predicción y luego optimización, pero que también es “inteligente”, pues mediante la construcción de funciones de pérdida contempla medir el error de decisión al predecir vectores de costos, aprovechando así la estructura del problema de optimización nominal. De aquí en adelante, se hará referencia a estas funciones de pérdida como Smart “Predict, then Optimize” (SPO).

3.3. Formulación lineal de la pérdida Smart “Predict, then Optimize”

Para diseñar los problemas lineales bajo el enfoque SMART “PREDICT-THEN-OPTIMIZE” (Elmachtoub & Grigas, 2022) es necesario mencionar los 4 ingredientes clave del marco:

1. Un problema de optimización nominal ($\mathcal{P}$) de la forma:

$$\min_{v \in V} c^\top v$$

donde V es la región y $v^*(c)$ es el oráculo de optimización.

2. Datos de entrenamiento de E muestras de la forma (x^i, c^i) con x^i es un vector de características que representa la información contextual asociada a c^i .
3. Una clase de hipótesis $\mathcal{H}$ de modelos de predicción de vectores de costos.
4. Finalmente, una función de pérdida ℓ que cuantifica el error de hacer una predicción

$\hat{c}$ cuando el valor real es c .

Considerando que se cuentan con los ingredientes 1 y 2, se debe de determinar un modelo de predicción $f^* \in \mathcal{H}$ que use ERM y considere la función de pérdida SPO, por lo que se tiene

$$\min_{f \in \mathcal{H}} \frac{1}{E} \sum_{i=1}^E \ell_{\text{SPO}+}(f(x^i), c^i) \quad (3.5)$$

$$\approx \min_{B \in \mathbb{R}^{d \times p}} \frac{1}{E} \sum_{i=1}^E \ell_{\text{SPO}+}(Bx^i, c^i) + \lambda \Omega(B) \quad (3.6)$$

En la Ecuación 3.6, en el caso de los predictores lineales, $\mathcal{H} = f : f(x) = Bx$, donde B es una matriz de pesos que se utiliza para combinar las características de entrada x^i en la función objetivo. La función $\text{SPO}+(Bx^i, c^i)$ representa el SPO LOSS, que es una medida de la discrepancia entre la predicción de costos Bx^i y el costo real c^i . El objetivo es minimizar la suma de los errores en la predicción de costos para todos los datos de entrenamiento, ponderados por $\frac{1}{E}$, donde E es el número de datos de entrenamiento. Así mismo, la Ecuación 3.7 incluye un término de regularización $\lambda \Omega(B)$, la cual representa una función que penaliza la complejidad del modelo. El parámetro de regularización λ controla el peso relativo entre la función objetivo y el término de regularización.

Una forma de resolver la última ecuación es reformulando mediante dualidad (Ver la Proposición 3.7). Esta proposición proporciona una reformulación del problema de optimización de riesgo mínimo regularizado (ERM) para el enfoque SPO+ cuando la región V es un politopo.

La proposición establece que si la región factible V , definida como $v : Av \geq b$, es un politopo, entonces el problema de ERM regularizado para SPO+ es equivalente a un problema de optimización específico.

$$\begin{aligned}
& \min_{B,p} \quad \frac{1}{E} \sum_{i=1}^E \left[-b^\top p^i + 2(v^*(c^i)x^{i\top}) \bullet B_z^*(c^i) \right] + \lambda \Omega(B) \\
& \text{s.a.} \quad A^\top p^i = 2Bx^i - c^i \quad \forall i \in \{1, \dots, E\} \\
& \quad p^i \in \mathbb{R}^m, p^i \geq 0 \quad \forall i \in \{1, \dots, E\} \\
& \quad B \in \mathbb{R}^{d \times p}
\end{aligned} \tag{3.7}$$

La Ecuación 3.7 se trata de un problema lineal, aunque $\Omega(\cdot)$ pudiera no serlo, y resuelve la ecuación 3.6 (Elmachtoub & Grigas, 2022)

3.3.1. Problema del camino más corto

El problema de programación lineal del camino más corto es el siguiente:

$$\begin{aligned}
& \min \quad \sum_{(i,j) \in A} c_{ij} v_{ij} \\
& \text{s.a.} \quad \sum_{j:(i,j) \in A} v_{ij} - \sum_{j:(j,i) \in A} v_{ji} = \begin{cases} 1 & \text{si } i = s \\ 0 & \text{si } i \neq s, t \\ -1 & \text{si } i = t \end{cases} \quad i \in N \\
& \quad v_{ij} \geq 0, \quad \forall (i,j) \in A
\end{aligned}$$

Entonces, considerando $\hat{c} = Bx$, se escribe este problema en el formato dual con base en la función de pérdida SPO+ ya antes vista:

$$\ell_{\text{SPO}^+}(Bx, c) := \max_{v \text{ camino s-t}} \{ (c^\top - 2(Bx)^\top) v \} + 2(Bx)^\top v^*(c)$$

por tanto, la función objetivo del problema en su forma primal es

$$\max \sum_{i,j \in A} \left(c_{ij} - 2 \sum_{i,j \in A} \left[\beta_0 + \sum_{k=1}^K \beta_k \cdot x_{ij}^k \right] \right) v_{ij}$$

Se declaran variables duales (p) y el problema de minimizar el riesgo empírico queda de la siguiente forma para ser resuelto en solver

$$\begin{aligned} \min_{\beta, p} \quad & \frac{1}{E} \sum_{e=1}^E p_s^e - p_t^e + 2 \sum_{i,j \in A} \left[\beta_0 + \sum_{k=1}^K \beta_k \cdot x_{ij,k}^e \right] v^*(c_{ij}^e) \\ \text{s.a.} \quad & p_i^e - p_j^e \geq c_{ij}^e - 2 \sum_{i,j \in A} \left[\beta_0 + \sum_{k=1}^K \beta_k \cdot x_{ij,k}^e \right] \quad \forall (i, j) \in A, \forall k \\ & p_i \text{ libre} \end{aligned}$$

3.3.2. Problema de la mochila 0 - 1

Para el problema de la mochila bajo el presente enfoque se comienza resolviendo el problema común, tal que

$$\begin{aligned} \max \quad & \sum_{i=1}^E c_i^\top v_i \\ \text{s.a.} \quad & w_i^\top v_i \leq M \\ & v_i \in \{0, 1\} \end{aligned}$$

Se hace el reemplazo c por $-c = c'$ y se resuelve en el solver

$$\begin{aligned}
& \text{mín} && \sum_{i=1}^N c_i' v_i \\
& \text{s.a.} && w_i' v_i \leq M \\
& && v_i \in \{0, 1\}
\end{aligned}$$

Del problema anterior se obtienen soluciones óptimas $v^*(c^i)$ y un costo real $z^*(c^i)$ por escenario, datos que forman parte de la construcción de la siguiente formulación

$$\begin{aligned}
& \text{mín} && \frac{1}{E} \sum_{i=1}^E l(\hat{c}_i, c_i) \\
\Leftrightarrow & \text{mín} && \frac{1}{E} \sum_{i=1}^E \max_{v \in S} \left\{ (c_i - 2(\hat{c}_i))' v \right\} + 2(\hat{c}_i)' v^*(c_i) \\
\Rightarrow & \text{mín} && \frac{1}{E} \sum_{i=1}^E [l_i + 2(\hat{c}_i)' v^*(c_i)] \\
& \text{s.a.} && l_i \geq (c_i - 2(\hat{c}_i))' v_i^s \quad \forall i, \forall v^s \in \text{Solution pool}
\end{aligned}$$

Donde, $\hat{c}_i = \beta_0 + \sum_{k=1}^K \beta_k \cdot x_k^i$.

De las formulaciones SPO+, al igual que de las regresiones lineales, se obtienen predicciones que permitirán estimar nuevos costos o utilidades y mediante un modelo de programación lineal, se tomarán decisiones con base en la nueva información.

3.4. Optimización aversa al riesgo

Las medidas a implementar en la construcción de los modelos son las siguientes:

3.4.1. Valor en riesgo condicional

El valor en riesgo condicional o Conditional Value-at-Risk (CVaR), también conocido como el déficit esperado, cuantifica las pérdidas tomando un promedio ponderado de las pérdidas “extremas” en la cola de la función de distribución más allá del punto de corte del valor en riesgo (VaR). $CVaR_\varepsilon(X)$ es igual a la expectativa condicional de X sujeto a $X \geq VaR_\varepsilon(X)$. Esta definición es la base para el nombre de valor en riesgo condicional. Esta medida de riesgo satisface la propiedad de subaditividad, es convexa y uniextremo, lo cual facilita la implementación de algoritmos de optimización y control.

Existen modelos de optimización matemática para resolver CVaR (Rockafellar, Uryasev et al., 2000). Para una formulación de 2.9 (Chicoisne et al., 2018) se tendrá que $CVaR_{\varepsilon \in [0,1]}$ (nivel), estará representado por una aproximación discreta z con realizaciones z_i de probabilidad p_i para $i \in \{1, \dots, E\}$.

$$CVaR(z) = \min_{t \in \mathbb{R}} t + \varepsilon^{-1} \sum_{i=1}^E p_i [z_i - t]_+ \quad (3.8)$$

donde $[z_i - t]_+$ es convexo lineal por partes. Entonces, el problema 2.9 con $z = c^T v$ queda expresado

$$\begin{aligned} \min_v \quad & CVaR_\varepsilon(z) \\ & v \in V \end{aligned} \quad (3.9)$$

Por tanto, en la Ecuación 3.9, en las realizaciones z , en realidad se tienen vectores de costos por escenario, es decir, c^i de probabilidad p_i para cualquier $i \in \{1, \dots, E\}$. Por tanto,

se tiene dicho problema de optimización en el que se está minimizando en v .

Finalmente escribiendo con $\eta_i = [(c^i)^\top v - t]_+$ se tiene

$$\begin{aligned}
 \min_{v, t, \eta_i} \quad & t + \varepsilon^{-1} \sum_{i=1}^E p_i(\eta_i) \\
 \text{s.a} \quad & \eta_i \geq (c^i)^\top v - t & \forall i \in \{1, \dots, E\} \\
 & \eta_i \geq 0 & \forall i \in \{1, \dots, E\} \\
 & t \in \mathbb{R} \\
 & v \in V \\
 & + \text{Restricciones propias del problema tratado}
 \end{aligned} \tag{3.10}$$

Donde, si $(c^i)^\top v - t$ es negativo, la segunda desigualdad queda sin efecto, y η_i sería el mínimo, es decir, 0. En caso contrario, si fuese positivo, eso significaría que la segunda desigualdad se cumpliría y de paso se verificaría la primera, por tanto, η_i se puede mover y estaría acotado por $(c^i)^\top v - t$, pero como se está minimizando, $(c^i)^\top v - t$ es lo menor que puede ser y por tanto, respondería a que sea $[(c^i)^\top v - t]_+$.

Se trabajará con $\varepsilon = (0.05, 0.1)$, lo que representa 95 % y 90 % de confiabilidad, respectivamente.

3.4.2. Pérdida media + varianza

La Pérdida media + varianza o por su nombre en inglés Mean Variance Risk (o simplemente Mean Var, término que se empleará de aquí en más) combina la media y la varianza en una medida que captura tanto el costo esperado como el riesgo asociado con una distribución de probabilidad.

$$\rho(z) := \min \mathbb{E}[z] + \alpha \text{Var}[z] \quad (3.11)$$

Ahora, aproximando la Ecuación 3.11, se tiene

$$\approx \min \frac{1}{E} \sum_{i=1}^E c_i^\top v + \frac{\alpha}{E} \sum_{i=1}^E \left(c_i^\top v - \frac{1}{E} \sum_{j=1}^E c_j^\top v \right)^2 \quad (3.12)$$

Considerando la Ecuación 3.12, se busca encontrar un equilibrio óptimo entre: la primera parte de la expresión $\frac{1}{E} \sum_{i=1}^E c_i^\top v$, que se enfoca en minimizar el costo promedio esperado y la segunda parte de la expresión $\frac{\alpha}{E} \sum_{i=1}^E \left(c_i^\top v - \frac{1}{E} \sum_{j=1}^E c_j^\top v \right)^2$, que se concentra en la reducción de la incertidumbre. Esta última parte de la ecuación busca reducir la variabilidad o dispersión en los costos. Al hacerlo, se minimiza la incertidumbre, lo que proporciona mayor estabilidad y predictibilidad en el resultado final. Vale destacar que esta parte representa la varianza de los costos, la cual está ponderada por un factor α . Al incrementar α , se otorga más importancia a la minimización de la varianza, lo que indica un deseo de no solo reducir el costo promedio esperado, sino también la variabilidad en los costos entre diferentes escenarios.

En relación con la resolución, se observa que la Ecuación 3.12 plantea un problema cuadrático, lo que implica un esfuerzo computacional considerable durante su solución. Por ende, una estrategia para simplificar su resolución consiste en sustituir la varianza por un valor absoluto, lo que conduce a una formulación lineal.

$$\min \mathbb{E}(c^\top v) + \alpha \mathbb{E} \|c^\top v - \mathbb{E}(c^\top v)\| \quad (3.13)$$

Considerando $|c^\top v - \mathbb{E}(c^\top v)| = \Delta_s$ se requiere la adición de restricciones adicionales en el problema. Esto da lugar a la siguiente formulación:

$$\begin{aligned}
 \text{mín} \quad & \frac{1}{E} \sum_{i=1}^E c_i^\top v + \frac{\alpha}{E} \sum_{i=1}^E \Delta_E \\
 \text{s.a.} \quad & \Delta_E \geq c_j^\top v - \frac{1}{E} \sum_{i=1}^E c_i^\top v & \forall j \in \{1, \dots, E\} \\
 & \Delta_E \geq \frac{1}{E} \sum_{i=1}^E c_i^\top v - c_j^\top v & \forall j \in \{1, \dots, E\} \\
 & + \text{Restricciones propias del problema tratado}
 \end{aligned} \tag{3.14}$$

Finalmente, queda la Ecuación 3.14, la cual es un problema lineal y con variables enteras, donde la primera y segunda restricción representan al módulo. Estas restricciones garantizan que la diferencia absoluta entre cada costo específico y el costo promedio sea mayor o igual a cero, y como se está minimizando, Δ_E tomará el mayor valor posible. Se tomarán valores α igual a 0.0005 y 0.001.

Una distinción significativa que vale la pena señalar entre esta medida y CVaR es que esta última no penaliza de manera simétrica la desviación por la cantidad de interés alrededor de la media, lo cual es deseable en contextos donde el enfoque de control modela una pérdida, por ejemplo.

3.5. K-medias

Llegado a este punto, es cuando se necesita un método para contextualizar las medidas de riesgos. Dicha metodología, es K-medias (o K-means), el cual es uno de los métodos no jerárquicos de clustering más populares, debido a la simplicidad del algoritmo y a la rapidez de selección del centro del cluster (centroide) (Gustriansyah et al., 2020).

Para poder clasificar en diversos grupos o clusters las características x del conjunto de entrenamiento, se utiliza el módulo *sklearn.cluster* de la biblioteca de SCIKIT-LEARN, en específico el método KMEANS, el cual es un algoritmo que agrupa e intenta separar en grupos de igual varianza, minimizando un criterio conocido como INERCIA o suma de cuadrados dentro de dicho clúster (similitud intracluster), mientras que intenta maximizar suma de cuadrados entre clústeres (similitud interclúster). La inercia puede reconocerse como una medida de cuán internamente coherente son los clusters.

Sin embargo, uno de los principales problemas del método de K-means es como determinar el número óptimo de clústers w . Investigaciones realizadas por Subbalakshmi et al., 2015 han demostrado que la precisión del método puede ser mayor si se selecciona adecuadamente el valor inicial y el número de clústeres (Wei et al., 2016; H.-H. Wu et al., 2009). Existen varias formas de estimar el número óptimo de conglomerados w . En este estudio, el número óptimo de conglomerados w se medirá utilizando el método del índice de Davies-Bouldin (Davies & Bouldin, 1979) y no otro, como por ejemplo, Silhouette Score, dado que es computacionalmente menos costosa (Petrovic, 2006).

El índice de Davies-Bouldin (*davies_bouldin_score()*), es una métrica de validación de conglomerados que mide la calidad de la agrupación en función de qué tan separados están los grupos y que tan diferentes son entre sí, y se utiliza comúnmente para identificar el número óptimo de grupos (clústeres) en los datos.

Las mediciones basadas en el índice de Davies-Bouldin (DBI) buscan maximizar la

distancia entre clústeres g_i y g_j y, al mismo tiempo, minimizar la distancia dentro de cada clúster. Cuando la distancia interclúster es máxima, implica que la similitud entre clústeres es baja, lo que permite una mejor visualización de las diferencias entre ellos. Por otro lado, cuando la distancia intraccluster es mínima, significa que los objetos dentro de cada clúster son altamente similares entre sí (Davies & Bouldin, [1979](#); Gustriansyah et al., [2020](#)).

Capítulo 4

Comparativa Empírica: Optimización Aversa al Riesgo vs. Predecir y Optimizar

A continuación, se presentan los resultados obtenidos para ambos problemas de optimización: el camino más corto y el problema de la mochila. Estos fueron abordados utilizando el lenguaje de programación Python, en su versión 3.11. Para resolver estos desafíos, se empleó el Solver Gurobi, una herramienta reconocida por su robustez y eficiencia. Para obtener más referencias, consulte [Gurobi Optimization](#).

Se presentan gráficos de cajas y bigotes que ilustran la distribución de los resultados obtenidos mediante diversos modelos. Estos se encuentran segmentados en conjuntos de entrenamiento y prueba para una evaluación más detallada. Con el propósito de distinguir entre modelos que hacen uso de información contextual global y aquellos que incorporan información local, se ha establecido que los nombres de los primeros llevarán el sufijo *RA* (abreviatura de *Repeating attributes*). Este añadido se justifica por la práctica común de mantener un único vector de características en todo el conjunto de datos con el fin de simular situaciones de la vida real. En contraste, para los modelos que consideran información local, el nombre del modelo permanecerá inalterado, pero se podrá hacer referencia a ellos utilizando el término *NRA* (abreviatura de *Not Repeating Attributes*).

Esto en la nomenclatura tiene como objetivo proporcionar una distinción clara entre los distintos tipos de modelos empleados en el análisis de los resultados.

Por tanto, la nomenclatura será la siguiente:

- **SPO RA:** Modelo bajo el enfoque Smart “Predict-the-Optimize” con repetición del vector de atributos por escenario.
- **RL RA:** Modelo bajo el enfoque “Predict-then-Optimize” con repetición del vector de atributos por escenario.
- **CVaR 0.05:** Modelo bajo un enfoque estocástico, considerando como medida aversa al riesgo, el Valor en riesgo condicional con un 95 % ($\varepsilon = 0.05$) de confiabilidad.
- **CVaR 0.1:** Modelo bajo un enfoque estocástico, considerando como medida aversa al riesgo, el Valor en riesgo condicional con un 90 % ($\varepsilon = 0.1$) de confiabilidad.
- **Mean Var 0.0005:** Modelo bajo un enfoque estocástico, considerando como medida aversa al riesgo, la media y la varianza ponderada por $\alpha = 0.0005$.
- **Mean Var 0.001:** Modelo bajo un enfoque estocástico, considerando como medida aversa al riesgo, la media y la varianza ponderada por $\alpha = 0.001$.

Además, se exhibirán en cuadros los tiempos promedios de ejecución en unidad de segundos de dichos modelos, teniendo en cuenta el conjunto de testeo. Se destacarán en **rojo** los tiempos máximos y en **azul** los tiempos mínimos para cada configuración, proporcionando así una visión más detallada y comprensible.

4.1. Resultados computacionales sobre el problema del camino más corto

La ejecución de los experimentos para SPP se realizaron en un sistema computacional que cuenta con sistema operativo Windows 11 y con las prestaciones de hardware: AMD Ryzen 5 5600H 3.30 GHz con Radeon Graphics, 8GB RAM y 512GB SSD.

El caso base para este estudio será definido con los siguientes parámetros: Escenarios = 1000, Grilla = $5 \cdot 5$, Atributos = 10 y Grado polinómico = 8. Estos parámetros serán modificados sistemáticamente para analizar cómo responden los modelos a dichas variaciones.

1^{RA} CONFIGURACIÓN: Grilla = $\{3 \cdot 3, 5 \cdot 5, 7 \cdot 7\}$, Esc = 1000, $k = 10$, Deg = 8.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Grilla	$3 \cdot 3$	0.008	~ 0	0.008	1.666	~ 0	1.666
	$5 \cdot 5$	0.044	~ 0	0.044	16.5	~ 0	16.5
	$7 \cdot 7$	0.046	~ 0	0.046	23.178	~ 0	23.178

Tabla 4.1: Tiempos computacionales [s] de RL RA y SPO+ RA según el tamaño de la red para SPP.

		Modelos			
		CVaR 0.05	CVaR 0.1	Mean Var 0.0005	Mean Var 0.001
Grilla	$3 \cdot 3$	0.015	0.015	~ 0	~ 0
	$5 \cdot 5$	0.276	0.381	0.009	0.002
	$7 \cdot 7$	0.447	0.726	0.004	0.003

Tabla 4.2: Tiempos computacionales [s] de CVaR ε Y Mean Var α según el tamaño de la red para SPP.

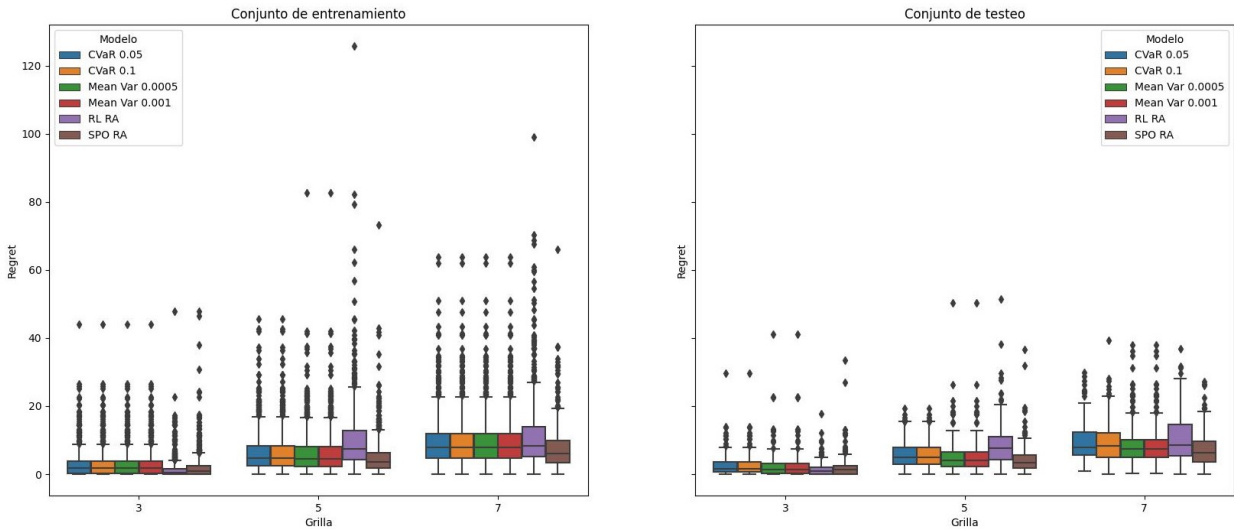


Figura 4.1: Distribución del arrepentimiento en función del caso base modificando el tamaño de la red para SPP.

De la Figura 4.1 se puede observar que:

- A medida que aumenta el tamaño de la red, se observa un incremento en el arrepentimiento, consistentes con el crecimiento de la grilla, y una mayor asimetría positiva en las cajas.
- La mediana e IQR casi no se ven afectados respecto de ambos conjuntos. RL es el único que muestra resultados variables.
- Al analizar los resultados del conjunto de pruebas, se observa que: SPO+ RA exhibe los mejores resultados y RL RA los peores; Mean Var α y CVaR ε mantienen una relativa estabilidad, incluso, Mean Var α ve reducida sus IQR con respecto al entrenamiento.
- Orden de modelos con menor arrepentimiento según sus predicciones:
 - **Grilla 3 · 3:** RL RA > SPO+ RA > Mean Var α > CVaR ε
 - **Grilla 5 · 5:** SPO+ RA > Mean Var α > CVaR ε > RL RA
 - **Grilla 7 · 7:** SPO+ RA > Mean Var α > CVaR 0.01 > CVaR 0.05 > RL RA

2^{DA} CONFIGURACIÓN: Grilla = $5 \cdot 5$, Esc = $\{200, 1000, 5000\}$, $k = 10$, $\text{deg} = 8$.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Esc	200	0.007	~ 0	0.007	0.906	~ 0	0.906
	1000	0.044	~ 0	0.044	16.5	~ 0	16.5
	5000	0.118	~ 0	0.118	111.425	~ 0	111.425

Tabla 4.3: Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de escenarios para SPP.

		Modelos			
		CVaR 0.05	CVaR 0.1	Mean Var 0.0005	Mean Var 0.001
Esc	200	0.022	0.021	~ 0	0.001
	1000	0.276	0.381	0.009	0.002
	5000	1.162	2.033	~ 0	0.001

Tabla 4.4: Tiempos computacionales [s] de CVaR ε Y Mean Var α según la cantidad de escenarios para SPP.

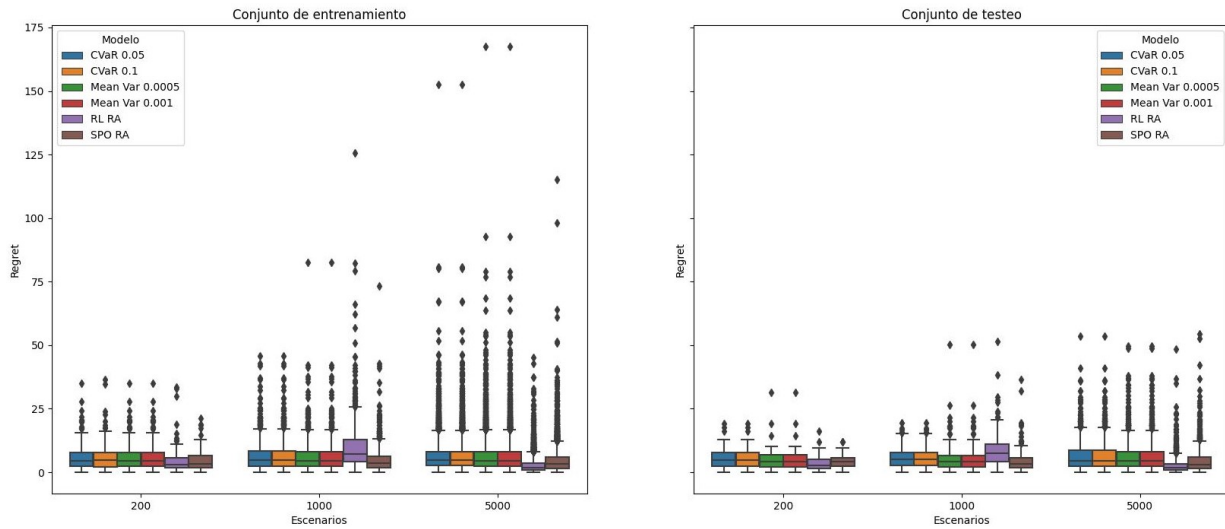


Figura 4.2: Distribución del arrepentimiento en función del caso base modificando la cantidad de escenarios para SPP.

De la Figura 4.2 se puede observar que:

- Con el incremento en el número de escenarios, se observa un leve aumento del arrepentimiento y una mayor estabilidad en la mediana para ambos conjuntos.
- Al examinar el conjunto de pruebas, se concluye que el modelo SPO RA presenta los mejores resultados. Este modelo mantiene constantes sus rangos intercuartílicos independientemente del aumento del tamaño del conjunto de datos. Los modelos Mean Var α y CVaR ϵ muestran consistencia, aunque este último presenta los resultados más desventajosos. Por otro lado, el modelo RL RA exhibe una tendencia variable, lo cual no es deseable.
- Orden de modelos con menor arrepentimiento según sus predicciones:
 - **200 escenarios:** RL RA > SPO+ RA > Mean Var α > CVaR ϵ
 - **1000 escenarios:** SPO+ RA > Mean Var α > CVaR ϵ > RL RA
 - **5000 escenarios:** RL RA > SPO+ RA > Mean Var α > CVaR ϵ

3^{RA} CONFIGURACIÓN: Grilla = $5 \cdot 5$, Esc = 1000, $k = 10$, Deg = $\{1, 2, 4, 6, 8\}$.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Deg	1	0.026	~ 0	0.026	12.228	~ 0	12.228
	2	0.039	~ 0	0.039	9.08	~ 0	9.08
	4	0.035	~ 0	0.035	8.115	~ 0	8.115
	6	0.027	~ 0	0.027	9.963	~ 0	9.963
	8	0.044	~ 0	0.044	16.5	~ 0	16.5

Tabla 4.5: Tiempos computacionales [s] de RL RA y SPO+ RA según el grado polinómico para SPP.

		Modelos			
		CVaR 0.05	CVaR 0.1	Mean Var 0.0005	Mean Var 0.001
Deg	1	0.178	0.192	0.001	~ 0
	2	0.139	0.161	~ 0	~ 0
	4	0.148	0.151	~ 0	0.001
	6	0.143	0.164	~ 0	~ 0
	8	0.276	0.381	0.009	0.002

Tabla 4.6: Tiempos computacionales [s] de CVaR ε Y Mean Var α según el grado polinómico para SPP.

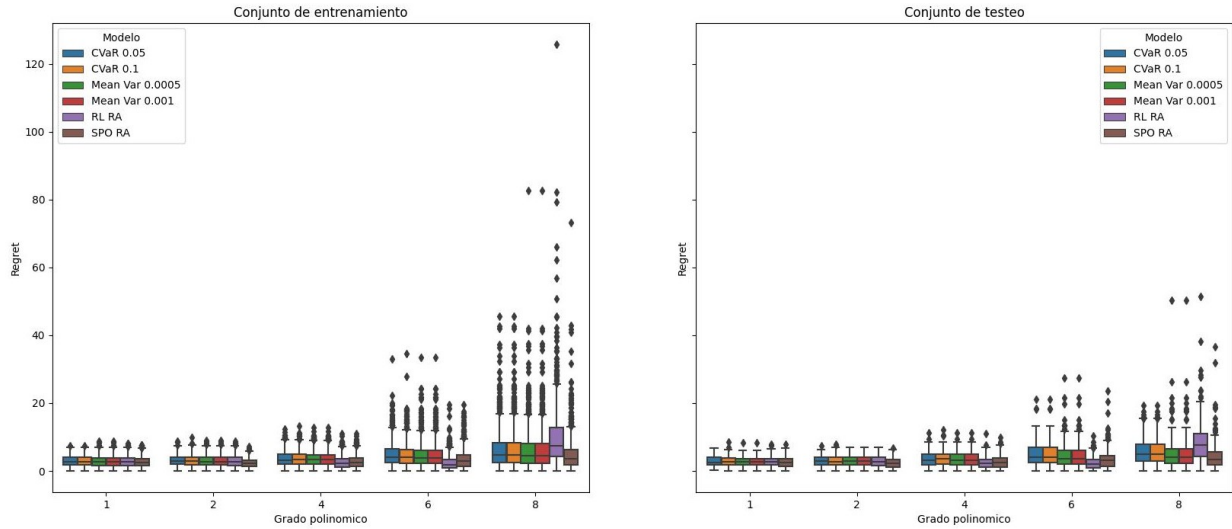


Figura 4.3: Distribución del arrepentimiento en función del caso base modificando el grado polinómico para SPP.

De la Figura 4.3 se puede observar que:

- A mayor grado polinómico, hay más arrepentimiento, ya que es más difícil para el modelo ajustar los datos.
- Al contrastar el conjunto de entrenamiento con el de testeo, se puede apreciar que la mayoría de los modelos fueron consistentes, manteniéndose dentro de los límites intercuartílicos y conservando la mediana. RL RA fue el único modelo que destacó por un cambio en su tendencia. Inicialmente, su arrepentimiento disminuía con el aumento del ruido; sin embargo, esta tendencia cambió, presentando un aumento exponencial en el grado 8.
- En el conjunto de testeo, se puede notar claramente que SPO+ RA se destaca por mantener un rendimiento consistentemente favorable, independientemente del grado polinómico empleado. En contraste, RL RA exhibe una tendencia variable en comparación con los otros modelos. Mean Var α y CVaR ε muestran una tendencia constante; sin embargo, este último presenta los resultados más desfavorables del grupo.

- Orden de modelos con menor arrepentimiento según sus predicciones:
 - **Grado 1:** SPO+ RA > Mean Var α > RL RA > CVaR 0.1 > CVaR 0.05
 - **Grado 2:** SPO+ RA > RL RA > Mean Var α > CVaR 0.1 > CVaR 0.05
 - **Grado 4:** RL RA > SPO+ RA > Mean Var α = CVaR 0.05 > CVaR 0.1
 - **Grado 6:** RL RA > SPO+ RA > Mean Var α > CVaR ε
 - **Grado 8:** SPO+ RA > Mean Var α > CVaR ε > RL RA

4^{TA} CONFIGURACIÓN: Grilla = $5 \cdot 5$, Esc = 1000, $k = \{5, 10, 20\}$, Deg = 8.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
k	5	0.014	~ 0	0.014	4.412	~ 0	4.412
	10	0.044	~ 0	0.044	16.5	~ 0	16.5
	20	0.037	~ 0	0.037	23.087	~ 0	23.087

Tabla 4.7: Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de atributos para SPP.

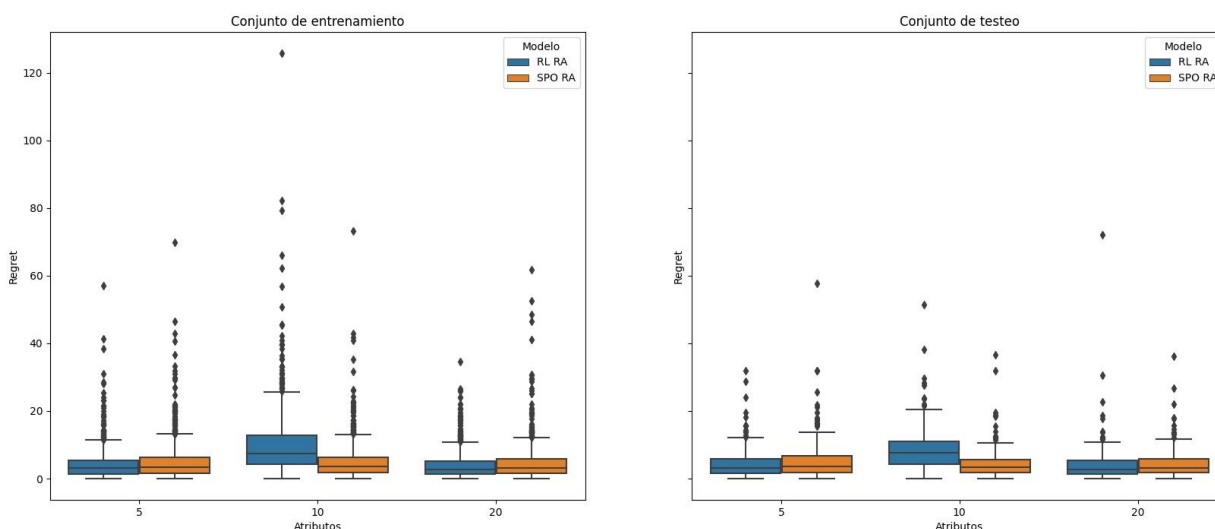


Figura 4.4: Distribución del arrepentimiento en función del caso base modificando en el número de atributos para SPP.

De la Figura 4.4 se puede observar que:

- Contrastando el conjunto de testeo con el de entrenamiento, se puede apreciar que SPO RA logra mantener casi estables sus resultados, pues no se aprecia variación significativa de los límites intercuartílicos, mediana y valor máximo. Sin embargo, RL RA presenta más variación en su IQR para un vector de tamaño 10, no pudiendo mostrar una tendencia clara con el aumento de información.
- Orden de modelos con menor arrepentimiento según sus predicciones:
 - 5 atributos: $RL\ RA > SPO+ RA$
 - 10 atributos: $SPO+ RA > RL\ RA$
 - 20 atributos: $RL\ RA > SPO+ RA$

	Puntajes Obtenidos																	
	Grid				Esc				Deg						k			
	3	5	7	Total	200	1000	5000	Total	1	2	4	6	8	Total	5	10	20	Total
RL RA	17	6	5	28	16	6	18	40	12	9	15	8	6	50	16	9	15	40
SPO RA	13	15	16	44	15	15	12	42	10	15	15	15	15	55	11	17	14	42
CVaR 0.05	8	12	9	29	8	12	7	27	14	10	8	9	12	53	-	-	-	-
CVaR 0.1	8	12	5	25	8	12	7	27	10	7	7	9	12	45	-	-	-	-
Mean Var 0.0005	9	13	12	34	11	13	9	33	12	9	8	11	13	53	-	-	-	-
Mean Var 0.001	9	13	12	34	11	13	9	33	12	9	8	11	13	53	-	-	-	-

Tabla 4.8: Puntuaciones según evaluación de gráficos según parámetro variable para SPP.

La Tabla 4.8 se construye a partir de la evaluación de los datos estadísticos de cada configuración detallada en el Apéndice A. En esta tabla, se muestran los puntajes totales, donde un puntaje más alto sugiere un mejor resultado. Con el propósito de mejorar la visualización, los resultados se presentan en un gradiente de color, en el cual el azul representa el mejor resultado y el rojo, el peor.

Al analizar los tiempos de ejecución, lo que corresponde a las Tablas 4.1, 4.2, 4.3, 4.4, 4.5, 4.6, y 4.7, se destacan los siguientes puntos:

- Tanto para configuraciones pequeñas como para las de gran tamaño, se observa que RL RA, SPO+ RA y Mean VaR α presentan tiempos de resolución casi nulos, cercanos a 0 segundos.
- Los modelos Mean VaR α con parámetros 0.0005 y 0.001 muestran una eficiencia notable, ya que, al predecir igual e incluso mejor que los modelos CVaR ε en la mayoría de los casos, logran hacerlo en cuestión de milisegundos, a diferencia de los casi segundos que requiere CVaR ε .
- En todas las configuraciones, los modelos CVaR ε exhiben los tiempos de ejecución más extensos.

4.1.1. Medidas de riesgos + contexto

En primer lugar, se presentan los resultados gráficos utilizando el Índice de Davies-Bouldin (DBI) para analizar el valor óptimo de clúster para el problema de la ruta más corta (SPP).

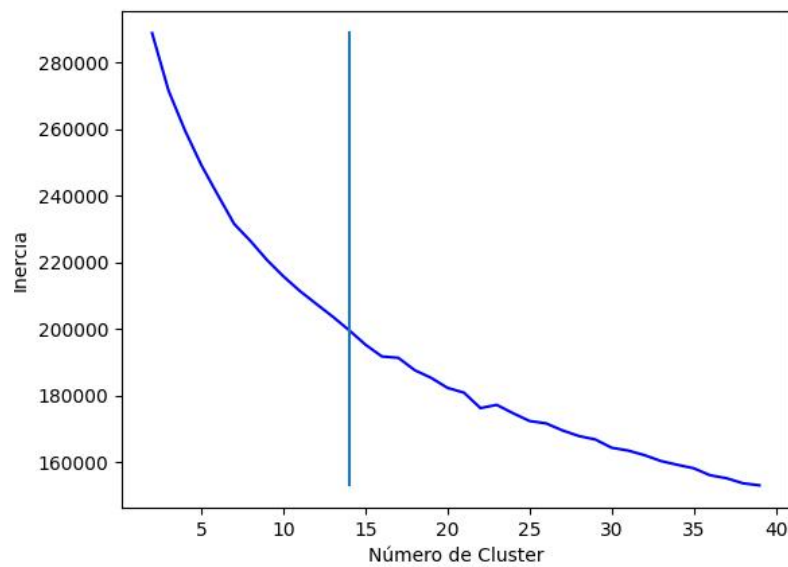


Figura 4.5: Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información global para SPP

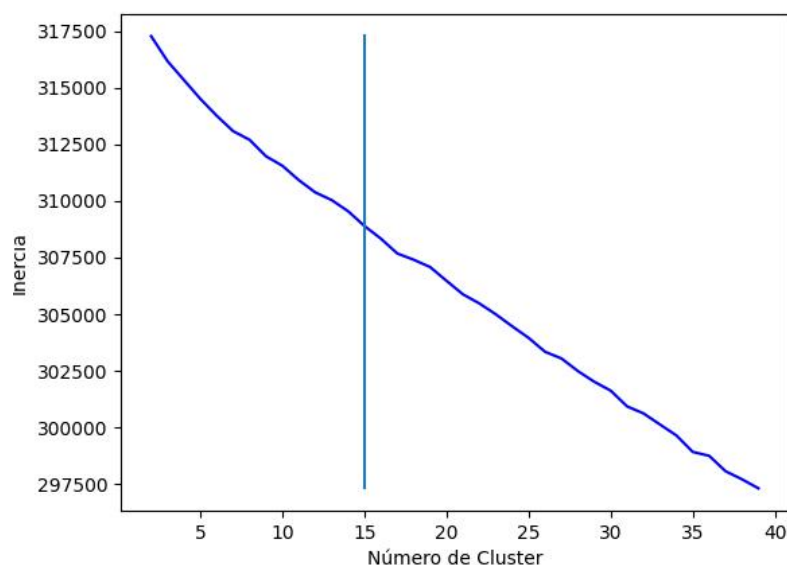


Figura 4.6: Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información local para SPP

Es importante señalar que de la Figura 4.6 se puede inferir que la tendencia de los resultados al mostrarse en una línea en lugar de una curva, para denotar el codo, puede deberse a la variabilidad de los datos al ser generados de forma sintética y sin repitencia, como si sucede en la Figura 4.5. Es por lo anterior que se obtiene el valor máximo del rango permitido del número de clusters, que es 15, pues en otro caso, perfectamente se podrían obtener tantos grupos como datos a clasificar se tengan. Por otra parte, se puede ver que para un caso donde se cuentan con datos repetidos, se visualiza un codo con un resultado $Q=14$.

A continuación, se presentan gráficos que ilustran la distribución del arrepentimiento de todos los modelos desarrollados en esta tesis: SPO, RL, Mean Var α y CVaR ε . Esto incluye sus versiones simples de medidas de riesgo, así como las variantes con clusterización y reducción de dimensionalidad del vector de características.

En tres de los seis gráficos subsiguientes, los cuales solo contemplan información parcial, se recuperó información previamente vista en la Figura 4.1. Los gráficos están divididos según la información contextual, ya sea global (RA) o local (NRA). Además, se han mantenido los parámetros del caso base, aunque se ha modificado el tamaño de la grilla.

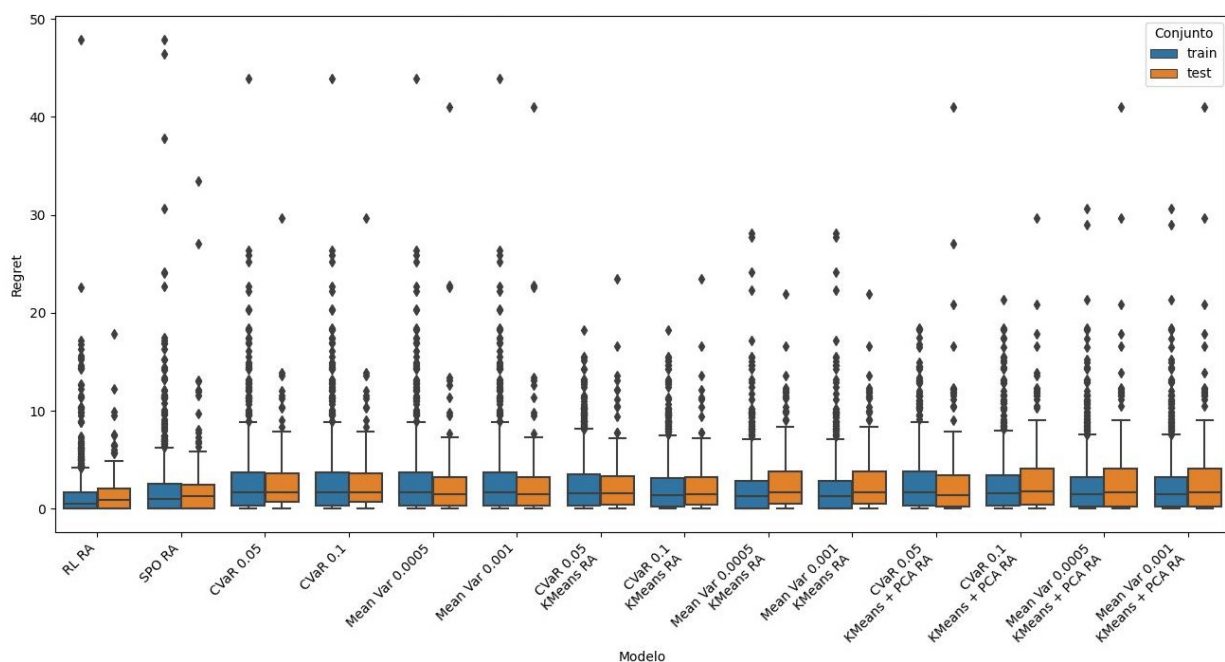


Figura 4.7: Distribución de arrepentimiento para una red $3 \cdot 3$ con información global en SPP.

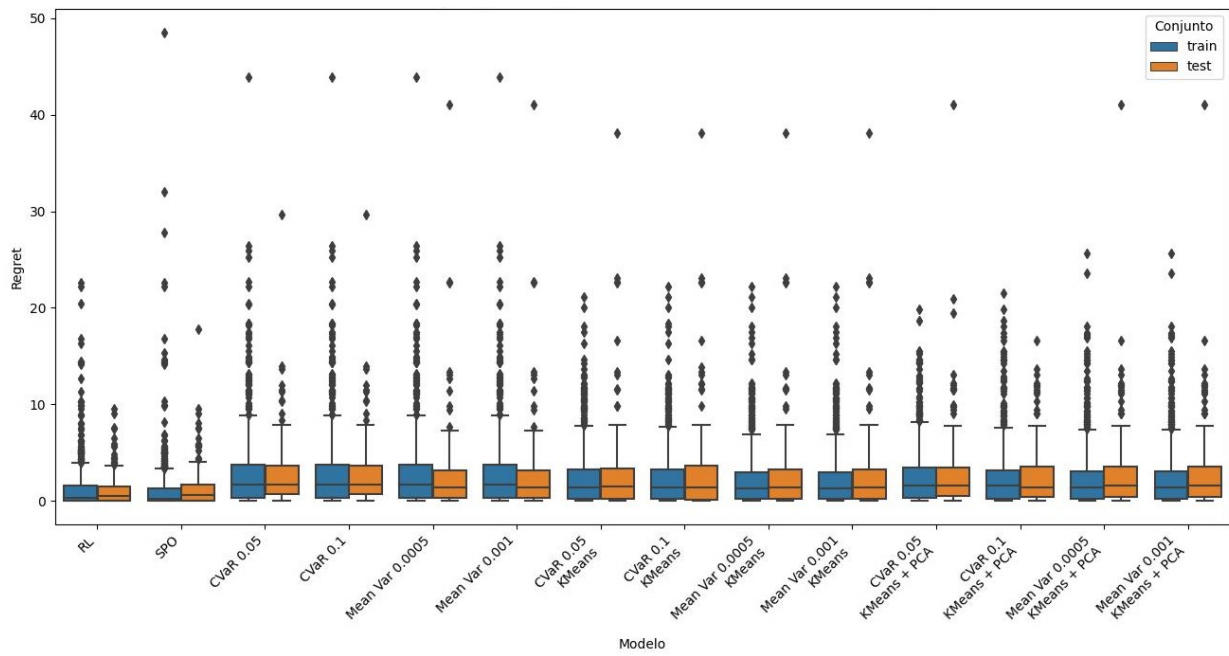


Figura 4.8: Distribución de arrepentimiento para una red $3 \cdot 3$ con información local en SPP.

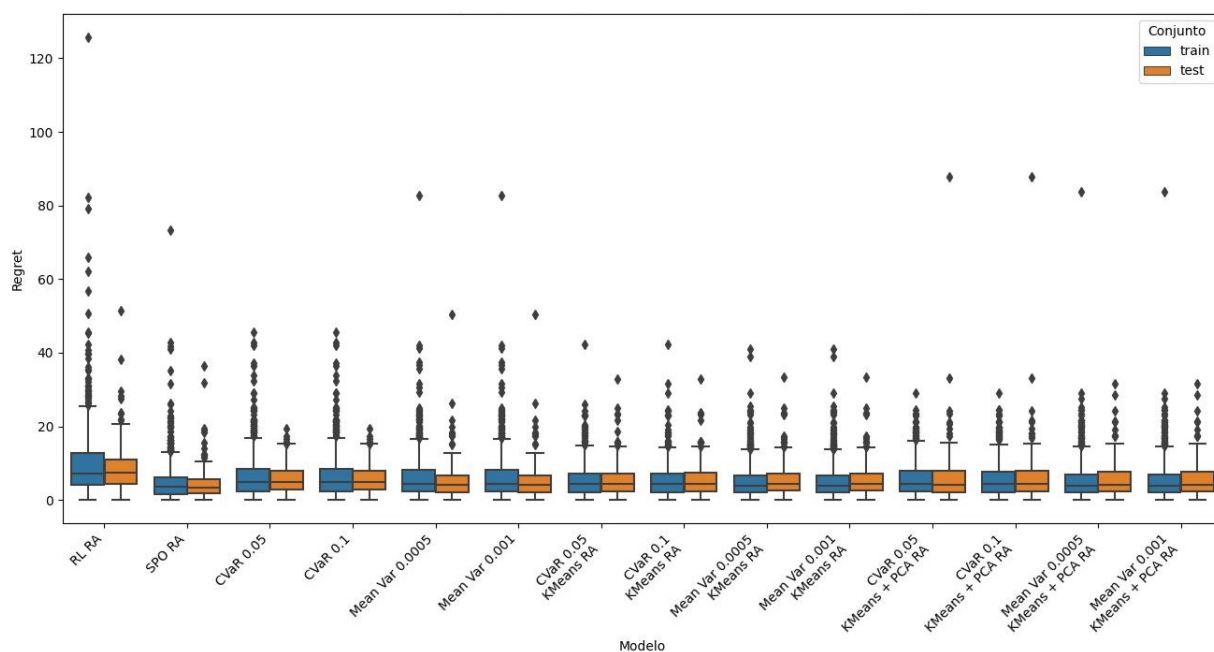


Figura 4.9: Distribución de arrepentimiento para una red $5 \cdot 5$ con información global en SPP.

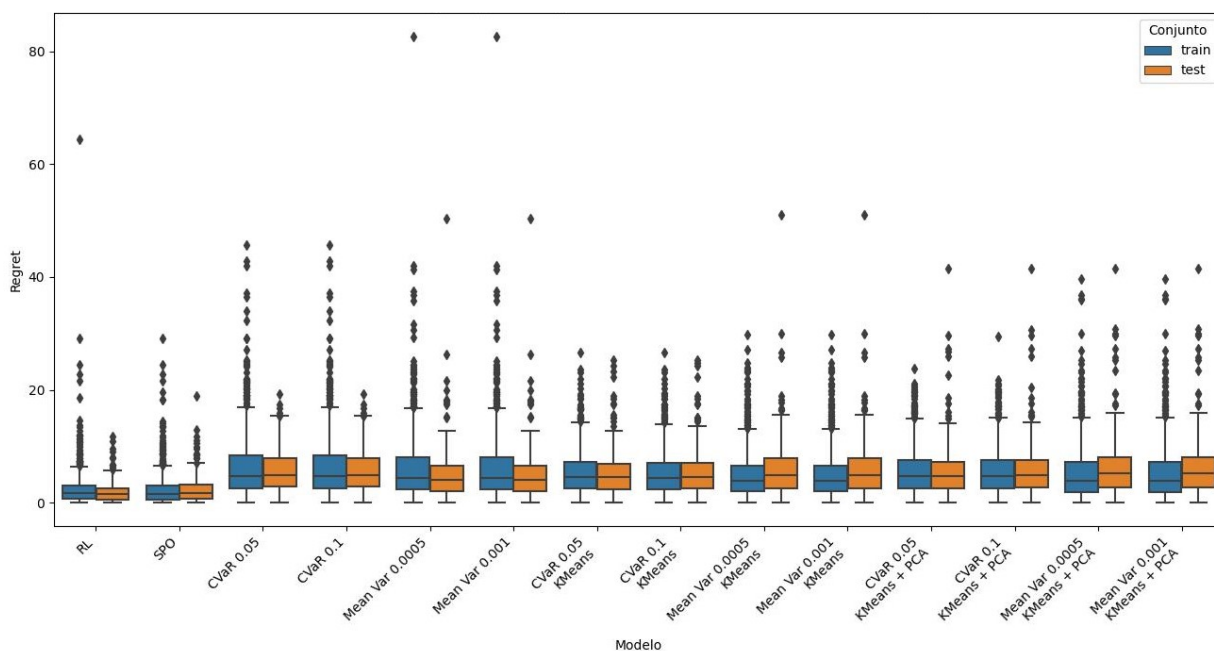


Figura 4.10: Distribución de arrepentimiento para una red $5 \cdot 5$ con información local en SPP.

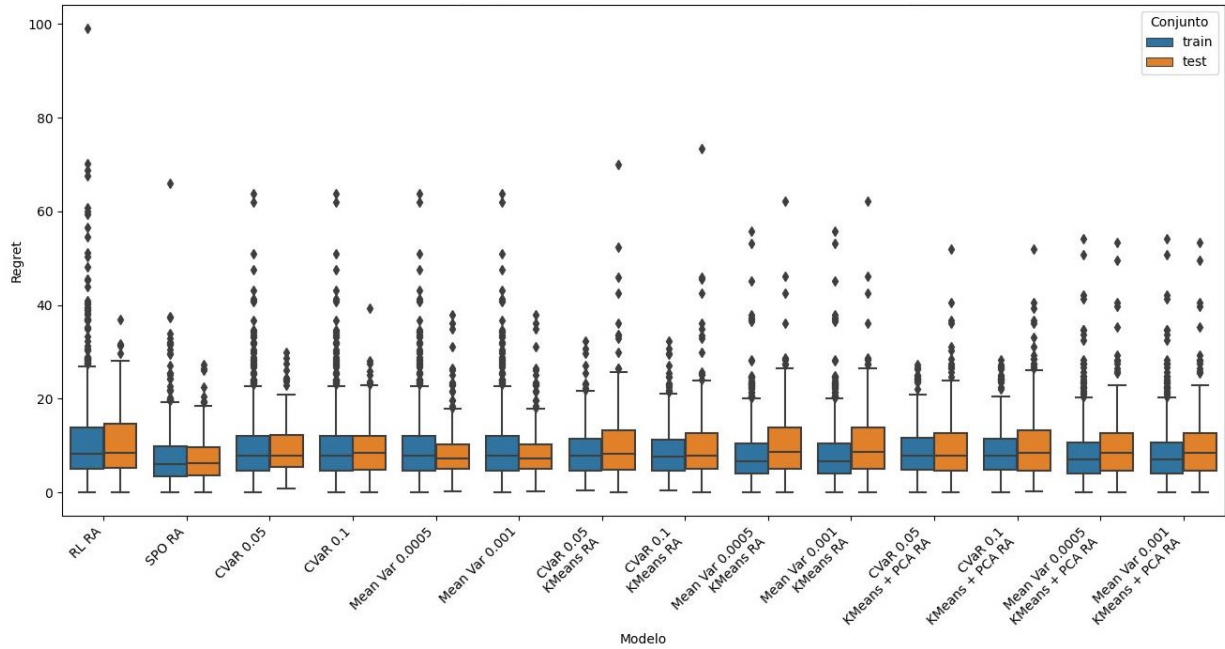


Figura 4.11: Distribución de arrepentimiento para una red 7 · 7 con información global en SPP.

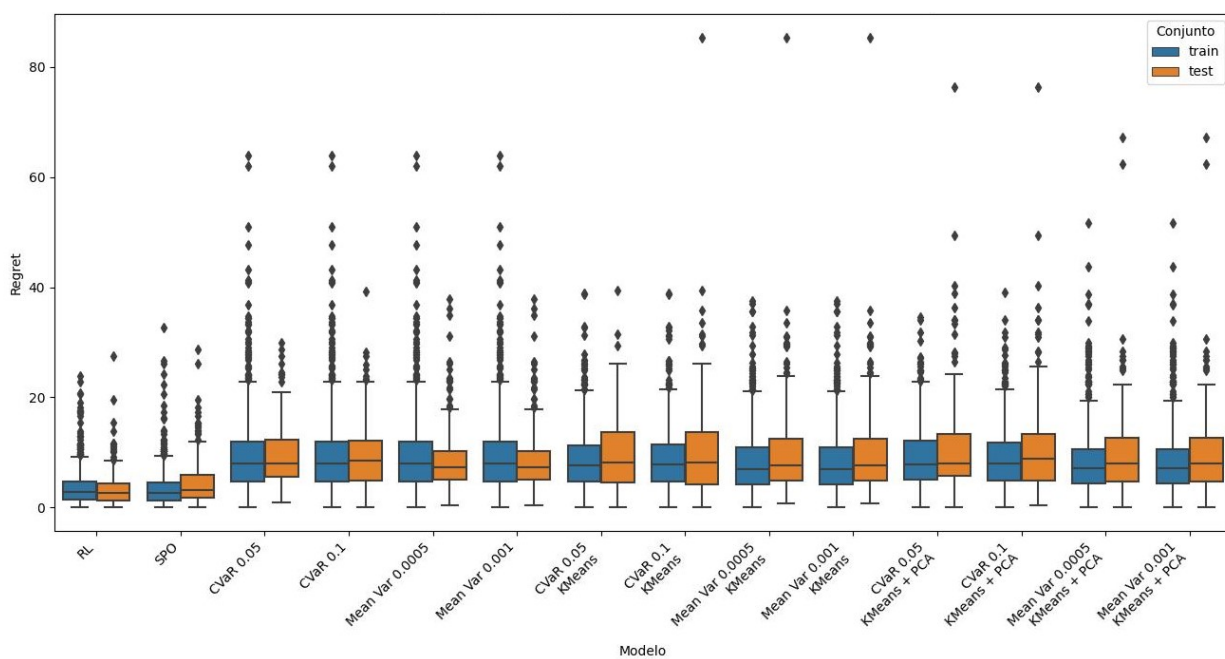


Figura 4.12: Distribución de arrepentimiento para una red $7 \cdot 7$ con información local en SPP.

		Modelos					
		RL RA			RL		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Grilla	3 · 3	0.008	~ 0	0.008	0.007	~ 0	0.007
	5 · 5	0.044	~ 0	0.044	0.068	~ 0	0.068
	7 · 7	0.046	~ 0	0.046	0.062	~ 0	0.062

Tabla 4.9: Tiempos computacionales [s] de RL RA vs. RL según tamaño de red para SPP.

		Modelos					
		SPO+ RA			SPO+		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Grilla	3 · 3	1.666	~ 0	1.666	1.048	~ 0	1.048
	5 · 5	16.5	~ 0	16.5	15.303	~ 0	15.303
	7 · 7	23.178	~ 0	23.178	35.177	~ 0	35.177

Tabla 4.10: Tiempos computacionales [s] de SPO+ RA vs. SPO+ según tamaño de red para SPP.

Al comparar las Figuras RA (4.7, 4.9, y 4.11) con las Figuras NRA (4.8, 4.10, y 4.12), se observa que las medidas de riesgo contextualizadas predicen de manera similar a los modelos originales ¹ de cada grupo, aunque con variaciones mínimas en el rango intercuartil y una mayor adición de varianza (asimetría positiva). Siendo las medidas con K-Means los únicos modelos que muestra una ínfima reducción en el arrepentimiento en graficos de 3 · 3 y 5 · 5. Para los graficos 7 · 7, aplicar K-Means y PCA no hubo impacto positivo.

En las figuras NRA, SPO muestra no predecir bien cuando el tamaño de grilla es muy grande, esto en comparación con RL.

Finalmente, al comparar los tiempos promedio de ejecución en el entrenamiento de los modelos RL RA y SPO+ RA con sus contrapartes del grupo NRA, se observa una notable diferencia en cuanto a su menor duración, a pesar de que la matriz de atributos tenga el mismo tamaño en ambos casos.

Por tanto, al evaluar los resultados de estos modelos, incluyendo las estimaciones de arrepentimiento y los tiempos de ejecución entre un grupo y otro, surge la pregunta de si vale la pena considerar toda la información disponible. La respuesta es; no. No se aprecia

¹Llámesese modelos originales a aquellos modelos aversos al riesgo que no consideran el contexto, es decir, los modelos CVaR 0.05, CVaR 0.1, Mean Var 0.05 y Mean Var 0.1

una ventaja significativa al considerar vectores de características distintos entre sí o, dicho de otra manera, al costear información local de cada arista.

4.1.2. Conclusiones para SPP

SPO+ RA demuestra ser el modelo más efectivo para manejar la incertidumbre en diversas configuraciones, mostrando un desempeño superior en situaciones con tamaños de grilla más grandes y en variados grados polinómicos. Este modelo mantiene un menor nivel de arrepentimiento y demuestra una mayor robustez frente a la incertidumbre creciente.

Por otro lado, RL rendimiento es menos consistente, especialmente en grados polinómicos altos y en grandes tamaños de grilla. En consecuencia, SPO+ RA se destaca como el modelo más capaz de enfrentar la incertidumbre en la mayoría de los escenarios evaluados, siendo la elección preferible para configuraciones con incertidumbres moderadas a altas en tamaño de grilla, número de escenarios y complejidad de los modelos.

Por tanto:

- SPO responde mejor a la incertidumbre y ofrece los mejores resultados en términos de arrepentimiento. Es ideal para escenarios donde la precisión es crucial, aunque con un mayor costo computacional.
- Mean Var α muestra estabilidad y buen desempeño en la mayoría de las configuraciones, siendo una opción robusta cuando se necesita un balance entre desempeño y costo computacional.
- RL es adecuado para escenarios con limitaciones de tiempo y recursos computacionales, aunque con un desempeño menos consistente.
- CVaR ϵ Aunque no es el mejor en términos de arrepentimiento, mantiene una estabilidad que puede ser útil en ciertos contextos.

4.2. Resultados computacionales sobre el problema de la mochila

Por otra parte, la ejecución de los experimentos KP01 se realizaron en un sistema computacional distinto al del camino más corto, dado que no fue posible ejecutar dichos experimentos en aquel computador dado los requerimientos de CPU, GPU y memoria RAM que se necesitaban. Este segundo equipo contó con un sistema operativo Windows 10 Pro, un procesador AMD Ryzen 5 1400 3.20 GHz Quad-Core, NVIDIA GeForce GTX 1050 Ti, 16GB RAM y 1TB SSD.

El caso base para este estudio será definido con los siguientes parámetros: $esc=750$, $cap=120$, $k=5$ y $deg=8$. Estos parámetros serán modificados sistemáticamente para analizar cómo responden los modelos a dichas variaciones.

1^{RA} CONFIGURACIÓN: $Cap = \{60, 120, 180\}$, $Esc = 750$, $k = 5$, $Deg = 8$.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Cap	60	0.027	0.001	0.028	74.356	0.001	74.357
	120	0.023	0.001	0.024	142.130	0.001	142.131
	180	0.024	0.001	0.025	144.490	0.001	144.491

Tabla 4.11: Tiempos computacionales [s] de RL RA y SPO+ RA según la capacidad para KP01.

		Modelos			
		CVaR 0.05	CVaR 0.1	Mean Var 0.0005	Mean Var 0.001
Cap	60	0.270	0.445	0.001	0.001
	120	0.123	0.137	0.001	~ 0
	180	0.241	0.243	0.001	0.001

Tabla 4.12: Tiempos computacionales [s] de CVaR Y Mean Var según la capacidad para KP01.

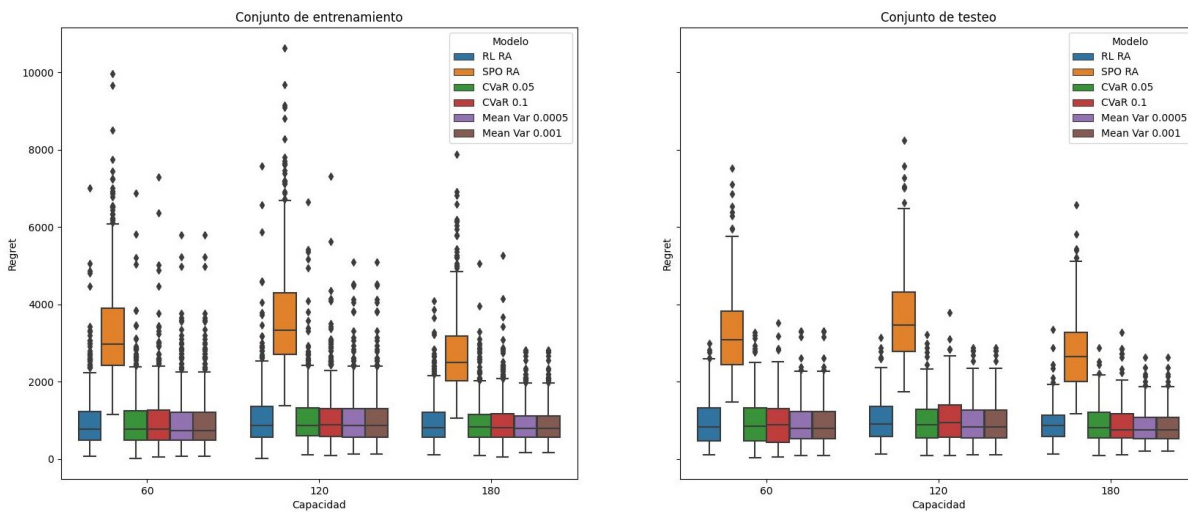


Figura 4.13: Distribución del arrepentimiento en función del caso base modificando la capacidad de la mochila para KP01.

Considerando la Figura 4.13, se pueden destacar los siguientes puntos:

- A una mayor capacidad de la mochila se observa una ínfima reducción de regret y estabilidad constante en los límites intercuartiles.
- SPO+ RA destaca por su comportamiento inusual, mostrando consistentemente el mayor nivel de arrepentimiento entre los modelos analizados. Además, este comportamiento tiende a empeorar en comparación con los otros modelos.
- Todos los modelos muestran una predicción prácticamente idéntica con respecto al conjunto de entrenamiento.

- Al evaluar los resultados en el conjunto de testeo, se deduce que Mean Var α muestra un desempeño superior en comparación con los otros modelos a medida que aumenta la capacidad de la mochila; CVaR ϵ y RL RA muestran una tendencia variable, aunque sugieren una posible mejora con el aumento de la capacidad.
- Orden de modelos con menor arrepentimiento según sus predicciones:
 - **Cap 60:** Mean Var α > CVaR 0.1 > CVaR 0.05 > RL RA $\ggg$ SPO+ RA
 - **Cap 120:** Mean Var α > CVaR 0.05 > RL RA > CVaR 0.1 $\ggg$ SPO+ RA
 - **Cap 180:** Mean Var α > CVaR 0.1 > CVaR 0.05 > RL RA $\gg$ SPO+ RA

2^{DA} CONFIGURACIÓN: Cap = 120, Esc = {200, 750, 1500}, k = 5, Deg = 8.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Esc	200	0.006	0.001	0.007	10.636	0.001	10.637
	750	0.023	0.001	0.024	142.130	0.001	142.131
	1500	0.049	0.001	0.050	715.073	0.001	715.074

Tabla 4.13: Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de escenarios para KP01.

		Modelos			
		CVaR 0.05	CVaR 0.1	Mean Var 0.0005	Mean Var 0.001
Esc	200	0.058	0.058	0.001	0.001
	750	0.123	0.137	0.001	~ 0
	1500	0.443	0.484	0.001	0.001

Tabla 4.14: Tiempos computacionales [s] de CVaR Y Mean Var según la cantidad de escenarios para KP01.

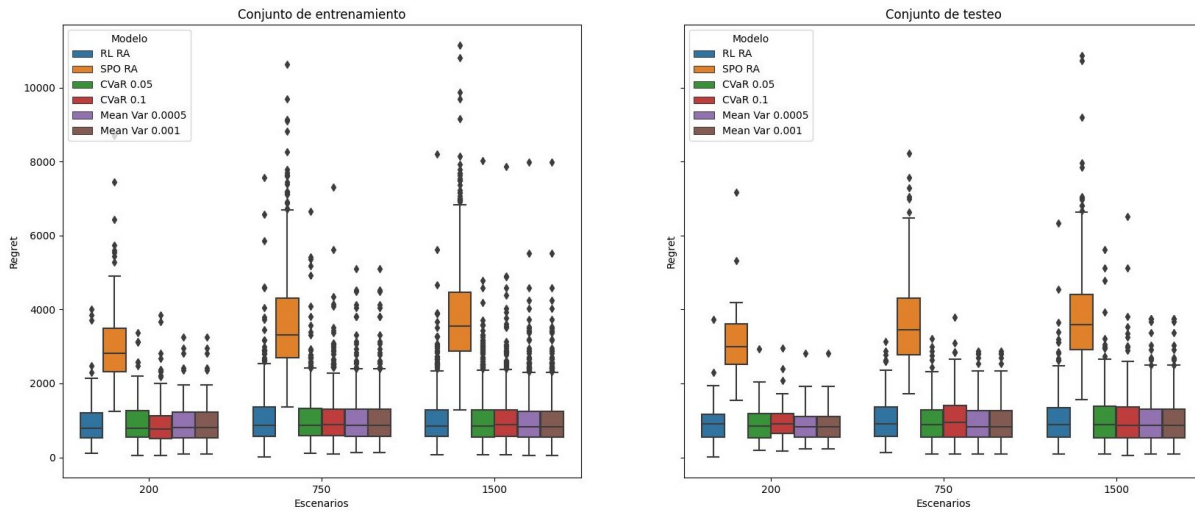


Figura 4.14: Distribución del arrepentimiento en función del caso base modificando la cantidad de escenarios para KP01.

Al considerar los resultados de la Figura 4.14, se puede identificar lo siguiente:

- En general, se observa que a medida que aumenta la cantidad de escenarios, hay un ligero aumento en la incertidumbre, una ligera tendencia hacia la asimetría positiva en las cajas y una mayor dispersión.
- Al comparar los conjuntos de testeo y entrenamiento, se observa que los límites intercuartílicos en las predicciones se mantienen constantes y que la mediana mejora ligeramente en la mayoría de los modelos a medida que aumenta la cantidad de escenarios considerados. A pesar de que SPO+ RA exhibe el mayor arrepentimiento y muestra un comportamiento atípico en comparación con los demás modelos, es el modelo más cuasiidéntico en términos generales.
- Al analizar el conjunto de testeo, se concluye que Mean Var α muestra una tendencia consistente de mejora a medida que aumenta el número de escenarios, siendo el modelo con los mejores resultados. Tanto RL RA como CVaR 0.1 parecieran mostrar un comportamiento variable, ya que empeoran en el segundo conjunto y en el tercero

mejoran. Por otro lado, aunque CVaR 0.05 no exhibe los mejores resultados, se mantiene consistente con el aumento del conjunto de datos.

- Orden de modelos con menor arrepentimiento según sus predicciones:
 - **200 escenarios:** Mean Var α > CVaR 0.05 > RL RA > CVaR 0.1 $\ggg$ SPO+ RA
 - **750 escenarios:** CVaR 0.05 > Mean Var α > RL RA > CVaR 0.1 $\ggg$ SPO+ RA
 - **1500 escenarios:** Mean Var α > RL RA > CVaR 0.1 > CVaR 0.05 $\ggg$ SPO+ RA

3^{RA} CONFIGURACIÓN: Cap = 120, Esc = 750, k = 5, Deg = {1, 2, 4, 6, 8}.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Deg	1	0.026	0.001	0.027	76.850	0.001	76.851
	2	0.024	0.001	0.025	47.393	~ 0	47.393
	4	0.026	0.001	0.027	136.006	0.001	136.007
	6	0.024	0.001	0.025	108.835	0.001	108.836
	8	0.023	0.001	0.024	142.132	0.001	142.133

Tabla 4.15: Tiempos computacionales [s] de RL RA y SPO+ RA según el grado polinómico para KP01.

		Modelos			
		CVaR 0.05	CVaR 0.1	Mean Var 0.0005	Mean Var 0.001
Deg	1	0.148	0.091	0.001	0.001
	2	0.126	0.089	0.001	0.001
	4	0.145	0.248	0.001	0.001
	6	0.143	0.250	0.001	0.001
	8	0.123	0.137	0.001	~ 0

Tabla 4.16: Tiempos computacionales [s] de CVaR Y Mean Var según el grado polinómico para KP01.

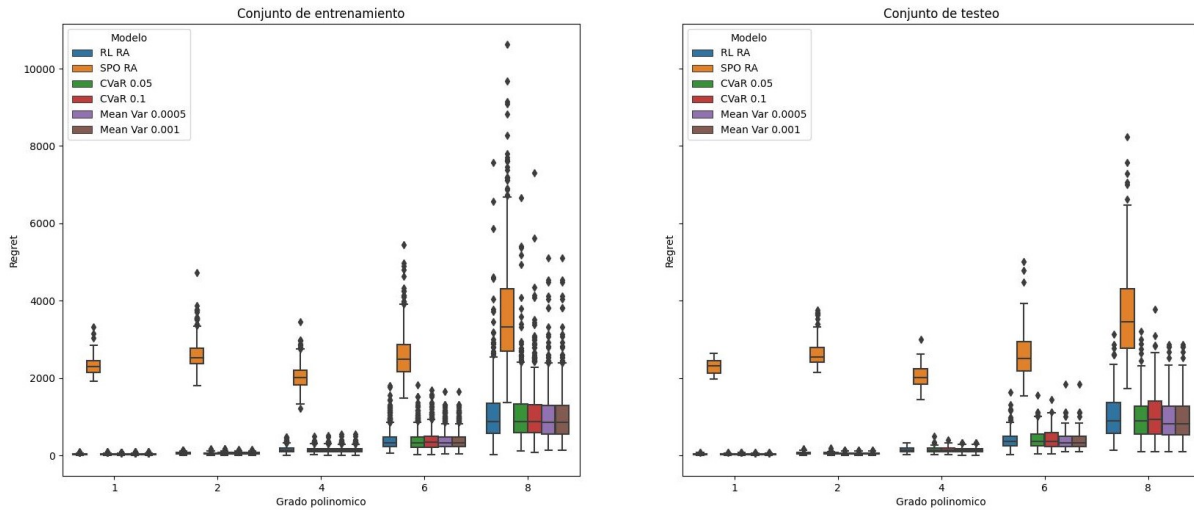


Figura 4.15: Distribución del arrepentimiento en función del caso base modificando el grado polinómico para KP01.

La Figura 4.15 muestra claramente que:

- A medida que aumenta el grado polinómico, también aumentan el arrepentimiento, la inexactitud en las predicciones, la asimetría positiva en las cajas y la dispersión.
- Respecto del conjunto de testeo con respecto al entrenamiento, los modelos que peores predicen son los modelos de CVaR ϵ .
- Nuevamente, destaca el comportamiento atípico de SPO+ RA. Sin embargo, su capacidad predictiva es notable, ya que en ambos conjuntos concentró el 50 % de los datos en rangos IQR casi idénticos y mantuvo la mediana constante para cada grado y conjunto.
- Al observar el conjunto de testeo Mean Var α se mantiene consistentemente en una posición favorable a medida que aumenta el grado polinómico, destacándolos como los modelos con mejores resultados. Por otro lado, CVaR ϵ y RL RA, aunque consistentes con el aumento de grado, exhiben los peores resultados.
- Orden de modelos con menor arrepentimiento según sus predicciones:

- **Grado 1:** Mean Var $\alpha > \text{CVaR } 0.1 > \text{RL RA} > \text{CVaR } 0.05 \ggg \text{SPO+ RA}$
- **Grado 2:** Mean Var $\alpha > \text{CVaR } 0.1 > \text{CVaR } 0.05 > \text{RL RA} \ggg \text{SPO+ RA}$
- **Grado 4:** Mean Var $\alpha > \text{CVaR } 0.05 > \text{RL RA} > \text{CVaR } 0.1 \ggg \text{SPO+ RA}$
- **Grado 6:** RL RA $> \text{Mean Var } \alpha > \text{CVaR } 0.05 > \text{CVaR } 0.1 \ggg \text{SPO+ RA}$
- **Grado 8:** Mean Var $\alpha > \text{CVaR } 0.05 > \text{RL RA} > \text{CVaR } 0.1 \ggg \text{SPO+ RA}$

4^{TA} CONFIGURACIÓN: Cap = 120, Esc = 750, $k = \{3, 5, 7\}$, Deg = 8.

		Modelos					
		RL RA			SPO+ RA		
		Train	Algoritmo	Total	Train	Algoritmo	Total
k	3	0.022	0.001	0.023	64.982	0.001	64.983
	5	0.023	0.001	0.024	142.130	0.001	142.131
	7	0.028	0.001	0.029	190.887	0.001	190.888

Tabla 4.17: Tiempos computacionales [s] de RL RA y SPO+ RA según la cantidad de atributos para KP01.

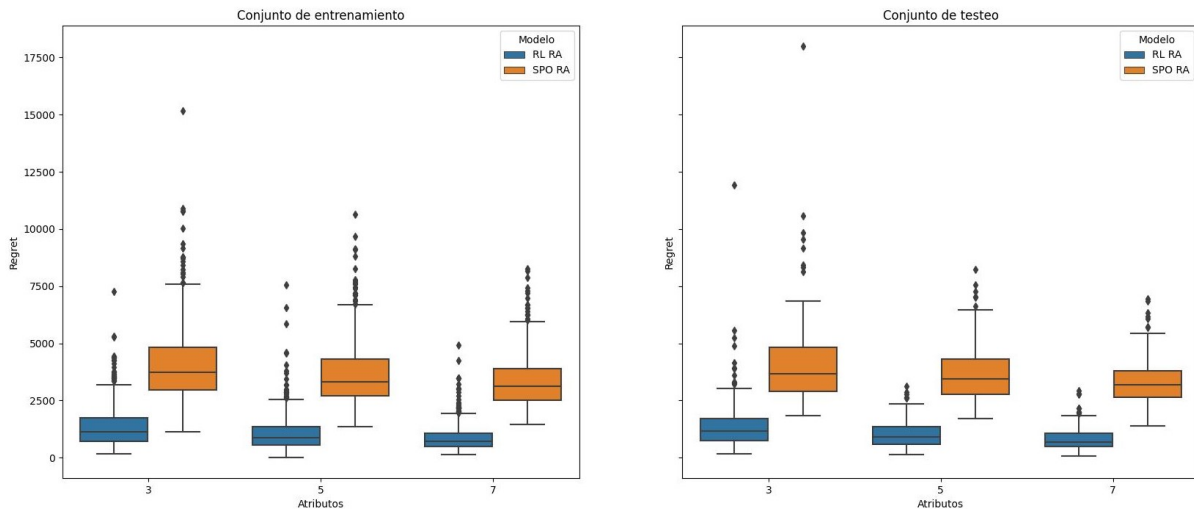


Figura 4.16: Distribución del arrepentimiento en función del caso base modificando en el número de atributos para KP01.

De acuerdo con la Figura 4.16, se pueden extraer las siguientes observaciones:

- Se observa que a medida que aumenta la cantidad de información contextual, disminuye la incertidumbre, la asimetría positiva en las cajas y la dispersión.
- Al comparar los conjuntos de testeo y entrenamiento, se nota que a medida que se considera más información, ambos modelos mantienen su precisión, conservando los límites intercuartiles y la mediana en relación con cada parámetro variable.
- Orden de modelos con menor arrepentimiento según sus predicciones:
 - 3 atributos: $RL\ RA > SPO+ RA$
 - 5 atributos: $RL\ RA > SPO+ RA$
 - 7 atributos: $RL\ RA > SPO+ RA$

	Puntajes Obtenidos																	
	Cap				Esc				Deg						k			
	60	120	180	Total	200	750	1500	Total	1	2	4	6	8	Total	3	5	7	Total
RL RA	9	10	8	27	6	10	11	27	7	7	11	11	10	46	17	17	18	52
SPO RA	2	2	2	6	2	2	2	6	6	2	2	2	2	14	9	9	9	27
CVaR 0.05	9	10	9	28	11	10	9	30	9	10	11	9	10	49	-	-	-	-
CVaR 0.1	8	9	11	28	12	9	8	29	8	9	8	9	9	43	-	-	-	-
Mean Var 0.0005	14	13	16	43	15	13	16	44	12	15	14	13	13	67	-	-	-	-
Mean Var 0.001	14	13	16	43	15	13	16	44	12	15	14	13	13	67	-	-	-	-

Tabla 4.18: Puntuaciones según evaluación de gráficos según parámetro variable para KP01.

La Tabla 4.18 tiene su origen en la evaluación aplicada a los datos estadísticos² de cada configuración que se encuentran en el Apéndice B. Al igual que en la sección anterior, en esta tabla, se muestran los puntajes totales en un gradiente de color, en el cual el azul representa el mejor resultado y el rojo, el peor.

Al analizar los tiempos de ejecución, lo que corresponde a las Tablas 4.11, 4.12, 4.13, 4.14, 4.15, 4.16, y 4.17, se destacan los siguientes puntos:

- Tanto para configuraciones pequeñas como para las de gran tamaño, se observa que RL RA, SPO+ RA y Mean VaR α presentan tiempos de resolución casi nulos, cercanos a 0 segundos.
- Los modelos de Mean VaR α demuestran una eficiencia destacable. No solo lideran los resultados en la evaluación de la medida de desempeño, el regret, sino que también son los que en su mayoría logran los tiempos más bajos dentro del grupo, realizando los cálculos en cuestión de milisegundos.
- En todas las configuraciones, los modelos CVaR ε exhiben los tiempos de ejecución más extensos.

²Los estadísticos de SPO para KP01 no se incluyeron en los cálculos de los promedios generales debido a que presentaban valores extremadamente altos. Por esta razón, se consideraron los estadísticos de este modelo como valores atípicos.

4.2.1. Medidas de riesgos + contexto

A continuación, se presentan los resultados gráficos utilizando el Índice de Davies-Bouldin (DBI) para analizar el valor cluster óptimo para KP01.

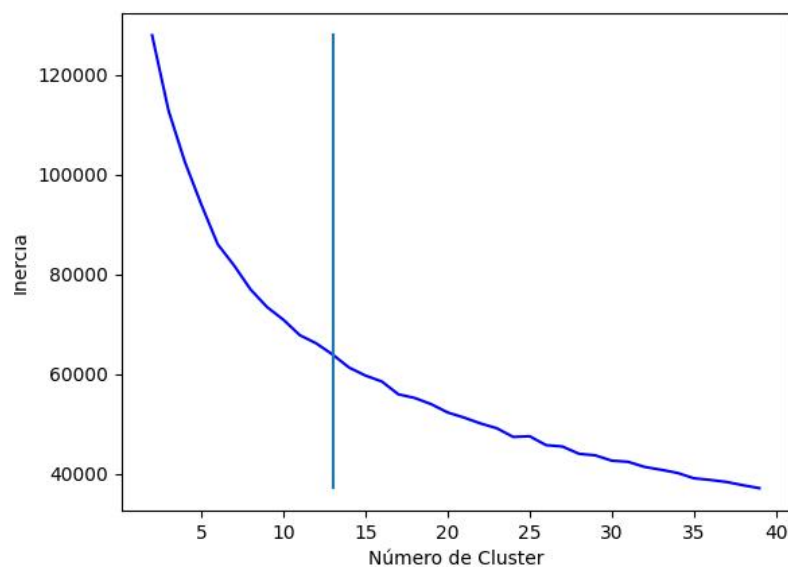


Figura 4.17: Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información global para KP01

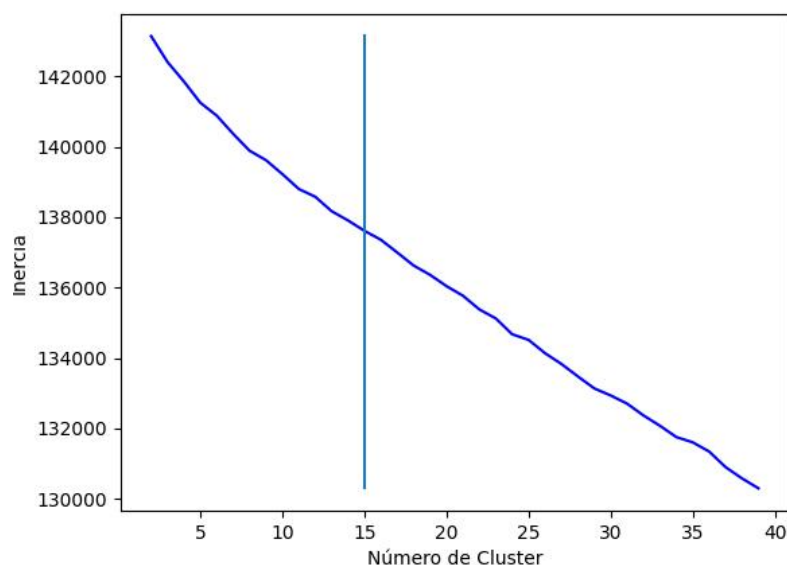


Figura 4.18: Número óptimo de grupos con índice de Davies-Bouldin para el caso base con información local para KP01

De los resultados de KP01, se observa la misma tendencia lineal o curva de tratarse de experimentos con información local o global. Sin embargo, el experimento con información global con caso base logró clasificar la información en 14 grupos.

Al igual que en el problema anterior, se presentan gráficos que ilustran la distribución del arrepentimiento de todos los modelos desarrollados en la presente tesis. En tres de los seis gráficos subsiguientes, se recuperó información previamente vista en la Figura 4.13. Sin embargo, tanto a estos como a los otros tres gráficos, se les añadieron los modelos aversos al riesgo que consideran los atributos o características x_i para su cálculo.

Los gráficos están divididos según la información contextual, ya sea global (RA) o local (NRA). Además, se han mantenido los parámetros del caso base, aunque se ha modificado la capacidad de la mochila.

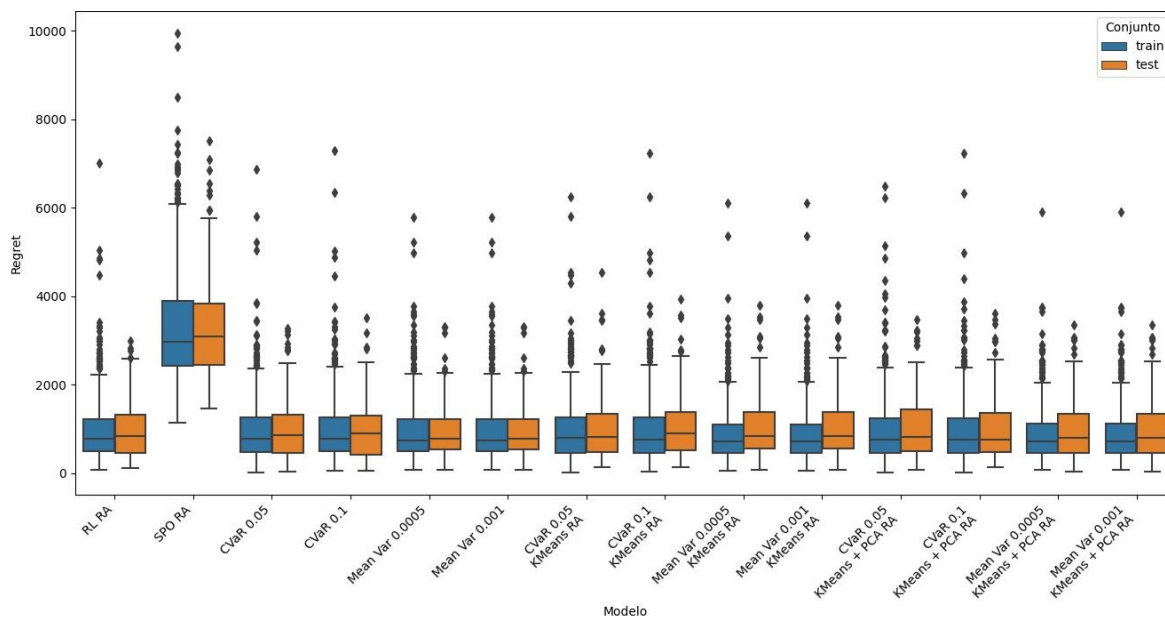


Figura 4.19: Distribución de arrepentimiento para una capacidad 60 con información global en KP01.

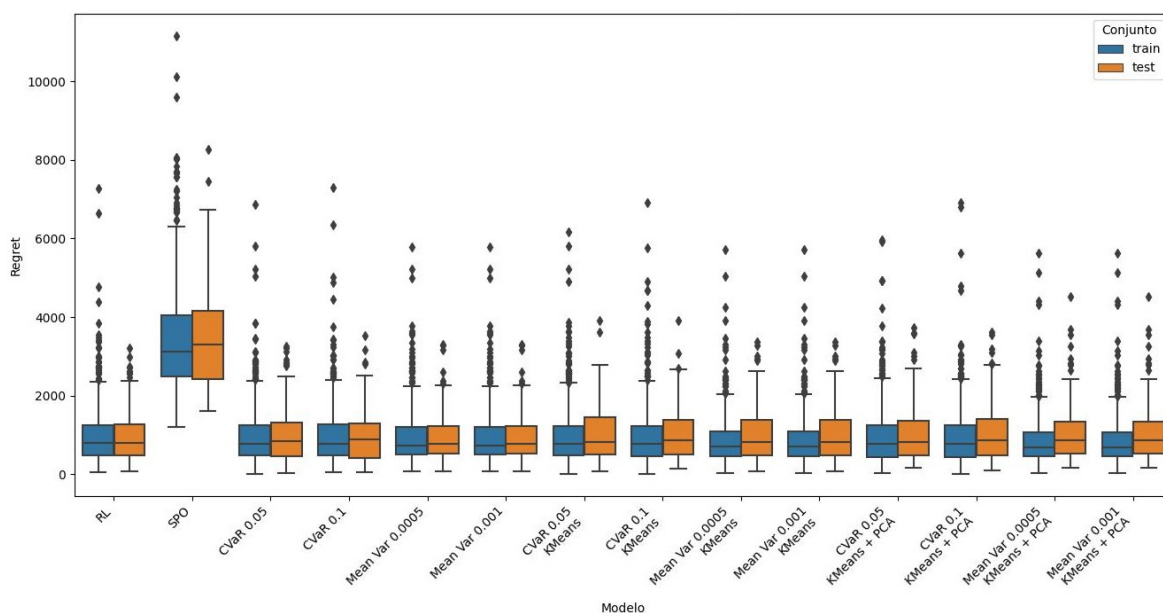


Figura 4.20: Distribución de arrepentimiento para una capacidad 60 con información local en KP01.

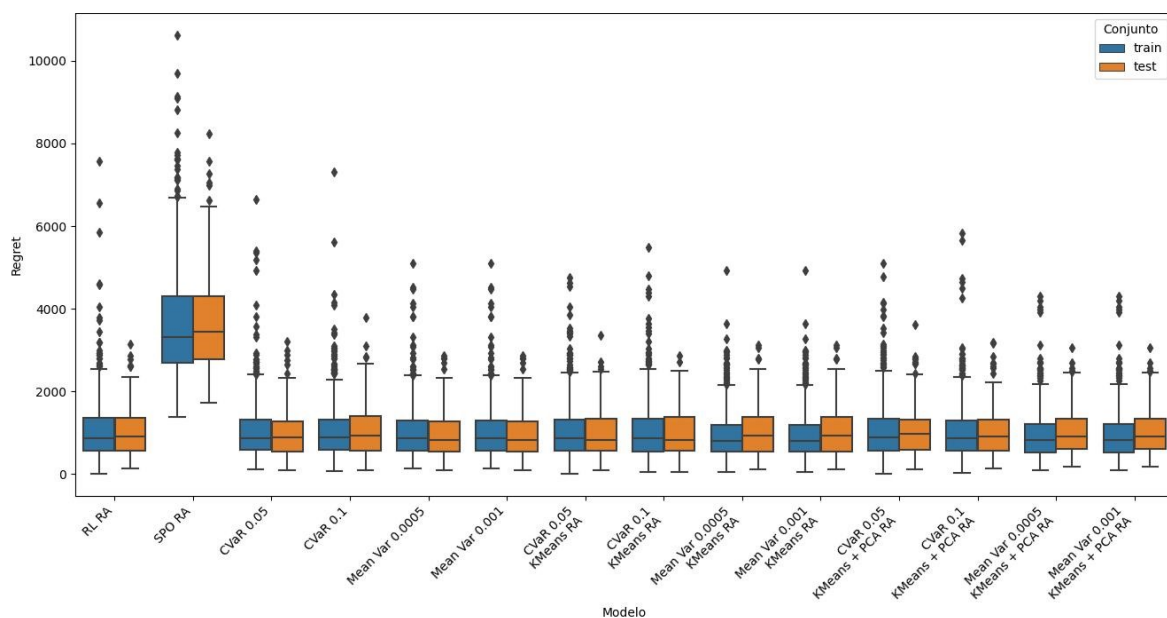


Figura 4.21: Distribución de arrepentimiento para una capacidad 120 con información global en KP01.

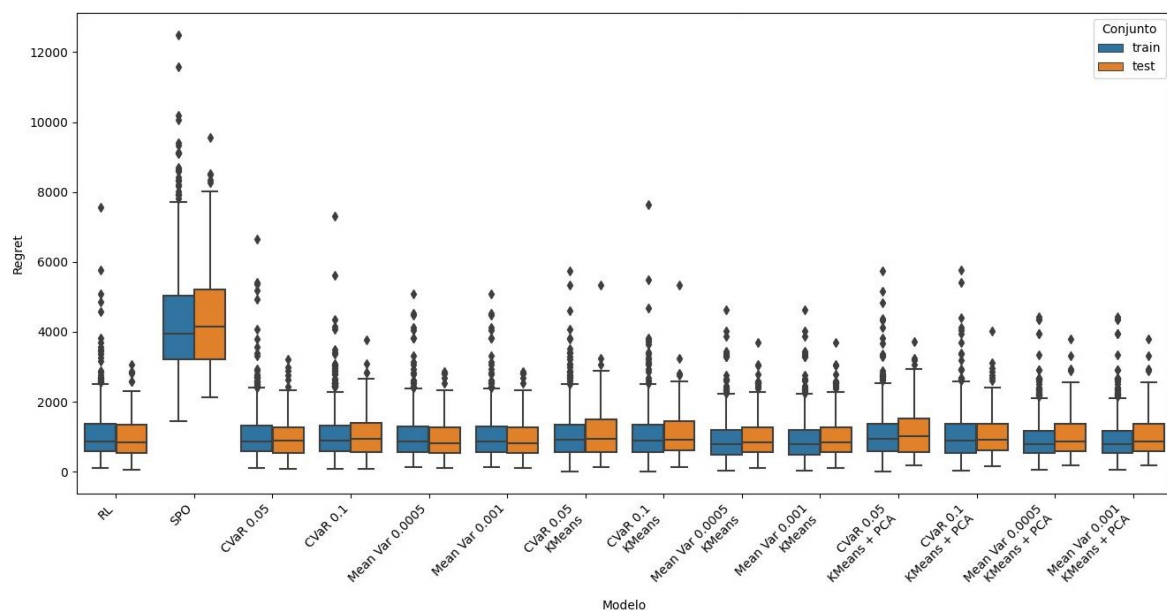


Figura 4.22: Distribución de arrepentimiento para una capacidad 120 con información local en KP01.

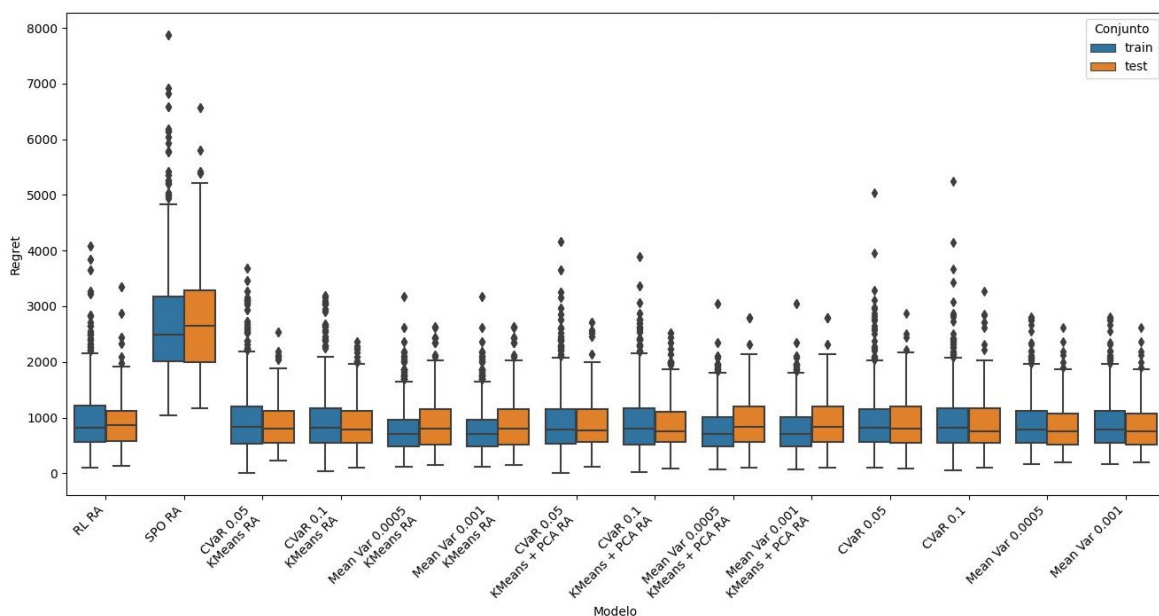


Figura 4.23: Distribución de arrepentimiento para una capacidad 180 con información global en KP01.

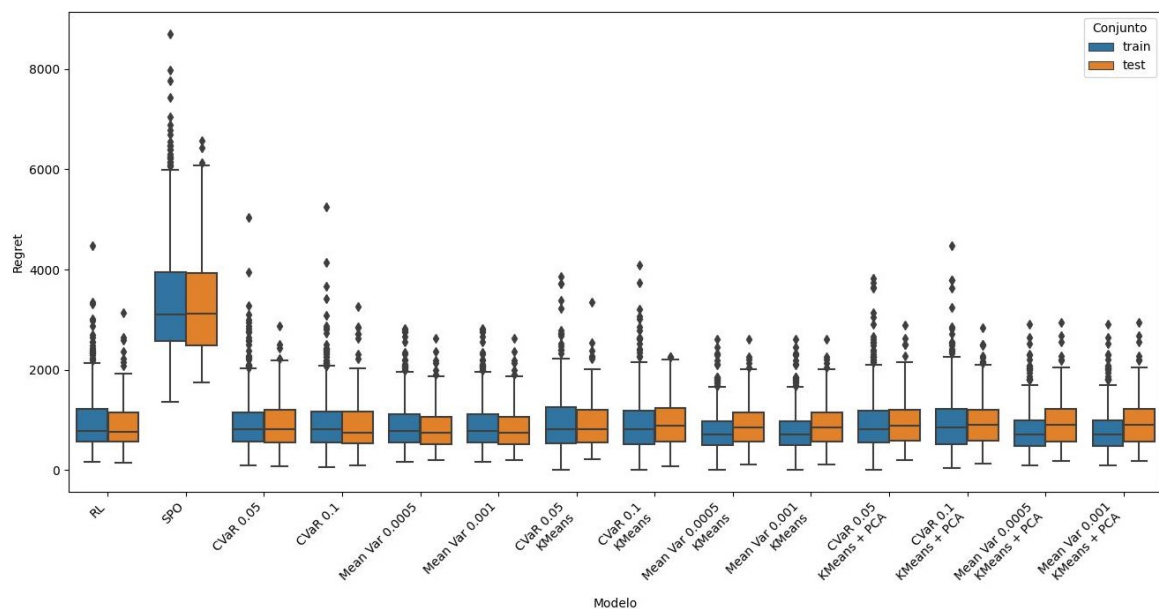


Figura 4.24: Distribución de arrepentimiento para una capacidad 180 con información local en KP01.

		Modelos					
		RL RA			RL		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Cap	60	0.027	0.001	0.028	0.126	0.001	0.127
	120	0.023	0.001	0.024	0.027	0.001	0.028
	180	0.024	0.001	0.025	0.028	0.001	0.029

Tabla 4.19: Tiempos computacionales [s] de RL RA vs. RL según capacidad para KP01.

		Modelos					
		SPO+ RA			SPO+		
		Train	Algoritmo	Total	Train	Algoritmo	Total
Cap	60	74.356	0.001	74.357	48.843	0.001	48.844
	120	142.130	0.001	142.131	80.488	0.001	80.489
	180	144.490	0.001	144.491	116.828	0.001	116.829

Tabla 4.20: Tiempos computacionales [s] de SPO+ RA vs. SPO+ según capacidad para KP01.

En ambos tipos de figuras, tanto en las Figuras RA (véase 4.19, 4.21, y 4.23) como en las NRA (véase 4.20, 4.22, y 4.24), se observó que la mediana, el rango intercuartílico (IQR), los límites superior e inferior, así como los valores extremos, fueron aproximadamente los mismos entre cada modelo respectivo. Además, al analizar exclusivamente los modelos aversos al riesgo con ajustes de cada figura contraparte, se notó que sus conjuntos de prueba tuvieron un rendimiento inferior en comparación con los conjuntos de entrenamiento.

Estos modelos ajustados también mostraron una mayor adición de varianza, evidenciada por una asimetría positiva. Notablemente, los modelos aversos originales fueron los que obtuvieron los mejores resultados, sin lugar a dudas.

Al comparar el comportamiento general entre ambos tipos de figuras, se destacó que el modelo SPO+ fue el más perjudicado al considerar información local, ya que los resultados de prueba no se encuentran dentro de los límites establecidos por el entrenamiento alcanzando un mayor arrepentimiento.

Una pequeña diferencia entre ambas figuras se observó en el modelo RL, donde al contemplar más información se redujo ligeramente la dispersión del bigote superior.

Finalmente, al comparar los tiempos promedio de ejecución en el entrenamiento de los modelos RL RA y SPO+ RA con sus contrapartes del grupo NRA, se observa una notable diferencia en cuanto a su menor duración, a pesar de que la matriz de atributos tenga el mismo tamaño en ambos casos.

Al evaluar los resultados de estos modelos y comparar las estimaciones de arrepentimiento y los tiempos de ejecución entre ambos grupos, surge la pregunta crucial de si realmente vale la pena tomar en cuenta toda la información disponible. La respuesta a esta interrogante es un contundente; no. Tanto si se contempla utilizar toda la información disponible para una capacidad de mochila grande o pequeña, la mejora resultante es prácticamente insignificante. Específicamente, para el caso de SPO+, resulta más ventajoso no considerar toda la información en todos los casos de capacidad; sin embargo, sus resultados son tan atípicos que no justifican utilizar este modelo para problemas de este tipo.

4.2.2. Conclusiones para KP01

- El experimento revela que los modelos de Mean Var α es el modelo más robusto y eficiente para abordar la incertidumbre en el problema de la mochila bajo diversas configuraciones (alta capacidad y numerosos escenarios) y mantiene un resultado consistente con el incremento del grado polinómico.
- RL aunque variable, también muestra adaptabilidad en diferentes contextos (aumento de escenarios y atributos).
- SPO no es recomendado debido a su alto arrepentimiento y comportamiento atípico.
- CVaR 0.1 y CVaR 0.05 se deben usar con cautela; CVaR 0.1 es menos confiable con mayores grados polinómicos, mientras que CVaR 0.05 exhibe mejoras con ciertos ajustes.

Capítulo 5

Conclusiones generales

El presente trabajo se desarrolló con el fin propósito de experimentar con dos enfoques diferentes: Optimización aversa al riesgo y Aprendizaje centrado en decisiones. El objetivo era determinar qué modelos podrían manejar mejor la incertidumbre asociada con la falta de conocimiento sobre los coeficientes del vector de costos en funciones objetivo, particularmente en problemas como el del camino más corto y el problema de la mochila. Se modificaron diversas variables en un caso base establecido para observar cómo estos cambios incluían en los resultados.

En el problema del camino más corto, se analizaron variables como el tamaño de la red, la cantidad de escenarios, el grado polinómico y la cantidad de atributos. Para el problema de la mochila, se consideraron variables similares, reemplazando el tamaño de la red por la capacidad de la mochila.

Retomando la hipótesis inicial, para el problema del camino más corto, se confirmó que un modelo DFL es más efectivo en enfrentar la incertidumbre, proporcionando mejores resultados en diversas configuraciones.

Por otro lado, para el problema de la mochila, un modelo con medida de riesgo Mean

Variance mostró resultados superiores. Sin embargo, los resultados menos favorables de SPO podrían deberse a la escasa cantidad de puntos muestreados y/o a la naturaleza combinatoria y no lineal del problema. Esto último requeriría una mayor investigación para obtener una explicación definitiva.

Es crucial enfatizar que, aunque los modelos aversos al riesgo generalmente no mostraron variaciones significativas en términos de arrepentimiento entre ellos, sí fueron altamente sensibles a los parámetros de la función objetivo, tales como ε o α , afectando de manera notable los tiempos de ejecución. Esto subraya la relevancia de ajustar adecuadamente estos parámetros al aplicar modelos de optimización en contextos de incertidumbre.

Es pertinente mencionar que la implementación efectiva de los resultados obtenidos mediante estas metodologías está condicionada por la disponibilidad de herramientas computacionales adecuadas que faciliten la investigación.

Uno de los principales desafíos en el aprendizaje orientado a la toma de decisiones radica en los recursos computacionales necesarios para entrenar el modelo integrado de predicción y optimización. Este enfoque requiere resolver el problema de optimización subyacente para cada muestra de datos en cada escenario, lo cual implica un coste computacional significativo, incluso para problemas de escala reducida. Este hecho representó una limitación considerable durante el trabajo con el problema de la mochila aplicado a conjuntos de datos de gran tamaño, especialmente porque se trata de un problema clasificado como NP-difícil, por lo cual, fue necesario emplear dos sistemas computacionales.

5.1. Trabajo futuro

Para futuros trabajos, se presentan diversos desafíos significativos. Primordialmente, es esencial realizar una revisión y, posiblemente, una reformulación de la estrategia de optimización empleada en el problema de la mochila bajo el método SPO+. La necesidad

de esta acción se debe a que los resultados obtenidos divergen considerablemente de los antecedentes documentados (Mandi et al., [2023](#)).

En segundo lugar, en relación con las medidas de riesgo contextualizadas, resulta esencial implementar técnicas de aprendizaje automático más aptas dada la diversidad y el volumen de los datos, con el objetivo de segmentar eficientemente la información contextual en distintos grupos. Una técnica prometedora a considerar es DBSCAN (Density-Based Spatial Clustering of Applications with Noise), un método de clustering no paramétrico que se basa en la densidad y exhibe características únicas y avanzadas, resultando especialmente eficaz en la detección de clusters de diversas formas y tamaños, así como en el manejo de outliers de manera efectiva (Ester et al., [1996](#)).

Se considera que las medidas de riesgo contextualizadas podrían beneficiarse considerablemente de una mayor cantidad de información clave. Aunque los métodos previamente empleados constituyeron un primer acercamiento, posiblemente no fueron completamente adecuados para abordar la totalidad de las complejidades presentes en el conjunto de datos.

Finalmente, un último desafío dentro del presente estudio, sería fundamental estipular resultados teóricos que establezcan una relación con los enfoques empleados.

Bibliografía

- Acerbi, C., & Tasche, D. (2002). On the coherence of expected shortfall. *Journal of banking & finance*, 26(7), 1487-1503.
- Agrawal, S., Ding, Y., Saberi, A., & Ye, Y. (2010). Correlation robust stochastic optimization. *Proceedings of the twenty-first annual ACM-SIAM symposium on Discrete Algorithms*, 1087-1096.
- Agrawal, S., Ding, Y., Saberi, A., & Ye, Y. (2012). Price of correlations in stochastic optimization. *Operations Research*, 60(1), 150-162.
- Ahmadi-Javid, A. (2011). An information-theoretic approach to constructing coherent risk measures. *2011 IEEE International Symposium on Information Theory Proceedings*, 2125-2127.
- Ahuja, R. K., Magnanti, T. L., & Orlin, J. B. (1995). *Network flows: theory, algorithms and applications*. Prentice hall.
- Amos, B., & Kolter, J. Z. (2017). Optnet: Differentiable optimization as a layer in neural networks. *International Conference on Machine Learning*, 136-145.
- Artzner, P., Delbaen, F., Eber, J.-M., & Heath, D. (1999). Coherent measures of risk. *Mathematical finance*, 9(3), 203-228.
- Averbakh, I., & Lebedev, V. (2004). Interval data minmax regret network optimization problems. *Discrete Applied Mathematics*, 138(3), 289-301.
- Bas, E. (2011). A capital budgeting problem for preventing workplace mobbing by using analytic hierarchy process and fuzzy 0–1 bidimensional knapsack model. *Expert Systems with Applications*, 38(10), 12415-12422.

- Bellman, R. (1958). On a routing problem. *Quarterly of applied mathematics*, 16(1), 87-90.
- Cambronero, C. G., & Moreno, I. G. (2006). Algoritmos de aprendizaje: knn & kmeans. *Inteligencia en Redes de Comunicación, Universidad Carlos III de Madrid*, 23.
- Cao, J., Yin, B., Lu, X., Kang, Y., & Chen, X. (2018). A modified artificial bee colony approach for the 0-1 knapsack problem. *Applied Intelligence*, 48, 1582-1595.
- Chang, P.-T., & Lee, J.-H. (2012). A fuzzy DEA and knapsack formulation integrated model for project selection. *Computers & Operations Research*, 39(1), 112-125.
- Chicoisne, R., Ordóñez, F., & Espinoza, D. (2018). Risk averse shortest paths: A computational study. *INFORMS Journal on Computing*, 30(3), 539-553.
- Cominetti, R. (2015). Equilibrium routing under uncertainty. *Mathematical Programming*, 151, 117-151.
- Dagnino, J., et al. (2014). Regresión lineal. *Revista chilena de anestesia*, 43(2), 143-149.
- Davies, D. L., & Bouldin, D. W. (1979). A cluster separation measure. *IEEE transactions on pattern analysis and machine intelligence*, (2), 224-227.
- Dijkstra, E. W. (1959). A note on two problems in connexion with graphs:(Numerische Mathematik, 1 (1959), p 269-271).
- Elmachtoub, A. N., & Grigas, P. (2022). Smart “predict, then optimize”. *Management Science*, 68(1), 9-26.
- Ester, M., Kriegel, H.-P., Sander, J., & Xu, X. (1996). Density-based spatial clustering of applications with noise. *Int. Conf. knowledge discovery and data mining*, 240(6).
- Feng, Y., Wang, G.-G., Deb, S., Lu, M., & Zhao, X.-J. (2017). Solving 0–1 knapsack problem by a novel binary monarch butterfly optimization. *Neural computing and applications*, 28, 1619-1634.
- Feng, Y., Wang, G.-G., Li, W., & Li, N. (2018). Multi-strategy monarch butterfly optimization algorithm for discounted {0-1} knapsack problem. *Neural Computing and Applications*, 30(10), 3019-3036.
- Ferber, A., Wilder, B., Dilkina, B., & Tambe, M. (2020). Mipaal: Mixed integer program as a layer. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(02), 1504-1511.

- Föllmer, H., & Schied, A. (2002). Convex measures of risk and trading constraints. *Finance and stochastics*, 6, 429-447.
- Gao, J., Zhao, Q., Ren, W., Swami, A., Ramanathan, R., & Bar-Noy, A. (2014). Dynamic shortest path algorithms for hypergraphs. *IEEE/ACM Transactions on networking*, 23(6), 1805-1817.
- Gholamy, A., Kreinovich, V., & Kosheleva, O. (2018). Why 70/30 or 80/20 relation between training and testing sets: A pedagogical explanation.
- Gowharji, W., & Whitefoot, K. S. (2021). Influence of omitted variables in consumer choice models on engineering design optimization solutions. *Journal of Mechanical Design*, 143(12), 121703.
- Gustriansyah, R., Suhandi, N., & Antony, F. (2020). Clustering optimization in RFM analysis based on k-means. *Indonesian Journal of Electrical Engineering and Computer Science*, 18(1), 470-477.
- Hao, J., & Ho, T. K. (2019). Machine learning made easy: a review of scikit-learn package in python programming language. *Journal of Educational and Behavioral Statistics*, 44(3), 348-361.
- Horowitz, E., & Sahni, S. (1974). Computing partitions with applications to the knapsack problem. *Journal of the ACM (JACM)*, 21(2), 277-292.
- Ingargiola, G. P., & Korsh, J. F. (1973). Reduction algorithm for zero-one single knapsack problems. *Management science*, 20(4-part-i), 460-463.
- Kolesar, P. J. (1967). A branch and bound algorithm for the knapsack problem. *Management science*, 13(9), 723-735.
- Kotary, J., Fioretto, F., Van Hentenryck, P., & Wilder, B. (2021). End-to-end constrained optimization learning: A survey. *arXiv preprint arXiv:2103.16378*.
- Kwon, C. (2011). Conditional value-at-risk model for hazardous materials transportation. *Proceedings of the 2011 winter simulation conference (WSC)*, 1703-1709.
- LAU, H. C., & LIM, M. K. (2004). Multi-period multi-dimensional Knapsack problem and its application to available-to-promise.

- Levi, R., Roundy, R., & Shmoys, D. B. (2004). Primal-dual algorithms for deterministic inventory problems. *Proceedings of the thirty-sixth annual ACM symposium on Theory of computing*, 353-362.
- Li, Z., Ma, L., & Zhang, H. (2012). Genetic mutation bat algorithm for 0-1 knapsack problem. *Comput Eng Appl*, 48(4), 50-53.
- Lim, S., Sommer, C., Nikolova, E., & Rus, D. (2013). Practical route planning under delay uncertainty: Stochastic shortest path queries. *Robotics: Science and Systems*, 8(32), 249-256.
- Madkour, A., Aref, W. G., Rehman, F. U., Rahman, M. A., & Basalamah, S. (2017). A survey of shortest-path algorithms. *arXiv preprint arXiv:1705.02044*.
- Mandi, J., Bucarey, V., Tchomba, M. M. K., & Guns, T. (2022). Decision-focused learning: through the lens of learning to rank. *International Conference on Machine Learning*, 14935-14947.
- Mandi, J., & Guns, T. (2020). Interior point solving for lp-based prediction+ optimisation. *Advances in Neural Information Processing Systems*, 33, 7272-7282.
- Mandi, J., Kotary, J., Berden, S., Mulamba, M., Bucarey, V., Guns, T., & Fioretto, F. (2023). Decision-focused learning: Foundations, state of the art, benchmark and future opportunities. *arXiv preprint arXiv:2307.13565*.
- Markowitz, H. (1952). Portfolio Selection. *The Journal of Finance*, 7(1), 77-91. Consultado el 3 de julio de 2023, desde <http://www.jstor.org/stable/2975974>
- Merkle, R., & Hellman, M. (1978). Hiding information and signatures in trapdoor knapsacks. *IEEE transactions on Information Theory*, 24(5), 525-530.
- Mulamba, M., Mandi, J., Diligenti, M., Lombardi, M., Bucarey, V., & Guns, T. (2020). Contrastive losses and solution caching for predict-and-optimize. *arXiv preprint arXiv:2011.05354*.
- Muter, I., & Sezer, Z. (2018). Algorithms for the one-dimensional two-stage cutting stock problem. *European Journal of Operational Research*, 271(1), 20-32.

- Nemhauser, G. L., & Ullmann, Z. (1969). Discrete dynamic programming and capital allocation. *Management Science*, 15(9), 494-505.
- Niepert, M., Minervini, P., & Franceschi, L. (2021). Implicit MLE: backpropagating through discrete exponential family distributions. *Advances in Neural Information Processing Systems*, 34, 14567-14579.
- Nikolova, E., Brand, M., & Karger, D. R. (2006). Optimal route planning under uncertainty. *Icaps*, 6, 131-141.
- Peeta, S., Salman, F. S., Gunnec, D., & Viswanath, K. (2010). Pre-disaster investment decisions for strengthening a highway network. *Computers & Operations Research*, 37(10), 1708-1719.
- Petrovic, S. (2006). A comparison between the silhouette index and the davies-bouldin index in labelling ids clusters. *Proceedings of the 11th Nordic workshop of secure IT systems*, 2006, 53-64.
- Pfeiffer, J., & Rothlauf, F. (2007). Analysis of greedy heuristics and weight-coded eas for multidimensional knapsack problems and multi-unit combinatorial auctions. *Proceedings of the 9th annual Conference on Genetic and Evolutionary Computation*, 1529-1529.
- Pisinger, D. (2002). Heuristics for the container loading problem. *European journal of operational research*, 141(2), 382-392.
- Pogančič, M. V., Paulus, A., Musil, V., Martius, G., & Rolinek, M. (2020). Differentiation of blackbox combinatorial solvers. *International Conference on Learning Representations*.
- Reniers, G. L., & Sörensen, K. (2013). An approach for optimal allocation of safety resources: Using the knapsack problem to take aggregated cost-efficient preventive measures. *Risk analysis*, 33(11), 2056-2067.
- Richardson, M. (2009). Principal component analysis. URL: <http://people.maths.ox.ac.uk/richardsonm/SignalProcPCA.pdf> (last access: 3.5. 2013). Aleš Hladnik Dr., Ass. Prof., Chair of Information and Graphic Arts Technology, Faculty of Natural Sciences and Engineering, University of Ljubljana, Slovenia ales.hladnik@ntf.uni-lj.si, 6(16), 4.

- Rockafellar, R. T., Uryasev, S., et al. (2000). Optimization of conditional value-at-risk. *Journal of risk*, 2, 21-42.
- Song, Y., & Shen, S. (2016). Risk-averse shortest path interdiction. *INFORMS Journal on Computing*, 28(3), 527-539.
- Subbalakshmi, C., Krishna, G. R., Rao, S. K. M., & Rao, P. V. (2015). A method to find optimum number of clusters based on fuzzy silhouette on dynamic data set. *Procedia Computer Science*, 46, 346-353.
- Taillandier, F., Fernandez, C., & Ndiaye, A. (2017). Real estate property maintenance optimization based on multiobjective multidimensional knapsack problem. *Computer-Aided Civil and Infrastructure Engineering*, 32(3), 227-251.
- Tranmer, M., & Elliot, M. (2008). Multiple linear regression. *The Cathie Marsh Centre for Census and Survey Research (CCSR)*, 5(5), 1-5.
- van Staden, H. E., Deprez, L., & Boute, R. N. (2022). A dynamic “predict, then optimize” preventive maintenance approach using operational intervention data. *European Journal of Operational Research*, 302(3), 1079-1096.
- Von Neumann, J., & Morgenstern, O. (1947). Theory of games and economic behavior, 2nd rev.
- Wagner, H. M., & Whitin, T. M. (1958). Dynamic version of the economic lot size model. *Management science*, 5(1), 89-96.
- Wei, J. T., Lin, S.-Y., Yang, Y.-Z., & Wu, H.-H. (2016). Applying data mining and RFM model to analyze customers' values of a veterinary hospital. *2016 International Symposium on Computer, Consumer and Control (IS3C)*, 481-484.
- Wilder, B., Dilkina, B., & Tambe, M. (2019). Melding the data-decisions pipeline: Decision-focused learning for combinatorial optimization. *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(01), 1658-1665.
- Wu, H., Zhou, Y., & Luo, Q. (2018). Hybrid symbiotic organisms search algorithm for solving 0-1 knapsack problem. *International Journal of Bio-Inspired Computation*, 12(1), 23-53.

- Wu, H.-H., Chang, E.-C., & Lo, C.-F. (2009). Applying RFM model and K-means method in customer value analysis of an outfitter. *Global Perspective for Competitive Enterprise, Economy and Ecology: Proceedings of the 16th ISPE International Conference on Concurrent Engineering*, 665-672.
- Yaari, M. E. (1969). Some remarks on measures of risk aversion and on their uses. *Journal of Economic Theory*, 1(3), 315-329. <https://EconPapers.repec.org/RePEc:eee:jetheo:v:1:y:1969:i:3:p:315-329>
- Yaari, M. E. (1987). The dual theory of choice under risk. *Econometrica: Journal of the Econometric Society*, 95-115.
- Yang, B., Guo, C., & Jensen, C. S. (2013). Travel cost inference from sparse, spatio temporally correlated time series using markov models. *Proceedings of the VLDB Endowment*, 6(9), 769-780.
- Yu, J., Gang y Yang. (1998). Sobre el problema robusto de la ruta más corta. *Informática & investigación de operaciones*, 457-468.
- Zhang, L., & Homem-de-Mello, T. (2017). An optimal path model for the risk-averse traveler. *Transportation science*, 51(2), 518-535.

Apéndice A

Resultados estadísticos de SPP

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	7.52	4.28	11.05	6.77	0.00	20.50	13.73	51.42
SPO RA	3.34	1.83	5.64	3.81	0.00	10.50	6.69	36.56
CVaR 0.05	4.94	2.81	7.85	5.04	0.00	15.50	10.46	19.32
CVaR 0.1	4.94	2.81	7.85	5.04	0.00	15.50	10.46	19.32
Mean Var 0.0005	4.09	2.13	6.58	4.45	0.00	13.00	8.55	50.34
Mean Var 0.001	4.09	2.13	6.58	4.45	0.00	13.00	8.55	50.34
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
4.82		4.93						

Tabla A.1: Datos estadísticos de Caso Base (Grid: 5 · 5, Esc: 1000, k: 10, Deg: 8)

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	0.86	0.00	2.07	2.07	0.00	5.00	2.93	17.80
SPO RA	1.25	0.06	2.45	2.39	0.00	6.00	3.61	33.48
CVaR 0.05	1.67	0.69	3.64	2.95	0.00	8.00	5.05	29.66
CVaR 0.1	1.67	0.69	3.64	2.95	0.00	8.00	5.05	29.66
Mean Var 0.0005	1.44	0.30	3.02	2.73	0.00	7.50	4.77	41.32
Mean Var 0.001	1.44	0.30	3.02	2.73	0.00	7.50	4.77	41.32
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
1.39		2.64						

Tabla A.2: Datos estadísticos con grilla 3 · 3

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	8.54	5.28	14.63	9.35	0.48	28.00	18.17	36.89
SPO RA	6.32	3.60	9.60	6.00	0.12	18.50	12.38	27.23
CVaR 0.05	7.92	5.51	12.33	6.82	0.88	21.00	13.30	29.93
CVaR 0.1	8.42	4.87	12.15	7.28	0.08	23.00	15.64	39.29
Mean Var 0.0005	7.33	4.99	10.21	5.22	0.33	18.00	12.45	37.89
Mean Var 0.001	7.33	4.99	10.21	5.22	0.33	18.00	12.45	37.89
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
7.64		6.65						

Tabla A.3: Datos estadísticos con grilla 7 · 7

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	2.58	1.42	4.99	3.57	0.04	9.00	5.39	16.00
SPO RA	4.06	2.44	5.67	3.23	0.00	9.00	5.77	12.06
CVaR 0.05	4.93	2.33	7.82	5.49	0.00	13.00	7.51	19.08
CVaR 0.1	4.93	2.33	7.82	5.49	0.00	13.00	7.51	19.08
Mean Var 0.0005	4.30	2.21	6.81	4.60	0.00	10.00	5.40	31.29
Mean Var 0.001	4.30	2.21	6.81	4.60	0.00	10.00	5.40	31.29
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
4.18		4.50						

Tabla A.4: Datos estadísticos con 200 escenarios

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	1.76	0.75	3.42	2.67	0.00	6.80	4.13	48.43
SPO RA	3.10	1.46	5.85	4.39	0.00	11.50	7.11	54.30
CVaR 0.05	4.60	2.40	8.54	6.14	0.00	17.00	10.86	53.61
CVaR 0.1	4.60	2.40	8.54	6.14	0.00	17.00	10.86	53.61
Mean Var 0.0005	4.48	2.26	7.99	5.73	0.00	15.80	10.07	49.72
Mean Var 0.001	4.48	2.26	7.99	5.73	0.00	15.80	10.07	49.72
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
3.84		5.13						

Tabla A.5: Datos estadísticos con 5000 escenarios

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	2.64	1.79	3.67	1.88	0.00	6.80	4.92	7.86
SPO RA	2.45	1.36	3.62	2.26	0.00	7.00	4.74	7.94
CVaR 0.05	2.57	1.69	4.12	2.43	0.16	7.00	4.41	6.63
CVaR 0.1	2.79	1.77	3.72	1.95	0.00	6.50	4.55	8.61
Mean Var 0.0005	2.77	1.73	3.65	1.92	0.00	6.30	4.38	8.29
Mean Var 0.001	2.77	1.73	3.65	1.92	0.00	6.30	4.38	8.29
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
2.67		2.06						

Tabla A.6: Datos estadísticos con grado polinómico 1

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	2.76	1.61	3.95	2.34	0.00	7.01	4.67	7.01
SPO RA	2.16	1.22	3.41	2.19	0.00	6.60	4.41	6.71
CVaR 0.05	2.96	1.90	4.04	2.14	0.00	6.60	4.46	7.43
CVaR 0.1	2.76	1.84	3.97	2.13	0.00	7.10	4.97	7.77
Mean Var 0.0005	2.85	1.82	3.97	2.15	0.00	6.94	4.79	6.94
Mean Var 0.001	2.85	1.82	3.97	2.15	0.00	6.94	4.79	6.94
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
2.72		2.18						

Tabla A.7: Datos estadísticos con grado polinómico 2

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	2.20	1.12	3.47	2.35	0.00	7.20	4.85	11.03
SPO RA	2.46	1.22	3.92	2.70	0.00	7.90	5.20	9.86
CVaR 0.05	3.25	1.83	4.87	3.04	0.00	8.90	5.86	11.23
CVaR 0.1	3.59	2.06	4.91	2.85	0.00	8.90	6.05	12.09
Mean Var 0.0005	3.25	1.83	4.87	3.04	0.00	8.90	5.86	11.23
Mean Var 0.001	3.25	1.83	4.87	3.04	0.00	8.90	5.86	11.23
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
3.00		2.84						

Tabla A.8: Datos estadísticos con grado polinómico 4

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	2.05	0.97	3.28	2.31	0.00	7.00	4.69	10.37
SPO RA	3.06	1.43	4.56	3.13	0.00	9.30	6.17	23.51
CVaR 0.05	4.13	2.42	6.85	4.43	0.09	13.50	8.98	21.18
CVaR 0.1	4.13	2.42	6.85	4.43	0.09	13.50	8.98	21.18
Mean Var 0.0005	3.60	1.99	6.01	4.02	0.00	12.00	7.98	27.50
Mean Var 0.001	3.60	1.99	6.01	4.02	0.00	12.00	7.98	27.50
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
3.43		3.72						

Tabla A.9: Datos estadísticos con grado polinómico 6

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	3.10	1.23	5.83	4.60	0.00	12.50	7.90	31.98
SPO RA	3.52	1.81	6.73	4.92	0.00	13.80	8.88	57.64
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
3.31		4.76						

Tabla A.10: Datos estadísticos con 5 atributos

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	2.78	1.30	5.41	4.11	0.00	11.00	6.89	72.05
SPO RA	3.13	1.84	5.82	3.98	0.00	11.70	7.72	36.10
$\bar{x}_{M_e}$		$\bar{x}_{IQR}$						
2.96		4.05						

Tabla A.11: Datos estadísticos con 20 atributos

Apéndice B

Resultados estadísticos de KP01

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	899.50	572.75	1359.00	786.25	132.00	2355.00	1436.75	3139.00
SPO RA	3453.50	2776.00	4312.50	1536.50	1723.00	6470.00	3205.50	8232.00
CVaR 0.05	884.00	545.00	1280.00	735.00	90.00	2335.00	1510.00	3214.00
CVaR 0.1	935.00	554.25	1403.00	848.75	94.00	2165.00	1222.25	3779.00
Mean Var 0.0005	820.50	540.25	1263.25	723.00	97.00	2338.00	1518.00	2872.00
Mean Var 0.001	820.50	540.25	1263.25	723.00	97.00	2338.00	1518.00	2872.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
871.90				763.20				

Tabla B.1: Datos estadísticos de Caso Base (Cap: 120, Esc: 750, k: 5, Deg: 8)

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	834.00	462.50	1315.50	853.00	108.00	2580.00	1619.00	2994.00
SPO RA	3081.00	2435.75	3823.50	1387.75	1469.00	5750.00	2893.25	7518.00
CVaR 0.05	848.00	456.00	1321.75	868.75	29.00	2485.00	1590.25	3263.00
CVaR 0.1	887.00	419.50	1299.25	879.75	42.00	2510.00	1588.25	3519.00
Mean Var 0.0005	784.50	528.50	1226.75	698.25	80.00	2265.00	1486.75	3305.00
Mean Var 0.001	784.50	528.50	1226.75	698.25	80.00	2265.00	1486.75	3305.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
827.60				799.00				

Tabla B.2: Datos estadísticos con capacidad 60

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	865.50	572.50	1124.50	552.00	126.00	2420.00	1742.00	3344.00
SPO RA	2641.50	2003.25	3279.50	1276.25	1175.00	5100.00	2648.75	6565.00
CVaR 0.05	808.50	546.75	1206.25	659.50	78.00	2175.00	1437.50	2875.00
CVaR 0.1	750.50	540.00	1160.50	620.50	99.00	2030.00	1310.50	3265.00
Mean Var 0.0005	749.00	518.25	1065.75	547.50	198.00	1865.00	1119.50	2624.00
Mean Var 0.001	749.00	518.25	1065.75	547.50	198.00	1865.00	1119.50	2624.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
784.50				585.40				

Tabla B.3: Datos estadísticos con capacidad 180

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	902.00	542.75	1156.50	613.75	18.00	1935.00	1303.25	3723.00
SPO RA	2990.00	2522.00	3603.75	1081.75	1542.00	4199.00	1575.25	7163.00
CVaR 0.05	854.50	524.25	1189.00	664.75	193.00	2040.00	1182.25	2940.00
CVaR 0.1	900.50	651.75	1177.50	525.75	174.00	1720.00	1020.25	2949.00
Mean Var 0.0005	832.00	546.75	1101.50	554.75	230.00	1925.00	1140.25	2808.00
Mean Var 0.001	832.00	546.75	1101.50	554.75	230.00	1925.00	1140.25	2808.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
864.20				582.75				

Tabla B.4: Datos estadísticos con 200 escenarios

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	885.00	538.00	1335.00	797.00	81.00	2485.00	1607.00	6325.00
SPO RA	3587.50	2904.00	4398.75	1494.75	1556.00	6500.00	3449.25	10867.00
CVaR 0.05	878.00	533.50	1381.00	847.50	94.00	2648.00	1706.50	5615.00
CVaR 0.1	864.00	528.00	1357.75	829.75	48.00	2590.00	1712.25	6510.00
Mean Var 0.0005	860.00	521.75	1309.75	788.00	87.00	2480.00	1605.00	3749.00
Mean Var 0.001	860.00	521.75	1309.75	788.00	87.00	2480.00	1605.00	3749.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
869.40				810.05				

Tabla B.5: Datos estadísticos con 1500 escenarios

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	28.00	21.00	40.00	19.00	3.00	68.00	46.00	79.00
SPO RA	2311.50	2129.50	2439.50	310.00	1972.00	2643.00	361.00	-
CVaR 0.05	28.50	21.00	40.00	19.00	6.00	63.10	38.10	78.00
CVaR 0.1	27.00	20.00	40.00	20.00	3.00	70.00	47.00	74.00
Mean Var 0.0005	25.00	18.00	34.00	16.00	4.00	54.00	34.00	84.00
Mean Var 0.001	25.00	18.00	34.00	16.00	4.00	54.00	34.00	84.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
26.70				18.00				

Tabla B.6: Datos estadísticos con grado polinómico 1

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	51.00	36.00	75.00	39.00	9.00	133.00	85.00	173.00
SPO RA	2550.50	2406.00	2782.00	376.00	2143.00	3330.00	811.00	3761.00
CVaR 0.05	48.00	33.00	66.75	33.75	7.00	101.00	60.25	181.00
CVaR 0.1	49.00	31.00	66.00	35.00	4.00	117.00	78.00	140.00
Mean Var 0.0005	46.00	33.00	63.00	30.00	1.00	105.00	74.00	126.00
Mean Var 0.001	46.00	33.00	63.00	30.00	1.00	105.00	74.00	126.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
48.00				33.55				

Tabla B.7: Datos estadísticos con grado polinómico 2

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	134.50	91.25	190.25	99.00	18.00	320.00	203.00	-
SPO RA	2012.00	1839.00	2242.50	403.50	1451.00	2165.00	310.50	2999.00
CVaR 0.05	140.50	102.00	186.50	84.50	14.00	274.00	175.50	492.00
CVaR 0.1	137.50	92.25	196.25	104.00	15.00	319.00	200.00	405.00
Mean Var 0.0005	135.00	90.25	172.75	82.50	3.00	280.00	194.50	321.00
Mean Var 0.001	135.00	90.25	172.75	82.50	3.00	280.00	194.50	321.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
136.50				90.50				

Tabla B.8: Datos estadísticos con grado polinómico 4

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	356.00	245.75	496.00	250.25	27.00	857.00	579.75	1628.00
SPO RA	2498.00	2184.75	2947.75	763.00	1530.00	3935.00	1642.00	5019.00
CVaR 0.05	360.00	248.00	548.00	300.00	29.00	1000.00	671.00	1552.00
CVaR 0.1	365.00	231.75	582.75	351.00	34.00	1095.00	710.00	1436.00
Mean Var 0.0005	323.50	232.25	500.00	267.75	89.00	831.00	474.25	1835.00
Mean Var 0.001	323.50	232.25	500.00	267.75	89.00	831.00	474.25	1835.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
345.60				287.35				

Tabla B.9: Datos estadísticos con grado polinómico 6

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	1165.00	737.75	1726.75	989.00	181.00	3025.00	1855.00	11935.00
SPO RA	3683.50	2908.00	4813.75	1905.75	1833.00	6870.00	3131.25	17994.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
2424.25				1447.38				

Tabla B.10: Datos estadísticos con 3 atributos

Modelos	Estadísticos							
	M_e	Q_1	Q_3	IQR	Min	Max	Bigotes	Último atípico
RL RA	668.50	477.25	1054.75	577.50	61.00	1900.00	1261.50	2931.00
SPO RA	3182.50	2635.00	3795.50	1160.50	1397.00	5500.00	2942.50	6944.00
$\bar{x}_{M_e}$				$\bar{x}_{IQR}$				
1925.50				869.00				

Tabla B.11: Datos estadísticos con 7 atributos