



Universidad
de O'Higgins

Escuela de Ingeniería
Ingeniería Civil Industrial

Aplicación de Aprendizaje Enfocado en la Decisión en optimización de dotación de personal en retail.

Katherine Barahona Norambuena
Profesor(a) guía: Víctor Bucarey

Memoria para optar al título y/o grado de Ingeniera Civil Industrial

Rancagua, Chile
Enero, 2024

Dedicatoria

A mis padres, María y Juan.

Agradecimientos

Quiero expresar mis sinceros agradecimientos a todos aquellos que formaron parte de este proceso académico. En especial, quiero agradecer a mis padres por su comprensión, apoyo incondicional, y motivación para seguir adelante a pesar de los desafíos.

A mis hermanos, por su apoyo y sus palabras de aliento. A mis amigos por estar siempre dispuestos a escuchar y animarme en momentos difíciles.

Por último, agradezco a mi profesor guía Víctor Bucarey por orientarme y apoyarme en este proceso de investigación.

Índice

RESUMEN	5
INTRODUCCIÓN	6
PLANTEAMIENTO DEL PROBLEMA	8
HIPÓTESIS.....	9
PREGUNTAS DE INVESTIGACIÓN	10
OBJETIVO GENERAL	11
OBJETIVOS ESPECÍFICOS	11
MARCO TEÓRICO Y REVISIÓN DE LITERATURA	12
MARCO METODOLÓGICO	19
RESULTADOS	28
ALCANCES Y LIMITACIONES	36
CONCLUSIÓN.....	38
REFERENCIAS	39

Resumen

En el sector minorista, las ventas de una tienda están relacionadas con el tráfico de clientes y el personal que lo atiende, por lo mismo, las ventas mejoran si se tiene el personal suficiente para satisfacer a los clientes. No obstante, los factores que logran determinar cuál es una cantidad de personal óptima en cada bloque horario por turno es un problema complejo de resolver. Una de las razones radica en la incertidumbre que hay en la relación entre las ventas y la cantidad de personal requerido, siendo un desafío en la gestión de turnos para una empresa.

El propósito de esta memoria es resolver un problema de dotación de personal en el cual el impacto que tiene la fuerza laboral en las ventas es incierto. Para estimar este parámetro desconocido, se realizará una comparación de dos metodologías de aprendizaje: enfocado en la predicción y enfocado en la decisión. En el primer caso, el modelo de aprendizaje se entrena minimizando el error de predicción y en el segundo caso minimizando el impacto en el proceso de decisión. Este último método está basado en lo que plantean Elmachetoub y Grigas (2022) en su artículo “Smart: Predict, then Optimize”.

Para la metodología enfocada en la decisión, se evaluaron distintas funciones de pérdida para estimar los parámetros necesarios para predecir las ventas. Además, se evaluó la eficacia en términos de error de predicción y error de decisión. Los principales hallazgos obtenidos concluyen que al aplicar metodologías de Aprendizaje Enfocado en la Decisión se puede tener mayores ganancias económicas.

Los resultados obtenidos respaldan la hipótesis de que el enfoque de Aprendizaje Enfocado en la Decisión supera al enfoque de Aprendizaje Enfocado en la Predicción para la toma de decisiones de personal. Por ende, la manera en la que se predice un parámetro incierto para un modelo de optimización influye y tiene un impacto económico. Aunque este estudio tiene ciertas limitaciones, contribuye en la comprensión de una nueva metodología que predice inteligentemente parámetros desconocidos, logrando asociarla a un desafío real de la industria.

Palabras clave: aprendizaje enfocado en la predicción, aprendizaje enfocado en la decisión, aprendizaje supervisado, modelos de dotación de personal, aprendizaje de máquina

Introducción

Las empresas de retail se enfrentan a desafíos en la gestión de operaciones, especialmente en lo que respecta a la decisión de contratación de mano de obra. La gestión del personal es importante, ya que el personal contratado está directamente relacionado con el servicio ofrecido diariamente por el retail. Sin embargo, determinar el nivel de dotación de personal es complejo por lo inestable que es el tráfico de clientes.

Según Chuang et al. (2016), las ventas del sector minorista explican su rendimiento conforme al tráfico de clientes y la mano de obra que se posee, de este modo, se puede decir que las ventas están influenciadas por distintos factores, siendo uno de éstos el personal contratado. Las tiendas minoristas de retail poseen bastante incertidumbre respecto a las ventas que pueden tener por día, y los costos laborales son uno de los mayores gastos en este tipo de empresas (Yung et al., 2020). Por ello, tener la capacidad de gestionar y programar de manera eficiente los turnos de trabajo y la mano de obra requerida para satisfacer a la demanda de clientes, es un problema importante a investigar para obtener decisiones acertadas de personal.

Elmachtoub y Grigas (2022) mencionan que cuando se aplican herramientas analíticas en la investigación de operaciones para tomar decisiones, en general, son abordados como una combinación de un problema de aprendizaje automático y de optimización. En este caso, el aprendizaje automático o aprendizaje de máquinas (por sus siglas en inglés, ML) se utiliza para crear un modelo de predicción que estima algún parámetro incierto necesario para tomar una decisión óptima. Este método utiliza los datos históricos para entrenar un algoritmo, aprender y estimar resultados basados en estos datos (Mohri et al., 2018). En él existen distintos escenarios de aprendizaje, uno de estos es el aprendizaje supervisado. Dentro del aprendizaje supervisado se encuentra la labor de predicción, donde se diseña un modelo de predicción en el cuál existe una respuesta asociada o etiqueta a cada observación de los datos. Un ejemplo de tarea en este escenario es la regresión lineal, la cual puede predecir parámetros a partir del entrenamiento del modelo.

Por otro lado, se utiliza la optimización para crear un modelo para adoptar decisiones eficaces. Los modelos de optimización regularmente buscan maximizar los beneficios de tener

cierta configuración de dotación de personal teniendo en cuenta los costos que se incurren en estos. No obstante, los parámetros que se involucran en este tipo de problema muchas veces no son conocidos, por lo que es necesario estimarlos de alguna manera.

Para estimarlas se utiliza el aprendizaje automático y estas estimaciones son utilizadas en el modelo de optimización, sin embargo, no toma la relación entre el error de los parámetros estimados y su efecto en el problema de optimización, si no que la toma de manera independiente (Mandi et al., 2022). Este método mide la calidad del resultado estimado en función del error de predicción, por lo que se quiere lograr un error de predicción bajo, a este enfoque se le llamará Aprendizaje Enfocado en la Predicción (por sus siglas en inglés, PFL).

En consecuencia, surge una herramienta que “está diseñada fundamentalmente en generar modelos de predicción cuyo objetivo es minimizar el error de decisión, no el error de predicción” (Elmachtoub y Grigas, 2022, p. 9), denominada Aprendizaje Enfocado en la Decisión (por sus siglas en inglés, DFL). La principal diferencia entre PFL y DFL es cómo se entrena el modelo, ya que en ambos modelos hay aprendizaje automático, se estiman predicciones y con ello se toman decisiones.

Esta investigación se centrará en cuantificar el impacto de los dos enfoques en la toma de decisiones de dotación de personal por turno. Este estudio es fundamental para disponer de métodos que aborden estos desafíos de manera eficaz.

La metodología de esta investigación consiste utilizar modelos de aprendizaje en un problema de decisión de dotación de personal y analizar tanto la precisión como la eficacia de los resultados obtenidos. La medida de cuantificación será el error de decisión (*regret*) que mide el exceso de costo (o pérdida de beneficio) que uno incurre por tener errores en las predicciones.

Planteamiento del problema

La industria minorista se encuentra constantemente ante desafíos respecto a la utilización efectiva de recursos. Uno de los problemas más relevantes es el asignar de manera adecuada al personal para satisfacer la demanda de clientes que está en constante cambio. Dado que el personal contratado posee interacción directa con el cliente, se sabe que la dotación de personal junto con el tráfico de clientes afecta las ventas (Chuang et al., 2016). Por esto, esta decisión se convierte en un factor esencial para tener eficiencia operativa, y con ello mayores beneficios económicos en la organización.

El impacto de la dotación de personal sobre las ventas es un parámetro esencial para la toma de decisión del diseño de turnos laborales, sin embargo, es desconocido a priori. Por lo mismo, es fundamental diseñar métodos de estimación de este parámetro.

Esta memoria tiene como objetivo comparar distintas metodologías de aprendizaje para estimar el parámetro desconocido de las ventas al contratar cierta cantidad de personal.

Hipótesis

Dada la problemática expuesta, se plantea la siguiente hipótesis:

En el contexto de determinar la dotación de personal en retail, la utilización de Aprendizaje Enfocado en la Decisión para determinar la relación entre la fuerza laboral y las ventas tiene un impacto económico favorable en el proceso de toma de decisiones por sobre el Aprendizaje Enfocado en la Predicción. Esto quiere decir que las ganancias que se dejan de percibir cuando las ventas son estimadas con DFL son menores que cuando estas son estimadas por PFL.

Preguntas de investigación

Las preguntas de investigación de esta memoria son:

- ¿Existe una mejora utilizando el Aprendizaje Enfocado en la Decisión versus usar solo Aprendizaje Enfocado en la Predicción, en la toma de decisiones de dotación de personal?
- ¿Cuánto disminuye el error de decisión utilizando el Aprendizaje Enfocado en la Decisión?
- ¿Cómo adoptar a un contexto semi supervisado el Aprendizaje Enfocado en la Decisión?

Objetivo general

Analizar si el uso de Aprendizaje Enfocado en la Decisión mejora el proceso de toma de decisiones para el problema de dotación de personal.

Objetivos específicos

- Implementar modelo de Aprendizaje Enfocado en la Predicción.
- Adaptar modelo de Aprendizaje Enfocado en la Decisión a problema de dotación de personal.
- Comparar métricas de calidad del proceso de decisión, utilizando estimaciones basadas Aprendizaje Enfocado en la Decisión y Aprendizaje Enfocado en la Predicción.
- Realizar simulaciones numéricas para analizar la robustez del aprendizaje.

Marco teórico y revisión de literatura

I. Revisión de conceptos

Este estudio se centra en la implementación y posterior análisis de un enfoque para diseñar mejores modelos de predicción y toma de decisiones, siendo un tema relevante para el área de administración empresarial. En esta sección, se definirán de forma general distintos conceptos y términos para la comprensión de la investigación.

a. Industria retail

Se entiende por industria retail o minorista, a las empresas que venden productos al detalle a consumidores y que abarcan desde cadenas de tiendas a supermercados (Guerrero-Martínez, 2012). Por otro lado, la dotación de personal se entenderá como el “personal permanente contratado de un servicio como honorario asimilado a grado y jornadas laborales permanentes” (Dirección de presupuestos, 2011, p.5).

b. Aprendizaje automático y modelos de predicción

El aprendizaje automático se define como “métodos computacionales que utilizan la experiencia para mejorar el rendimiento o hacer predicciones precisas” (Mohri et al., 2018, p.1). Refiriéndose con “experiencia” a la información disponible, como los datos. El aprendizaje automático o aprendizaje de máquina (ML, por sus siglas en inglés) tiene como objetivo “generar predicciones precisas para elementos invisibles y diseñar algoritmos eficientes y robustos para producir predicciones, incluso para problemas a gran escala” (Mohri et al., 2018, p.3). Las tareas que se pueden realizar con ML varían desde clasificación, clustering, regresión, entre otros. Estas tareas pertenecen a escenarios de aprendizaje, que cambian respecto al tipo de datos disponible que se tiene para entrenar el modelo, donde se destacan tres principalmente: aprendizaje supervisado, aprendizaje no supervisado y aprendizaje semi supervisado.

El aprendizaje supervisado es cuando se tiene disponible “un conjunto de elementos etiquetados como datos de entrenamiento y hace predicciones para todos los puntos invisibles” (Mohri et al., 2018, p.6), es decir, para cada elemento, existe una etiqueta o respuesta asociada.

Los problemas comunes en este caso son los de regresión y clasificación. Contrario a este, el aprendizaje no supervisado “recibe exclusivamente datos de entrenamiento sin etiquetar y hace predicciones para todos los puntos invisibles” (Mohri et al., 2018, p.6), un ejemplo es el clustering. Por otro lado, el aprendizaje semi supervisado es cuando se tienen datos que poseen etiquetas y otros que no, y se hacen las predicciones en base a ello (Mohri et al., 2018).

En esta memoria se realizará Aprendizaje Enfocado en la Predicción (PFL, por sus siglas en inglés) utilizando una regresión lineal múltiple. La regresión lineal es una herramienta de aprendizaje supervisado con un enfoque simple para predecir una etiqueta “Y” cuantitativa con respecto a una variable predictiva “X”, suponiendo una relación lineal entre ambas (James et al., 2021). De forma similar se define la regresión lineal múltiple, pero ésta posee más de una variable predictiva y ambas corresponden a “una herramienta útil para predecir una respuesta cuantitativa” (James et al., 2021, p.59). Una predicción k-ésima para una observación específica de los datos (y_k) es explicada en base a variables descriptivas o características ($X_1, X_2, \dots, X_k$), y el modelo de regresión lineal múltiple toma la forma:

$$y_k = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon_k,$$

donde β_k son los coeficientes que asocian a las variables descriptivas y la predicción, los que son estimados en base al modelo, y ε_k corresponde al error de predicción o residuo k-ésimo, que es la diferencia entre la respuesta observada y el valor de la predicción.

Al estimar a través de un modelo de predicción es importante evaluar la precisión de los pronósticos. Por lo que precisar algunas métricas relevantes es conveniente para tener mayor comprensión de su efectividad. Se destacan las siguientes:

- RSS: La suma residual de cuadrados es “la diferencia entre el i-ésimo valor de respuesta observado y el i-ésimo valor de respuesta predicho por el modelo lineal” (James et al., 2021, p.62), siendo:

$$RSS = \varepsilon_1^2 + \varepsilon_2^2 + \dots + \varepsilon_n^2$$

- RSE: El error estándar residual “es una estimación de la desviación estándar de ε . En términos generales es la cantidad promedio en que la respuesta se desviará de la línea de regresión verdadera” (James et al., 2021, p.69). Y se define como:

$$RSE = \sqrt{\frac{1}{n-p} RSS} = \sqrt{\frac{1}{n-p} \sum_{i=1}^n (y_i - \hat{y}_i)^2}$$

Siendo n el número total de observaciones de la muestra y p la cantidad de parámetros del modelo.

- R^2 : Este estadístico es “la proporción de la varianza explicada – y por eso siempre toma el valor entre 0 y 1” (James et al., 2021, p.70).

$$R^2 = 1 - \frac{\sum_t (Y_t - \hat{Y}_t)^2}{\sum_t (Y_t - \bar{Y}_t)^2}$$

c. Modelos de optimización

Un problema de optimización busca encontrar la mejor solución, maximizando o minimizando una función objetivo a partir de ciertas restricciones, en definitiva, es un problema de decisión. Éste busca optimizar un sistema con los datos que describen una función objetivo, por ejemplo, maximizar los beneficios netos de la venta de productos. El problema está sujeto a restricciones que forman la región factible de posibles soluciones al problema. Como mencionan Mandi et al. (2023), los procesos de toma de decisiones usualmente pueden describirse mediante problemas de optimización con restricciones. Estos pueden escribirse de manera general como:

$$\begin{aligned} x^* &= \operatorname{argmin}_x f(x) \\ \text{s. t: } g(x) &\leq 0 \\ h(x) &= 0. \end{aligned}$$

Donde x^* es la solución óptima a encontrar, en este caso minimizando una función objetivo que cumpla con las restricciones asociadas a las funciones g y h . En general, “*argmin*” es un conjunto de soluciones óptimas, por lo que esta igualdad es un abuso de notación. En esta memoria se dirá que se tiene un oráculo que selecciona un elemento del “*argmin*”.

d. Modelos de optimización lineal mixta

Los modelos de optimización lineal mixta (por sus siglas en inglés, MILP), se definen como modelos que tienen variables continuas y enteras de forma lineal, aplicadas a una función

objetivo y sus respectivas restricciones. Muchos problemas reales se pueden modelar con un MILP, como en áreas de transporte, ingeniería, procesos y planificaciones. Se pueden expresar de manera general como:

$$\begin{aligned} \min \quad & c^T x + d^T y \\ \text{s.t.} \quad & Ax + By \leq b \\ & x \in R^{n_1}; \quad y \in \mathbb{Z}^{n_2} \end{aligned}$$

Con x un vector de n_1 variables continuas, y un vector de variables enteras de n_2 , c, d como vectores de coeficientes, A y B como matrices, y b es un vector de componentes.

Es importante mencionar que el modelo descrito corresponde a un caso especial del modelo que fue detallado en la sección anterior.

e. Aprendizaje Enfocado en la Decisión

Esta investigación considera un concepto nuevo que es el Aprendizaje Enfocado en la Decisión (por sus siglas en inglés, DFL), o “Smart: Predict, Then Optimize” (SPO) como lo definen Elmachtoub y Grigas (2022). Esta metodología “aprovecha directamente la estructura de un problema de optimización, es decir, su objetivo y sus restricciones para diseñar mejores modelos de predicción” (Elmachtoub y Grigas, 2022, p.1). En este sentido, el foco es minimizar el error de decisión, en vez del error de predicción como lo hace un modelo de aprendizaje automático. El error de decisión o *regret* que se produce a partir de la predicción, se refiere a que

mide el error al predecir el vector de costos de un problema de optimización nominal con restricciones lineales, convexas o enteras. La pérdida corresponde a la brecha de suboptimalidad con respecto al vector de costos verdadero, debido a la aplicación de una decisión posiblemente incorrecta inducida por el vector de costos predicho (Elmachtoub y Grigas, 2022, p.10).

Asimismo, el *regret* se definió por otros autores como, “la diferencia entre el valor objetivo óptimo de información completa y el valor objetivo obtenido mediante la decisión prescriptiva” (Mandi et al., 2023, p.5), y lo definen como:

$$\text{Regret}(x^*(\hat{c}), c) = f(x^*(\hat{c}), c) - f(x^*(c), c)$$

Siendo $\hat{c}$ las predicciones, c los parámetros reales y $x^*(\hat{c})$ la solución óptima que se obtiene al resolver el vector de costos predicho, cuyo valor objetivo $f(x^*(\hat{c}), c)$ se aproxima al valor óptimo real $f(x^*(c), c)$.

II. Revisión de literatura

Este estudio se basa principalmente en dos grandes etapas: la primera es la predicción de ventas de una tienda retail a partir del entrenamiento de un modelo de ML, mientras que la segunda es enlazar estas predicciones a un modelo de optimización y medir el error de decisión que se produce posterior a ello. Por otro lado, se aplicará una metodología que entrenará el modelo con respecto al error de decisión, minimizando este error. En base a esto, esta sección presenta investigaciones y referencias para abordar el tema de estudio.

Esta investigación surge de la afirmación que las ventas de una tienda retail se ven afectadas por los cambios existentes en la dotación de personal. Perdikaki et al. (2012) estudiaron el efecto entre el tráfico de clientes, la mano de obra y el rendimiento de las ventas. Evalúan la venta a partir de una tasa de conversión definida como la relación entre las transacciones o compras y el tráfico de clientes, junto con la relación entre las ventas y las transacciones. Los resultados fueron que el volumen de ventas muestra rendimientos de escala decrecientes con respecto al tráfico de clientes, y la mano de obra modera el impacto del tráfico en las ventas. Por lo tanto, demuestra que la mano de obra es crucial en los ingresos de una empresa. Christiansen (2020) por otra parte, desarrolla un modelo de predicción de tráfico de clientes para optimizar la dotación de personal que se ajuste a la demanda en tiendas físicas de retail. Esto a partir de la identificación de variables exógenas que afecten el tráfico, y una herramienta de predicción robusta "Prophet". La herramienta permite realizar modelos de ajuste y pronóstico de series con sets de datos donde haya observaciones faltantes, nulas, cambios de tendencia y tendencias no lineales. El modelo permite componentes de una serie temporal como tendencia, estacionalidad, ruido, en este caso, incluye estacionalidad semanal, mensual y anual, y regresores que indican si hay días feriados o con lluvia.

Hay estudios que modelan, a partir de programación lineal mixta, la optimización de personal necesaria para resolver un problema de fuerza laboral. Miranda y Palacios (2017) abordan un problema asociado a la dotación de personal en reposición de productos para una empresa proveedora de tiendas retail. Utilizan modelos de optimización lineal entera mixta para diseñar jornadas laborales diarias y semanales para la empresa SC Johnson, resultando en un aumento de puntos de venta. Si bien, estos estudios emplean metodologías de optimización, no proponen minimizar el error de tomar una decisión en base a la solución propuesta. Sin embargo, la perspectiva común es impulsar a un mejor desempeño en las operaciones basándose en proporcionar una planificación correcta para mejorar las ventas.

Por otra parte, Henao et al. (2015) analizan el impacto de añadir empleados polivalentes en la industria minorista, entendiendo como polivalente a trabajadores que hagan múltiples funciones. Ellos plantean que el uso de la programación de personal para reducir el exceso de personal o la falta de éste en servicios de este tipo de industria se ve perjudicado por la poca flexibilidad que hay en las tareas que realiza el personal. Para este análisis utiliza un modelo de programación lineal entera mixta para determinar qué empleados están capacitados para ciertas actividades y tareas, y realizar asignaciones en la programación del personal. Dando como resultado que los costos más bajos se obtienen en escenarios donde la oferta y la demanda de personal está en equilibrio con personal polivalente.

Nyman (2022) plantea un enfoque de análisis de sensibilidad para mejorar los resultados de la optimización de la fuerza laboral mediante una simulación de la capacitación de los empleados. En esta investigación se plantea un algoritmo basado en la simulación de mejorar los resultados a partir de una solución existente de optimización de la mano de obra, entregando nuevas habilidades que deben tener los trabajadores para reducir el tiempo ocioso de éstos.

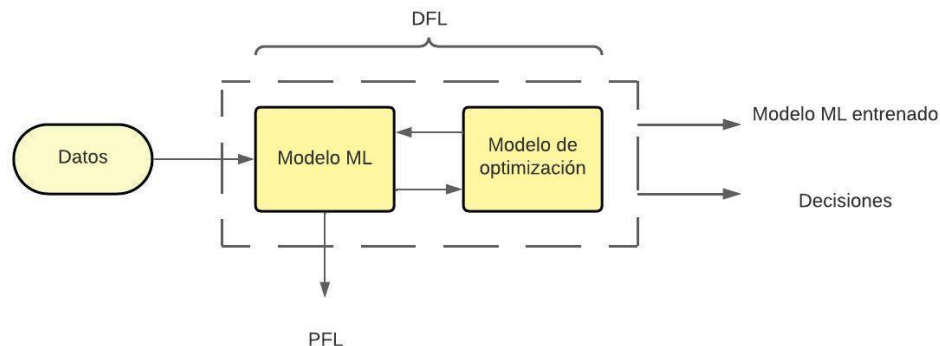
En definitiva, las tiendas minoristas experimentan fluctuaciones en el flujo de clientes, lo que dificulta determinar la cantidad de personal necesario para proporcionar un buen servicio al cliente. Esto destaca la importancia de estudiar en detalle cómo obtener la capacidad laboral adecuada para satisfacer la demanda. Siguiendo esta perspectiva, Chuang et al. (2016), propone una planificación laboral basada en el tráfico de clientes, utilizan métodos de estimación que

dependen del tráfico de clientes y mano de obra, y proponen una función de respuesta de ventas respecto a estos parámetros. Luego, utilizan esta función de ventas para desarrollar una heurística para determinar cuál debe ser su mano de obra, incorporando la pérdida de predicción. De igual modo Yung et al. (2020) exponen una metodología similar a esta, desarrollan un modelo econométrico para estimar el efecto de la mano de obra en las ventas por horas, y también realizan proyecciones de afluencia. Esto lo utilizan como entrada a un problema de optimización para la necesidad de personal hora a hora basándose en las proyecciones del tráfico de clientes, con el fin de encontrar una planificación laboral ideal que maximicen los beneficios de la tienda.

A diferencia de los trabajos anteriores, esta memoria de investigación tiene como finalidad medir el error de decisión, y con ello estudiar el impacto que tiene en la toma de decisiones. Esta investigación se relaciona con el estudio de Elmachoub y Grigas (2022), quienes basan su investigación en predecir parámetros desconocidos de un problema de optimización estocástica. En este caso, predicen un vector de costos de la función objetivo a partir de entrenamiento de datos, por lo que el modelo de predicción se utiliza en la estructura del problema de optimización. Además, proponen una nueva función de pérdida sustitutiva denominada “SPO+” para el entrenamiento de modelos de predicción; ésta es convexa y está adaptada al problema de optimización, por lo tanto, se puede optimizar de manera eficaz.

Marco metodológico

El propósito de este estudio se orienta en medir el impacto económico que tiene el aplicar distintos enfoques de aprendizaje para estimar un parámetro desconocido en un problema de optimización, con el fin de tomar decisiones respecto a la dotación de personal en una tienda retail. Esta se divide en dos grandes etapas: en líneas generales se tienen datos que simulan una situación realista en la industria minorista. En la primera etapa, estos datos serán entrenados en un modelo de aprendizaje automático (ML) para realizar predicciones, y en esta fase se utilizará la metodología PFL. Por otro lado, la segunda etapa se centra en la aplicación de un modelo de optimización que permite tomar decisiones en base a una situación. La segunda metodología que se analiza en esta memoria, DFL, incorpora ambos modelos, de ML y de optimización. El enfoque DFL entrena un modelo lineal de manera distinta, ya que aplica la estructura del modelo de optimización eficazmente en el entrenamiento del modelo de predicción. La metodología de investigación se puede visualizar en el Esquema 1.



Esquema 1. *Marco metodológico de investigación. Elaboración propia.*

I. Modelo de optimización nominal

Tras establecer el marco metodológico general, es necesario especificar el modelo de optimización desarrollado por Yung et al. (2020) ya que es fundamental en el contexto de esta investigación. Este modelo está basado en las siguientes consideraciones laborales: la semana de trabajo se divide en jornadas mañana y tarde, se consideran bloques como la combinación del

día de la semana con la jornada, la cantidad de trabajadores máxima por jornada es de 7 trabajadores, y el conjunto de 14 bloques se denominada H . Para cada bloque existe un conjunto de turnos:

- Full time (FT): Turnos que consisten en cinco días laborales (10 bloques) los cuales pueden ser 4 días y dos medias jornadas (un bloque diario por dos días), o cinco días enteros.
- Part time 30 (PT30): Turnos que son equivalentes a tres días laborales. Estos pueden conformarse por dos días completos y dos medias jornadas, o tres días laborales completos.
- Part time 10 (PT10): Turnos equivalentes a dos días laborales. Incluyendo siempre un día completo del fin de semana. Por ejemplo, un turno puede comprender un sábado completo y dos medias tardes lunes y martes.

Se denota como S al conjunto de todos los turnos factibles. Cada turno $s \in S$ se representa a través de un vector $a^s \in \{0,1\}^{14}$, entonces a_h^s es un parámetro que toma el valor 1 si el turno s asigna al trabajador al bloque h , 0 si no. Estos turnos son generados antes de la optimización y representan un total de 1407 vectores a^s .

Cada turno tiene un costo asociado al salario de los trabajadores el cual depende de la cantidad de bloques y días que se trabaja. El salario se denomina w_s y su valor es en pesos chilenos. El salario por turno será:

$$w_s = 3.900 \cdot \alpha_s + 1.600 \cdot \beta_s$$

Con α_s como la cantidad de horas por turno: para turnos FT son 45, para PT30 son 30, y para PT10 son 10, y β_s corresponde a la cantidad de días que tiene que ir el trabajador a la tienda.

El siguiente modelo de optimización lineal mixta tiene como propósito encontrar los turnos que maximizan la utilidad esperada por el retail, refiriéndose a utilidad como los beneficios económicos (ingresos menos costos) de tener $m \in \{0,1 \dots, M\}$ trabajadores en el bloque h . Esta utilidad se denomina $U_{h,m}$. El modelo de optimización es el siguiente:

$$MÁX \sum_{m=0}^M \sum_{h \in H} U_{h,m} \cdot Y_{h,m} - \sum_{s \in S} w_s \cdot J_s$$

$$s. t: \sum_{s \in S} a_h^s \cdot J_s = \sum_{m=0}^M m \cdot Y_{h,m} \quad \forall h \in H \quad (1)$$

$$\sum_{m=0}^M Y_{h,m} = 1 \quad \forall h \in H \quad (2)$$

$$J \in \mathbb{Z}_+^{|S|}, Y \in \{0,1\}^{14 \times M} \quad (3)$$

Donde $Y_{h,m}$ es una variable binaria que vale 1 si hay exactamente m trabajadores en el bloque h , 0 si no. J_s corresponde a la cantidad de turnos tipo s seleccionados. En cuanto a las restricciones, (1) asegura la consistencia de la relación de las variables J_s y las variables $Y_{h,m}$. Esto quiere decir que en cada bloque h tenga la cantidad de trabajadores m correcta según los turnos asignados. La expresión (2) obliga a $Y_{h,m}$ valer 1 solo para un valor m en el bloque h , es decir se asegura un número exacto de trabajadores para h . Finalmente, (3) corresponde a la naturaleza de las variables.

Adicionalmente, se agregan restricciones al modelo de optimización que hacen referencia a normas de trabajo que se consideran para el estudio, estas son:

$$\sum_{s \in PT10} J_s \leq \sum_{s \in PT30} J_s \quad (4)$$

$$\sum_{s \in PT10} J_s \leq \sum_{s \in FT} J_s \quad (5)$$

$$\sum_{s \in FT} Dom_s \cdot J_s \leq \frac{1}{2} \sum_{s \in FT} J_s \quad (6)$$

Con Dom_s como una constante binaria que vale 1 si el turno s incluye trabajar un domingo, y 0 en otro caso, (4) restringe que la cantidad de turnos PT10 es menor o igual que los turnos PT30, (5) restringe la cantidad de turnos PT10 debe ser menor o igual que los turnos FT. Mientras que (6) corresponde a que no más de la mitad de los trabajadores FT trabajen los domingos. Esto en la práctica permite cumplir con la legislación de que los trabajadores FT tienen dos domingos libres al mes.

El modelo de optimización descrito permite saber qué tipo de turnos son los que maximizan los beneficios de la tienda, contribuyendo al objetivo de tomar decisiones de manera correcta. No obstante, en este estudio y en la realidad no se dispone a priori las utilidades que

tiene la tienda al contar con m trabajadores en el bloque h , por lo que se deben estimar. En este caso, para estimarlas se utilizarán dos metodologías, PFL Y DFL. Ambas metodologías tienen la finalidad de predecir un parámetro desconocido, pero se diferencian en la manera en que se entrena el modelo.

II. Metodología PFL

Este enfoque emplea un algoritmo ML para predecir un parámetro desconocido. En este caso específico, se utiliza un modelo de regresión lineal múltiple para estimar las ventas de una tienda minorista. Inicialmente se debe disponer de una base de datos, y luego se entrena el modelo de regresión. Con ello se predicen las ventas y se resuelve el problema de dotación de personal en el modelo de optimización descrito con anterioridad a través del software “Gurobi Optimizer”. Entregando así, la cantidad de trabajadores en el bloque h óptimo para generar mayores beneficios, y la cantidad de turnos tipo s seleccionados para ello.

a. Entrenamiento de modelo y predicción de ventas

En esta etapa se procede a la división de la base de datos en entrenamiento y testeo, en este caso, la base de datos fue simulada al no poseer datos de una tienda real. La instancia de entrenamiento consta del 80% de los datos y la instancia de testeo del 20% de los datos. Se entrena el modelo utilizando la librería de Python “Scikit-Learn” que permite aplicar algoritmos de ML, facilitando el aprendizaje e implementación. Específicamente se utiliza “LinearRegression” y se estiman las ventas respecto a los coeficientes, y al intercepto entregado por el modelo. Este entrenamiento se realiza para el 80% y el 20% de la instancia de datos.

b. Modelo de optimización nominal

Una vez estimadas las predicciones de ventas a través del modelo ML, se ingresan en el modelo y se calcula el *regret* ejecutando las utilidades reales ($U_{h,m}$) y las utilidades predichas ($\hat{U}_{h,m}$) de los datos de entrenamiento y de testeo en el modelo de optimización. Con ello resultan

los vectores $Y_{h,m}^*, J_s^*$ como vectores óptimos de haber conocido $U_{h,m}$ y los vectores $\hat{Y}_{h,m}, \hat{J}_s$ resultantes de $\hat{U}_{h,m}$.

III. Metodología DFL

En el marco de esta investigación, se adopta la metodología de Aprendizaje Enfocado en la Decisión desarrollada por Elmachtoub y Grigas en su artículo “Smart: Predict, then Optimize”. La elección de este método es idónea para abordar los objetivos de esta memoria, ya que ofrece una estructura sólida que permite predecir inteligentemente para la toma de decisiones.

En relación con el artículo de Elmachtoub y Grigas (2022), cuyo enfoque se basa en un modelo de optimización que busca minimizar una función objetivo, se destaca que uno de los parámetros clave de esta función aún no se conoce con certeza. Sin embargo, tienen observaciones relacionadas con dicho parámetro. Ellos dividen su investigación en dos etapas, primero se predice este parámetro utilizando datos históricos y datos característicos que expliquen el parámetro incierto. Luego, se optimiza minimizando la función objetivo, sin embargo, se plantean si se puede predecir de manera inteligente este parámetro, tomando en cuenta el problema de decisión en el que dicho parámetro será utilizado. Para ello entrenan un modelo tratando de minimizar el *regret*, siendo este la pérdida de beneficio por no estar prediciendo bien. El dilema surge al intentar minimizar el *regret*, ya que se enfrentan a la complicación de que esta función no es convexa, lo que dificulta su resolución. Por ende, desarrollan una función de pérdida convexa ($l_{spo} +$) que aproxima el *regret* siendo una cota superior de éste.

Esta memoria desarrolla esta metodología construyendo el estimador DFL con la expectativa de encontrar los coeficientes β que mejoren la regresión lineal que estima las ventas. Este estimador se denominará β^{loss} y este proceso se divide en tres etapas: transformar el modelo de optimización nominal a minimización, encontrar el conjunto de soluciones factibles (Y, J) , y encontrar β^{loss} al minimizar una función de pérdida específica en ese conjunto de soluciones.

a. Transformación de modelo de optimización nominal

Según el modelo de optimización expuesto y detallado al inicio de esta sección, se lleva a cabo una transformación para convertirlo en un problema de minimización. Esto se logra mediante la equivalencia que implica cambiar un problema de maximización a uno de minimización, multiplicando la función objetivo por -1 , sin afectar el conjunto de restricciones asociadas Ω .

$$\begin{aligned} \text{MÁX} \sum_{m=0}^M \sum_{h \in H} U_{h,m} \cdot Y_{h,m} - \sum_{s \in S} w_s \cdot J_s \quad & \leftrightarrow \quad \text{MIN} \sum_{s \in S} w_s \cdot J_s - \sum_{m=0}^M \sum_{h \in H} U_{h,m} \cdot Y_{h,m} \\ \text{s. t: } (Y, J) \in \Omega \quad & \quad \quad \text{s. t: } (Y, J) \in \Omega \end{aligned}$$

b. Conjunto de soluciones factibles

Para cada semana de la muestra de entrenamiento se ejecuta el nuevo modelo de optimización y se almacenan las soluciones (Y, J) en un conjunto P . Este conjunto es un diccionario que contiene la semana como clave y una tupla $((Y, J))$, es decir, $\{(Y^p, J^p)_{p \in P}\}$. Hay que recordar que este conjunto P describe los turnos y la cantidad de trabajadores que se requieren, y la cantidad de turnos tipo s utilizados.

c. Estimador β^{loss}

Una vez determinado el conjunto de soluciones factibles P , se procede a minimizar la función de pérdida en este conjunto. En este punto, se emplean dos funciones de pérdida: “Pointwise” propuesta por Mandi et al. (2022) y “SPO+” propuesta por Elmachetoub y Grigas (2022), con el fin de observar la existencia de mejoras comparando ambas funciones.

i. Estimador $\beta^{Pointwise}$

Mandi et al. (2022) proponen una función de pérdida llamada Pointwise que es convexa, y se asocia al problema que abarca esta memoria con $U_{h,m}^i$ como las utilidades reales de la tienda, y $\hat{U}_{h,m}^i$ como las utilidades predichas que, en este caso, se estiman a partir de una regresión lineal. El modelo de optimización que minimiza esta función de pérdida se define de la siguiente manera:

$$\text{MIN}_{\beta} \quad \frac{1}{N} \sum_{i=1}^N \text{loss}(U_{h,m}^i, \hat{U}_{h,m}^i(\beta))$$

$$\begin{aligned}
&\leftrightarrow \text{MIN}_{\beta} \frac{1}{N} \sum_{i=1}^N \frac{1}{|P|} \sum_{(Y^p, J^p)_{p \in P}} ||w^T \cdot J^p - \hat{O}^{iT} \cdot Y^p - w^T \cdot J^p + U^{iT} \cdot Y^p|| \\
&\leftrightarrow \text{MIN}_{\beta} \frac{1}{N} \sum_{i=1}^N \frac{1}{|P|} \sum_{(Y^p, J^p)_{p \in P}} ||U^{iT} \cdot Y^p - \hat{O}^{iT} \cdot Y^p|| \\
&\leftrightarrow \text{MIN}_{\beta} \frac{1}{N} \sum_{i=1}^N \frac{1}{|P|} \sum_{(Y^p, J^p)_{p \in P}} || \sum_{(h,m)} U_{h,m}^i \cdot Y_{h,m}^p - \hat{U}_{h,m}^i \cdot Y_{h,m}^p || \\
&\leftrightarrow \text{MIN}_{\beta} \frac{1}{N} \cdot \frac{1}{|P|} \sum_{i=1}^N \sum_{(Y^p, J^p)_{p \in P}} | \sum_{(h,m)} U_{h,m}^i \cdot Y_{h,m}^p - \left(\beta_0 + \sum_{k=1}^K \beta_k \cdot X_k + \beta_m \cdot \ln(m+1) \right) \cdot Y_{h,m}^p |
\end{aligned}$$

En donde, en la primera línea se hace explícito la dependencia de $\hat{U}_{h,m}^i$ de β y luego se omite para mayor claridad. De esta función se puede aplicar la siguiente regla:

$$\begin{aligned}
\text{MÍN } |x| &\leftrightarrow \text{MÍN } \alpha \\
\text{s. t:} &\quad \alpha \geq x \\
&\quad \alpha \geq -x
\end{aligned}$$

Aplicando esta regla al modelo descrito, el estimador $\beta^{\text{Pointwise}}$ se obtiene del siguiente modelo de optimización.

$$\begin{aligned}
&\text{MIN}_{\beta, \Delta ip} \frac{1}{N} \cdot \frac{1}{P} \sum_{i=1}^N \sum_{(Y^p, J^p)_{p \in P}} \Delta ip \\
&\text{s. t:} \quad \Delta ip \geq U_{h,m}^i \cdot Y_{h,m}^p - \left(\beta_0 + \sum_{k=1}^K \beta_k \cdot X_k + \beta_m \cdot \ln(m+1) \right) \cdot Y_{h,m}^p \\
&\quad \Delta ip \geq \left(\beta_0 + \sum_{k=1}^K \beta_k \cdot X_k + \beta_m \cdot \ln(m+1) \right) \cdot Y_{h,m}^p - U_{h,m}^i \cdot Y_{h,m}^p
\end{aligned}$$

Al resolver el modelo entregaran los parámetros β óptimos. Estos coeficientes β son los que se utilizan para estimar las ventas ($\hat{U}_{h,m}^i$) bajo este estimador de la metodología DFL.

ii. Estimador $\beta^{\text{SPO+}}$

Por otro lado, Elmachtoub y Grigas (2022) proponen una función de pérdida distinta. Esta función de pérdida se le denomina SPO+, y tiene como objetivo, al igual que “Pointwise”, encontrar los coeficientes betas óptimos que permitan estimar el parámetro incierto. La función de pérdida SPO+ asociada a esta memoria, se define como:

$$\text{loss SPO}^+ := \text{MAX}_{(Y,J) \in P} \{ -w \cdot J + (2\hat{U} - U) \cdot Y \} + 2w \cdot J^* - 2\hat{U} \cdot Y^* - z^*(U)$$

Recordando que w son los costos asociados, U las utilidades reales de la tienda, $\hat{U}$ las predicciones de las utilidades, que se estiman a partir de una regresión lineal, y $z^*(U)$ como el beneficio óptimo de tener la solución (Y^*, J^*) , siendo estos últimos los vectores óptimos al ingresar las utilidades reales al modelo de optimización de esta sección.

Se busca minimizar la pérdida función de pérdida SPO+, por lo que el modelo de optimización se define como:

$$\begin{aligned} & \min_{\beta} \frac{1}{N} \sum_{i=1}^N \text{loss spo}^+(U^i, \hat{U}^i(\beta)) \\ & \min_{\beta} \frac{1}{N} \sum_{i=1}^N \max_{(Y,J) \in P} \{-w \cdot J^i + (2\hat{U}^i(\beta) - U^i) \cdot Y^i\} + 2w \cdot J^{i*} - 2\hat{U}^i(\beta) \cdot Y^{i*} - z^*(U^i) \end{aligned}$$

Quedando explicito la dependencia de $\hat{U}^i$ de los coeficientes β , por lo que luego se omiten por simplicidad. Además, se denominará la expresión $\max_{(Y,J) \in P} \{-w \cdot J^i + (2\hat{U}^i - U^i) \cdot Y^i\}$ como θ_i , quedando el modelo como:

$$\min_{\beta} \frac{1}{N} \sum_{i=1}^N \theta_i + 2w \cdot J^{i*} - 2\hat{U}^i \cdot Y^{i*} - z^*(U_i)$$

Se puede notar que las expresiones $2w \cdot J^{i*}$ y $z^*(U^i)$ no dependen de los coeficientes β , por lo que la función objetivo queda reescrita de manera más simple:

$$\begin{aligned} & \min_{\beta} \frac{1}{N} \sum_{i=1}^N \theta_i - 2\hat{U}^i \cdot Y^{i*} \\ \text{s.t: } & \theta_i \geq \sum_{s \in S} -w_s \cdot J_s^{i*} + \sum_{(h,m)} (2\hat{U}_{h,m}^i - U_{h,m}^i) \cdot Y_{h,m}^{i*} \quad \forall i, p \end{aligned} \quad (1)$$

$$\theta_i + 2w \cdot J^{i*} - 2\hat{U}^i \cdot Y^{i*} - z^*(U_i) \geq 0 \quad \forall i \quad (2)$$

La restricción (1) corresponde a la definición de θ_i , mientras que la restricción (2) define que la función de pérdida SPO+ debe ser positiva.

En esta memoria, se considera el siguiente modelo lineal para la relación entre ventas y fuerza laboral

$$\hat{U} = \beta_0 + \sum_{k=1}^K \beta_k \cdot X_k + \beta_m \cdot \ln(m+1),$$

donde X_k es vector de atributos y m es la cantidad de trabajadores observados. Para capturar la relación de tasa marginal decreciente entre las ventas y la cantidad de trabajadores se utiliza en la regresión el $\ln$ de la cantidad de trabajadores.

Una vez descrito el modelo de optimización que minimiza esta función de pérdida, se estiman los coeficientes β óptimos, y se predicen las ventas o utilidades de la tienda.

IV. Error de decisión versus Error de predicción

Este estudio tiene como objetivo analizar si el uso del DFL mejora el proceso de toma de decisiones para el problema de dotación de personal en la industria de retail, por lo que comparar métricas de calidad de ambas metodologías aplicadas es crucial. Para esto, se comparan dos tipos de medidas de evaluación de los modelos aplicados en las metodologías, el error de predicción y el error de decisión.

En el caso de el error de predicción, se mide a través del error cuadrático medio o MSE de cada semana, para la instancia de entrenamiento y testeo, permitiendo observar qué tan bueno es el modelo prediciendo.

Mientras que el error de decisión o *regret*, que mide la pérdida de beneficio que se obtiene por tener errores en las predicciones de cada semana, para los datos de entrenamiento y testeo. Hay que tener presente que el *regret* se calcula como “lo que hubiese ganado si hubiese conocido exactamente $U_{h,m}$ ” menos “lo que realmente gané”. En este estudio se mide de la siguiente manera:

$$Regret = (U_{h,m}^T \cdot Y_{h,m}^* - w_s^T \cdot J_s^*) - (U_{h,m}^T \cdot \hat{Y}_{h,m} - w_s^T \cdot \hat{J}_s)$$

Estos errores de decisión son calculados para el uso de metodología PFL y para la metodología DFL, y se visualizan los resultados con el objetivo de hacer un análisis.

Resultados

Una vez descritas ambas metodologías que permiten estimar las ventas, se desea analizar el comportamiento de estos métodos para realizar un análisis, con la finalidad de aprobar o rechazar la hipótesis de esta memoria. Sin embargo, para este estudio no se dispone de información real de alguna tienda de retail, por lo que se realiza una simulación de datos.

Los atributos que conforman esta base de datos son, semana de trabajo (i), cantidad de trabajadores (m), día de trabajo de lunes a domingo (día), jornada de trabajo (hora), cinco variables explicativas (x_1 a x_5 , como factores que influyen en las ventas) y las utilidades o ventas de la tienda (ventas). Es importante mencionar que las variables explicativas fueron creadas a partir de valores aleatorios dentro del rango (0,1) de manera uniforme, redondeados a dos decimales por simplicidad.

Esta instancia posee todas las combinaciones posibles de turnos laborales para cada día y horario, permitiendo una asignación de cero a siete trabajadores por bloque, es decir, $\forall((día, horario), m): m \in \{0,7\}$. Convirtiéndose en un caso ideal de datos donde para cada observación posible existe una etiqueta de venta, trabajando así en el caso de aprendizaje supervisado. Mientras que cualquier otro tipo de instancia que simule mejor una situación real de configuración de turnos laborales, es decir, donde no se conozcan todas las utilidades posibles para $\forall((día, horario), m)$, el problema se volvería un desafío de aprendizaje semi supervisado. Las ventas de la base de datos fueron estimadas a partir de una regresión lineal múltiple de la forma:

$$U_{h,m} = (\varepsilon + \beta_0 + \sum_{k=1}^5 \beta_k \cdot X_k + \beta_m \cdot \ln(m + 1))^{deg}$$

Donde $U_{h,m}$ se define como las ventas en el bloque h con m trabajadores, los β_k siendo coeficientes estimados de manera aleatoria uniforme (0,1), el logaritmo natural del número de trabajadores en h se implementa para que la regresión entregue ventas no tan lineales. Además, se eleva a un grado (deg) a la ecuación para que el modelo sea más flexible con respecto a la relación entre variables.

La instancia que se trabajará tiene datos simulados de un año (52 semanas), y cada bloque posee desde cero hasta siete trabajadores, es decir, 112 observaciones por semana. Con ello, el 80% de los datos corresponden a 41 semanas y conforman la instancia de entrenamiento, en tanto, la de testeo corresponde a 11 semanas.

En esta sección se comparan la calidad de las decisiones tomadas al aplicar los enfoques PFL y DFL. Como ya se describió, con la instancia del 80% de los datos se entrena un modelo de regresión lineal minimizando el error cuadrático de predicción. También se entrenó un modelo lineal con respecto a las funciones de pérdida Pointwise y SPO+. Posterior a entrenar los modelos con las metodologías PFL y DFL, se comparan los coeficientes del modelo resultante para estimar las ventas ($\hat{U}_{h,m}$) de cada uno de ellos. Respecto a PFL los coeficientes que resultan de entrenar la regresión lineal, se denominan β^{RL} . Mientras que para la metodología DFL, al entrenar respecto a la función de pérdida Pointwise, los coeficientes β estimados se denominan $\beta^{Pointwise}$ y para la función de pérdida SPO+ son β^{SPO+} , estos se muestran en la Tabla 1.

Coeficientes	RL	Pointwise	SPO+
β_0	-4671279.31	0.0	0.0
β_1	1811454.84	3381125.27	2142577.97
β_2	1639126.76	3809832.04	2121940.73
β_3	1624214.78	3446305.35	2604837.78
β_4	547688.85	1357243.42	752296.47
β_5	744290.58	1557094.40	1434169.66
β_6	53332.31	0.0	0.0
β_7	48378.78	0.0	0.0
β_8	37077.37	-138927299.7	0.0
β_m	2778501.20	9230992.77	7904045.06

Tabla 1. *Comparación de coeficientes betas (β^{RL} , $\beta^{Pointwise}$ y β^{SPO+}). Elaboración propia.*

Estos hallazgos sugieren que los coeficientes β estimados por las metodologías DFL de Pointwise y SPO+ difieren, pero poseen magnitudes relativamente similares. Sin embargo, existe un valor extremadamente negativo en Pointwise correspondiente a la variable dummy “weekend mañana” (β_8), lo que indicaría un valor atípico que puede afectar la calidad de la estimación. En

comparación con los parámetros β estimados por la regresión lineal en la metodología PFL, se observa que hay diferencias notables en la magnitud versus DFL. Además, se puede observar que para las metodologías de DFL, los coeficientes asociados a las variables dummies ($\beta_6, \beta_7, y \beta_8$) que están relacionadas con respecto al bloque h como: día de semana–horario mañana, día de semana–horario tarde, y fin de semana–horario mañana, respectivamente, en su mayoría son coeficientes nulos. Asimismo, sucede con β_0 por lo que se asume que estas variables explicativas dentro de la predicción de ventas son poco características.

Para el caso de entrenar el modelo bajo el enfoque PFL, se realizó un análisis de estadístico p–valor, que muestra que los coeficientes $\beta_0, \beta_1, \beta_2, \beta_3, \beta_4, \beta_5, \beta_m$ pueden ser significativos, ya que su p–valor son cercanos a cero. En cambio, el resto de los parámetros $\beta_6, \beta_7, \beta_8$ pueden ser insignificantes porque los p–valores de son mayores a 0,05. Estos valores se presentan en la Tabla 2.

Coeficiente	P_valor
β_0	0.000000e+00
β_1	4.948213e-131
β_2	1.397997e-108
β_3	7.163264e-107
β_4	3.154745e-14
β_5	5.180834e-25
β_6	4.108016e-01
β_7	4.556186e-01
β_8	6.323847e-01
β_m	0.000000e+00

Tabla 2. *P_valor de coeficientes estimados con metodología PFL. Elaboración propia.*

A continuación, se mide la calidad de los modelos observando las métricas de evaluación de MSE y el *regret* para la instancia de entrenamiento y de testeo.

I. Métricas de evaluación PFL

Para la metodología PFL se realizan histogramas de frecuencia para visualizar la distribución de los errores cuadráticos medios para el conjunto de entrenamiento y de testeo, el resultado se puede ver en la Figura 1.

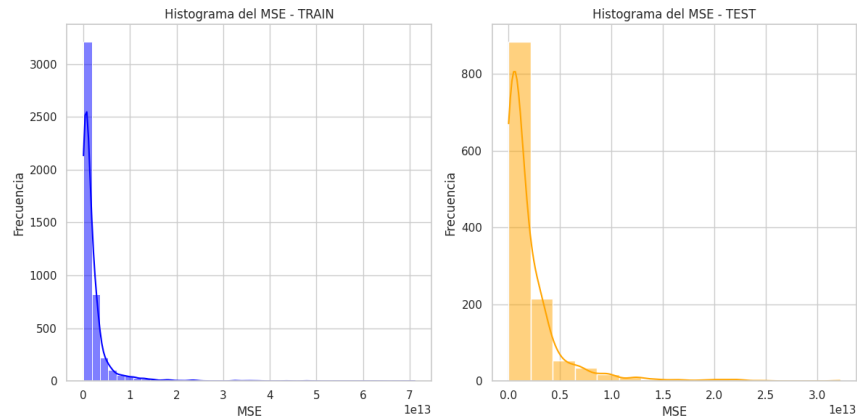


Figura 1. *Histogramas de frecuencia de MSE para datos de entrenamiento y testeo.*

Elaboración propia.

Se observa que el modelo comete errores tanto para la instancia de entrenamiento como en la de testeo, donde la mayor concentración está desde el rango de 0 a 3 en una escala $1e^{13}$.

Por otro lado, en la Figura 2 se muestra el *regret* de los datos de entrenamiento y testeo, respectivamente. Con el fin de realizar un posterior análisis comparativo con la metodología DFL.

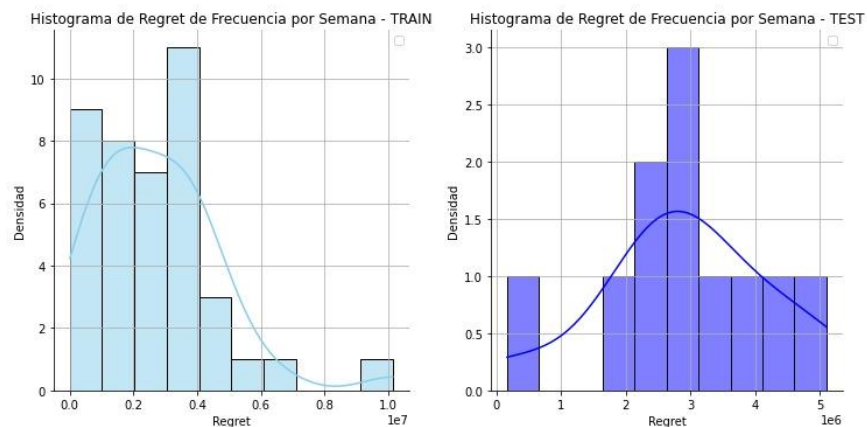


Figura 2. *Histograma de frecuencia del regret para datos de entrenamiento. Elaboración propia.*

Se puede observar en la Figura 2 que la pérdida de beneficio al estimar utilizando PFL están en un rango de miles de pesos a 10 millones. Con una concentración mayor desde los miles de pesos hasta los 4 millones.

II. Métricas de evaluación DFL

Se realiza comparación entre la metodología DFL y PFL. Donde los estimadores utilizando DFL corresponden a $\beta^{Pointwise}$ y β^{SPO+} , mientras que para PFL está el estimador de coeficientes β^{RL} . Dicha comparación se observa en la Figura 3.

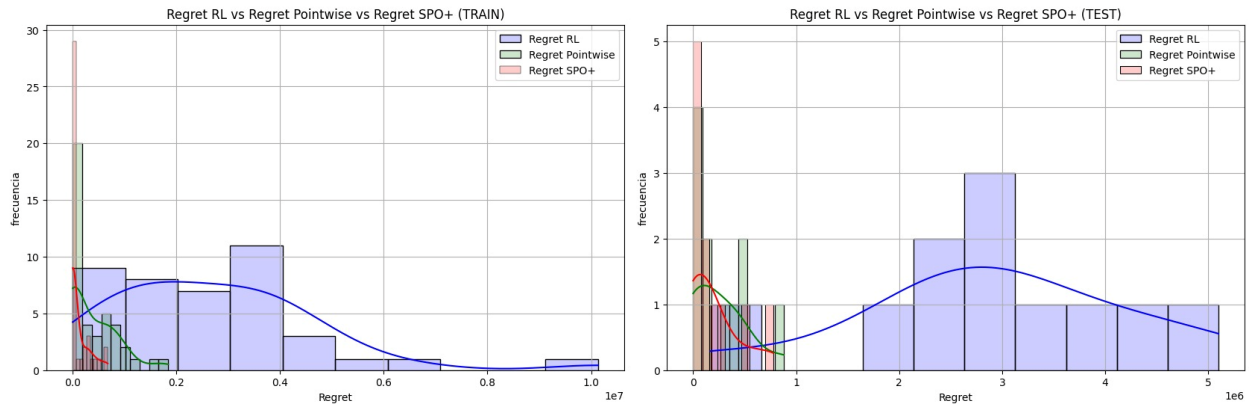


Figura 3. Comparación de regret para metodologías PFL y DFL para datos de entrenamiento y testeo. Elaboración propia.

De esta figura, se observa que el *regret* asociado al método SPO+ es mejor que el asociado al método Pointwise en términos de la metodología DFL, y ambos son mejores que el *regret* asociado a la metodología PFL. En términos asociados a los beneficios de la tienda, se observa que utilizando Pointwise versus la regresión lineal de PFL, se reducen las pérdidas económicas en aproximadamente 8 millones en el conjunto de entrenamiento, y cerca de 4 millones de pesos con el de testeo. Al aplicar el método SPO+ versus el PFL, se reducen las pérdidas en 9 millones en el conjunto de entrenamiento, y cerca de 4 millones en el de testeo.

En la Figura 4 y 5, se muestran histogramas de frecuencia de los errores cuadráticos medios de las ventas predichas en cuanto a las ventas reales de cada estimador entrenado por las funciones de pérdida. En el caso de la Figura 4 se observan los MSE para la metodología Pointwise para la muestra de entrenamiento y testeo, donde se ve una alta concentración de errores entre rango 0.0 a 0.25 en escala 1e16, sugiriendo que la mayoría de las predicciones tienen errores cuadráticos medios pequeños, y tienen una distribución consistente o bastante similar entre ambos conjuntos. En la Figura 5, se observan los MSE para la metodología SPO+, observando que para ambos conjuntos de datos existe una distribución de errores uniformes, ya

que no existen picos ni valles importantes hasta un MSE de 2.0 en escala de $1e14$. Sin embargo, se percibe un incremento en la concentración de errores cuadráticos medios desde 2.0 en escala $1e14$ en adelante, lo que podría significar que el modelo tiene un desempeño deficiente en ciertas predicciones.

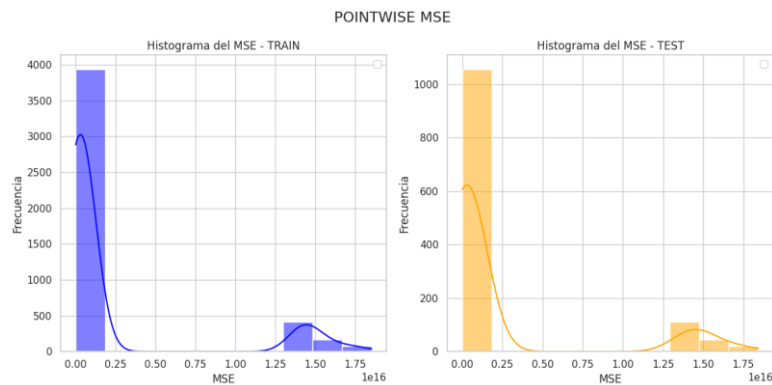


Figura 4. *MSE para metodología Pointwise, de datos de entrenamiento y testeo.*

Elaboración propia.

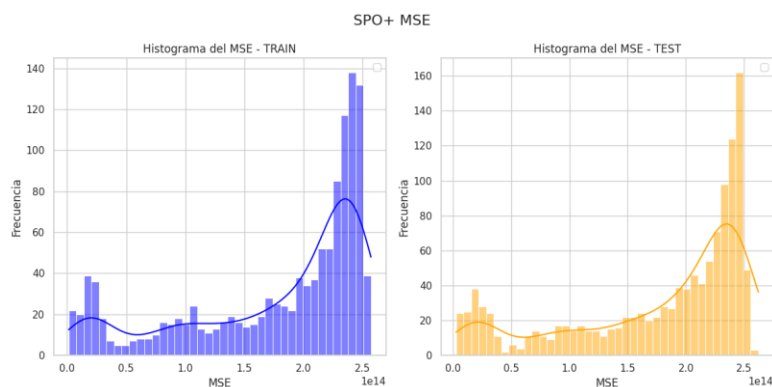


Figura 5. *MSE para metodología SPO+, de datos de entrenamiento y testeo. Elaboración propia.*

III. Análisis de robustez

Como siguiente paso, se quiere estudiar la robustez de los modelos aplicados en el estudio con el propósito de evaluar la estabilidad de los resultados al enfrentar variaciones en los datos. Por ende, se realizan las predicciones de ventas a partir de distintas combinaciones del modelo con el fin de ver si consistentemente el *regret* se comporta mejor para el método DFL versus el PFL. Para llevar a cabo el análisis de robustez para este estudio, se omiten ciertos coeficientes β para la estimación de las ventas. La primera combinación implica omitir

$\beta_0, \beta_1, \beta_6, \beta_7, \text{ y } \beta_8$ para la predicción, es decir, no se entrena el modelo con los atributos que explican dichos parámetros. Los tres últimos ($\beta_6, \beta_7, \text{ y } \beta_8$) corresponden a los coeficientes de las variables dummies. La segunda combinación es omitir solo coeficientes β_4 y β_5 .

Los resultados de aplicar la primera combinación se ven en la figura 6, y los de aplicar la segunda combinación se muestran en la figura 7.

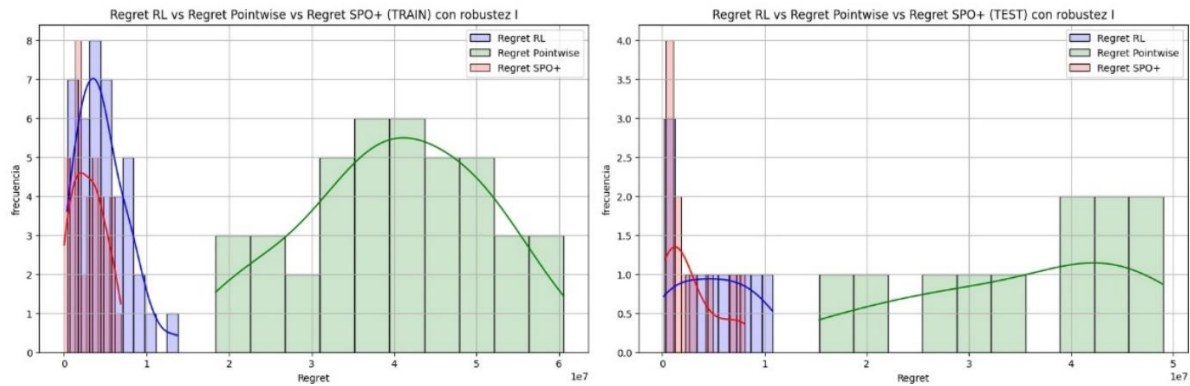


Figura 6. Comparación regret de PFL y DFL, con primera combinación. Elaboración propia.

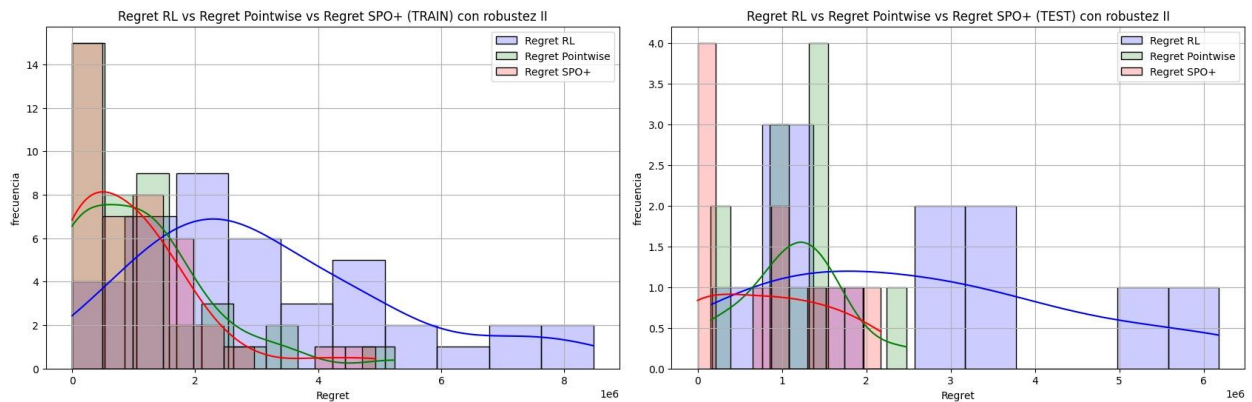


Figura 7. Comparación regret de PFL y DFL, con segunda combinación. Elaboración propia.

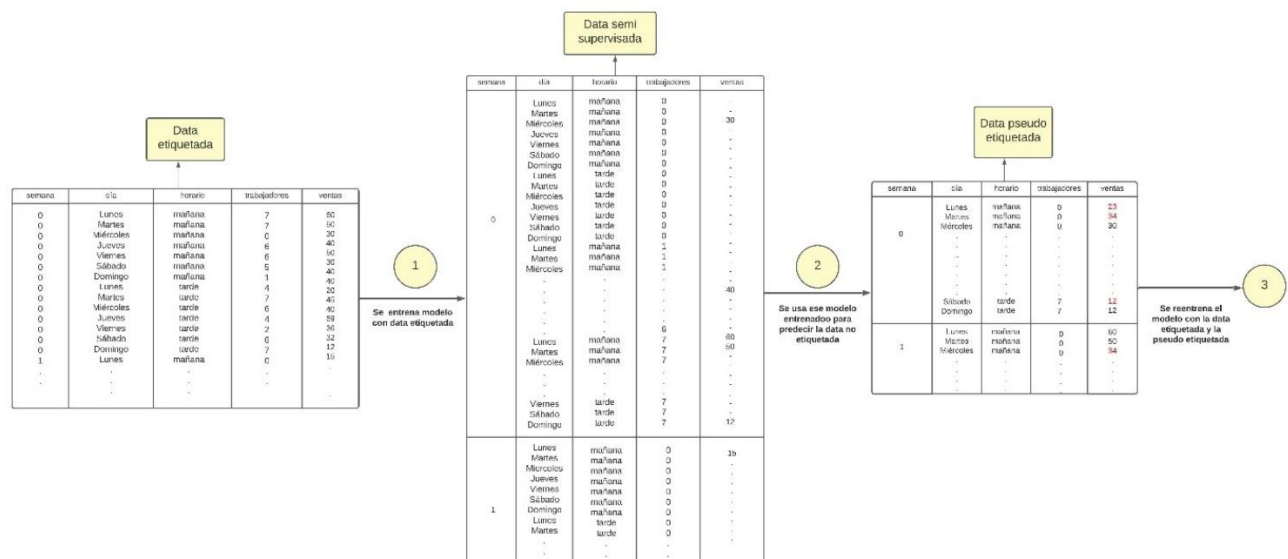
Se observa que para la primera y segunda combinación el método SPO+ es robusto, ya que el *regret* es mejor en términos económicos que el del método PFL. No obstante, el método Pointwise no es robusto en la Figura 6; esto puede deberse al hecho que se ocultan los coeficientes de las variables dummies y para este método el coeficiente β de la variable dummy

“weekend mañana” parece ser de interés para la predicción de ventas. En cambio, se puede ver que para la segunda combinación los métodos DFL si son consistentes en relación con que el *regret* obtiene menor pérdida de beneficio en comparación con el método PFL.

Alcances y limitaciones

Las limitaciones de este estudio recaen en cómo se obtuvieron los datos para la aplicación de metodologías DFL. En esta memoria, se creó una instancia para entrenar el modelo donde se conocen todas las ventas para todas las combinaciones de bloque h con la cantidad de trabajadores m . Sin embargo, en la vida real esto no es posible, ya que se puede conocer solo las combinaciones (h, m) existentes dentro de la semana, entendiendo esta combinación como $((\text{día}, \text{mañana/tarde}), \text{cantidad de personal})$. Por lo que estudiar otro caso donde solo se conozcan ciertas combinaciones se estaría convirtiendo en un método de aprendizaje semi supervisado.

El aprendizaje semi supervisado es un área de investigación relativamente joven, sin embargo, Van y Hoss (2020) proponen algoritmos para incorporar datos no etiquetados en el entrenamiento de un modelo. Mencionan que existen métodos para tratar este tipo de problemas de regresión. Uno de ellos son los métodos envolventes, que si bien son métodos que se han investigado poco, pueden aplicarse al entorno de regresión semi supervisada. Basado en esto, se propone un nuevo alcance para esta investigación en el cuál se tiene una instancia de datos para un aprendizaje semi supervisado. En el Esquema 2 se visualiza una manera de resolver este desafío.



Esquema 2. Esquema para resolver problema de dotación de personal por aprendizaje semi supervisado. Elaboración propia.

Este método envolvente consiste en entrenar primero la regresión con los datos etiquetados. Luego, utilizar las predicciones de ese modelo para generar datos etiquetados adicionales en las observaciones no etiquetadas. Finalmente, el modelo vuelve a entrenarse con los datos pseudo etiquetados además de los datos etiquetados ya existentes.

Para trabajos futuros respecto a la aplicación de Aprendizaje Enfocado en la Decisión a problemas de optimización de dotación de personal enfocado en utilizar aprendizaje de máquina semi supervisado, se recomienda probar otro tipo de modelo de aprendizaje. En el caso de este estudio, se alude a un modelo de regresión lineal para estimar un parámetro desconocido, sin embargo, podría estimarse por ejemplo, utilizando redes neuronales. Como mencionan Van y Hoss (2020), las investigaciones sobre el aprendizaje semi supervisado se ha enfocado más en este tipo de modelo, utilizando un método inductivo como el de preprocesamiento no supervisado.

Conclusión

En esta memoria se desarrolló un análisis comparativo de metodologías de aprendizaje distintas, PFL y DFL, para estimar el impacto de la dotación de personal en las ventas que es desconocido. Se aplicaron métodos que permitieron predecir de manera inteligente las ventas con el fin de poder obtener mejores decisiones al momento de optimizar los turnos laborales. De acuerdo con los resultados, se puede concluir que existe una mejora en términos económicos para una tienda al utilizar métodos de Aprendizaje Enfocado en la Decisión en comparación a utilizar métodos de Aprendizaje Enfocados en la Predicción. En este caso, se recomienda utilizar el estimador SPO+ ya que fue un enfoque robusto para todas las distribuciones en las ventas. Según este estudio, el enfoque DFL permite tener ganancias monetarias en un rango de 4 a 9 millones de pesos por sobre los enfoques PFL, en instancias sintéticas.

Por otro lado, se discutió la manera de aplicar este tipo de problema a un contexto de aprendizaje semi supervisado, con ello se aconseja utilizar un método inductivo de aprendizaje, específicamente el método de envoltura o pseudo etiquetado.

Es importante señalar que, aunque estos resultados sustentan la hipótesis de este estudio, existen limitaciones que es importante abordar en futuras investigaciones relacionadas con el tema. Algunas son la selección de los datos, que se limitó a simular una situación real perfecta en el contexto de turnos laborales y ventas de una tienda minorista. Otra limitación es el modelo de aprendizaje: si bien la regresión lineal es un modelo que funciona bien para predecir valores de salida al ingresar datos de entrada, esta metodología se puede extender a modelos de aprendizaje más complejos, tales como redes neuronales.

Esta memoria ha contribuido a resolver de manera certera un problema real al que se enfrenta la industria. Y da a conocer que aplicar buenos métodos para predecir parámetros desconocidos para la toma de decisiones influye económicamente en la gestión de una empresa.

Referencias

- Adam N. Elmachtoub, & Paul Grigas (2022). Smart “Predict, then Optimize”. *Management Science* 68(1):9–26. <https://doi.org/10.1287/mnsc.2020.3922>
- Christiansen Agar, M. (2020). *Desarrollo de un modelo de predicción de tráfico de clientes para optimizar la dotación de personal en tiendas físicas del Retail*. [Memoria para optar al título de Ingeniero Civil Industrial, Universidad de Chile] Disponible en <https://repositorio.uchile.cl/handle/2250/177936>
- Chuang, H.H.-C., Oliva, R. & Perdikaki, O. (2016). Traffic-Based Labor Planning in Retail Stores. *Prod Oper Manag*, 25: 96-113. <https://doi.org/10.1111/poms.12403>
- Dirección de presupuestos (2011). *Instrucciones para el envío de Informe de Dotación de personal: Personal de dotación*. Gobierno de Chile. http://www.dipres.gob.cl/598/articles-84135_doc_pdf.pdf
- Guerrero-Martínez, D.G. (2012). Factores clave de éxito en el negocio del retail. *Ingeniería Industrial*, 30(030), 189–205. <https://doi.org/10.26439/ing.ind2012.n030.223>
- Henao, C. A., Muñoz, J. C., & Ferrer, J. C. (2015). The impact of multi-skilling on personnel scheduling in the service sector: a retail industry case. *Journal of the Operational Research Society*, 66(12), 1949–1959. <https://doi.org/10.1057/jors.2015.9>
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An Introduction to Statistical Learning: With Applications in R*. Springer New York.
- Mandi, J., Bucarey, V., Tchomba, M.M.K. & Guns, T.. (2022). Decision-Focused Learning: Through the Lens of Learning to Rank. *Proceedings of the 39th International Conference on Machine Learning, in Proceedings of Machine Learning Research* 162:14935–14947 Available from <https://proceedings.mlr.press/v162/mandi22a.html>.
- 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. <https://doi.org/10.48550/arXiv.2307.13565>
- Miranda, J., y Palacios, F. (2017). Programación entera para el diseño de jornadas laborales de reponedores en la industria del retail. *Revista Ingeniería de Sistemas Volumen XXXI*. http://www.dii.uchile.cl/~ris/RIS2017/4_Jornadas_laborales_retail.pdf.
- Mohri, M., Rostamizadeh, A., & Talwalkar, A. (2018). *Foundations of Machine Learning*. Reino Unido: MIT Press.

- Nyman, E. (2022). *A sensitivity analysis approach to improve workforce optimization results by simulating employee training*. [Thesis submitted for examination for the degree of Master of Science in Technology, Aalto University School of Science]. <http://urn.fi/URN:NBN:fi:aalto-202202061783>
- Perdikaki, O., Kesavan, S., & Swaminathan, J. M. (2012). Effect of traffic on sales and conversion rates of retail stores. *Manufacturing & Service Operations Management*, 14(1), 145–162. <https://doi.org/10.1287/msom.1110.0356>
- Van Engelen, J.E. & Hoos, H.H. (2020). A survey on semi-supervised learning. *Mach Learn* 109, 373–440. <https://doi.org/10.1007/s10994-019-05855-6>
- Yung, D., Bucarey, V., & Olivares, M. (2020). Labor planning and shift scheduling in retail stores using customer traffic data. Available at SSRN 3695353.