



Master in Management + Analytics

Gestión de listas de espera para estudios endoscópicos en hospitales públicos: un abordaje mediante simulación y Machine Learning

Autoría: Bircher, Andrés Alfonso

Año: 2025

¿Cómo citar este trabajo?

Bircher, A. (2025) "*Gestión de listas de espera para estudios endoscópicos en hospitales públicos: un abordaje mediante simulación y Machine Learning*". [Tesis de maestría. Universidad Torcuato Di Tella]. Repositorio Digital Universidad Torcuato Di Tella <https://repositorio.utdt.edu/handle/20.500.13098/13666>

El presente documento se encuentra alojado en el **Repositorio Digital de la Universidad Torcuato Di Tella** bajo una licencia Creative Commons Atribución-No Comercial-Compartir Igual 4.0 Internacional
Dirección: <https://repositorio.utdt.edu>



UNIVERSIDAD
TORCUATO DI TELLA

Master in Management + Analytics

GESTIÓN DE LISTAS DE ESPERA PARA ESTUDIOS ENDOSCÓPICOS EN HOSPITALES PÚBLICOS: UN ABORDAJE MEDIANTE SIMULACIÓN Y MACHINE LEARNING.

Autor: Andrés Alfonso Bircher.

Tutor: Javier Marengo.

Mayo 2025

A mis padres, Analía y Andrés por su valentía

A mis abuelos, Ana Lía y Mingo por su corazón.

Agradecimientos

Deseo expresar mi profundo agradecimiento a la Dra. Sandra Arisio por su dedicación, acompañamiento y compromiso constante a lo largo de este trabajo. Su apoyo fue indispensable para que esta tesis fuera posible.

Agradezco también al Dr. Eduardo Marini por su generosidad y confianza al brindarnos acceso a su servicio y sumarse al proyecto con entusiasmo desde el inicio.

Finalmente, extendiendo mi reconocimiento a todo el equipo de la Unidad de Gastroenterología del Hospital General de Agudos J.M. Ramos Mejía, por su colaboración en la recolección de datos, la predisposición para responder consultas y el tiempo dedicado al seguimiento del proyecto.

Resumen

La gestión de la atención en sistemas hospitalarios resulta desafiante en múltiples aspectos. En los últimos años se vio un aumento en la demanda que no se pudo acompañar de un aumento proporcional en la capacidad operativa. Esta problemática tiene dos grandes aristas a abordar: la disponibilidad de recursos tecnológicos y la cantidad de personal disponible. En países donde la salud pública representa uno de los principales sistemas de atención esta situación cobra especial relevancia.

El Hospital General de Agudos José María Ramos Mejía de la Ciudad Autónoma de Buenos Aires cuenta con una Unidad de Gastroenterología que brinda atención ambulatoria y a pacientes internados, realizando estudios endoscópicos tanto diagnósticos como terapéuticos. Dichos estudios son de vital importancia en la prevención y detección temprana de múltiples patologías. Debido a distintos agentes externos, entre ellos la pandemia de COVID-19, la lista de espera para la realización de estudios en el quirófano ha aumentado de manera exponencial. Esto conlleva una dificultad adicional para los profesionales dado que, a medida que aumenta la demanda, más complejo resulta establecer las prioridades en los turnos que se asignan.

En este trabajo de Tesis, se realizó una asignación de prioridades para la realización de estudios endoscópicos mediante un abordaje con simulación y Machine Learning. Para ello se utilizaron datos médicos obtenidos en consultas y se seleccionaron los features más relevantes. Además, a través del análisis de series temporales logramos definir cuáles son los requisitos mínimos para que el servicio brindado sea óptimo. Por último, se definió una métrica que permite evaluar si el tiempo de respuesta de la institución es acorde a la patología de los pacientes para distintas situaciones operativas.

Abstract

Managing patient care in hospital systems presents multiple challenges. In recent years, there has been an increase in demand that has not been matched by a proportional increase in operational capacity. This issue has two major aspects to address: the availability of technological resources and the number of available staff. In countries where public healthcare is one of the main systems of care, this situation becomes particularly relevant.

The José María Ramos Mejía General Hospital in the Autonomous City of Buenos Aires has a Gastroenterology Unit that provides both outpatient and inpatient care, performing diagnostic and therapeutic endoscopic procedures. These procedures are of vital importance for the prevention and early detection of various pathologies. Due to several external factors, including the COVID-19 pandemic, the waiting list for procedures requiring an operating room has increased exponentially. This poses an additional challenge for healthcare professionals, as rising demand makes it increasingly difficult to establish priorities when assigning appointments.

In this thesis project, priority assignments for endoscopic procedures were carried out using a machine learning approach. Medical data obtained during consultations were used, and the most relevant features were selected. Furthermore, a metric was defined to evaluate whether the institution's response time aligns with the severity of patients' conditions under different operational scenarios. Additionally, through time series analysis, the minimum requirements for optimal service delivery were predicted.

Índice

1.	Introducción	7
1.1	Priorización de pacientes en el mundo	8
1.2	Servicio de Gastroenterología en el Hospital J.M. Ramos Mejía	11
1.3	Demanda del servicio: Complicaciones para priorizar	12
1.4	Objetivos	12
1.5	Desarrollo general	13
2.	Datos	14
2.1	Dataset principal	14
2.2	Dataset secundarios	15
2.3	Análisis Exploratorio	17
3.	Metodología	27
3.1	División del problema	27
3.2	Score mediante Machine Learning	28
3.3	Gestión de la nueva demanda	30
4.	Resultados	33
4.1	Lista de espera	33
4.1	Priorización	36
4.2	Demanda: Teoría de Colas	40
4.3	Simulación	41
4.4	Nivel de Servicio	43
5.	Discusión	45
5.1	Primer escenario: Quirófano largo	45
6.	Conclusiones	49
6.1	Cuando la IA Decide: La Importancia de la Ética en Salud	49
6.2	Recomendaciones	50
7.	Bibliografía	52
8.	Anexo	54

1. Introducción

La simulación de datos es el proceso de creación de modelos que permite imitar escenarios del mundo real para entender su comportamiento, analizar y predecir resultados. Este puede aplicarse para el estudio de distintos campos permitiendo optimizar recursos limitados, planificar y tomar decisiones basadas en datos.

En los últimos años los servicios de atención en salud se ven afectados por factores que determinan un aumento en la demanda de atención y factores que limitan la oferta. Este fenómeno es global y transversal a los sistemas de cobertura, si bien se evidencia una disparidad en detrimento de aquellos sectores con coberturas dependientes del estado. En este contexto la pandemia por COVID-19 exacerbó una problemática ya presente (Organización Mundial de la Salud, 2024; Organización Panamericana de la Salud, 2020).

Los estudios endoscópicos, problemática que abordaremos en esta tesis, en su mayoría son de carácter programado y preventivo. La video colonoscopia (VCC) es el estándar de oro para la prevención del cáncer colorrectal. Esta patología, a nivel mundial, es la 3° neoplasia en incidencia y la 2° en mortalidad. A nivel nacional representa la 2° neoplasia tanto en incidencia como en mortalidad siendo Argentina un país con incidencia creciente y mortalidad estable en el tiempo (Kirschbaum & Yonamine, 2020; Organización Mundial de la Salud, 2023; Yonamine & Kirschbaum, 2022).

Existen diversos factores que dificultan el acceso a estos procedimientos como la cobertura en salud, el grado de instrucción, el nivel socioeconómico, etc (Báscolo et al., 2020; Kaarboe & Carlsen, 2014; Shim et al., 2024).

Durante la pandemia los estudios electivos en general fueron altamente limitados por factores como redistribución de personal y transformación de unidades de atención hacia la emergencia. Esto representó un aumento exponencial en la demanda insatisfecha y en el desarrollo de patologías prevenibles en la población. Asimismo, tiempos de espera excesivos con listas voluminosas. En este contexto una priorización eficiente de los pacientes resulta clave para dar un servicio de calidad. (Cámara Argentina de Especialidades Medicinales, 2024; Ministerio de Salud de Argentina, 2022; Shim et al., 2024).

Los algoritmos de Machine Learning (ML) son capaces de identificar patrones permitiendo clasificar a los pacientes según la urgencia de atención que requieren. Por ejemplo, a partir de datos como los síntomas, el historial médico o los resultados de pruebas, el ML podría identificar aquellos pacientes con mayor riesgo de complicaciones, ayudando a los médicos a priorizar los casos más críticos (SimPy Documentation Release 4.0.0 Team SimPy, 2020)

Es por lo previamente expuesto que la simulación (como herramienta fundamental) y el Machine Learning (como soporte a la decisión médica) se presentan

como potenciales valiosas herramientas para la gestión permitiendo colaborar con la priorización y brindando una estrategia para entender cómo licuar la demanda insatisfecha y gestionarla eficientemente.

1.1 Priorización de pacientes en el mundo

La asignación de prioridades a los pacientes para la realización de intervenciones quirúrgicas o estudios médicos es un tema común a las diversas instituciones médicas en todo el mundo. No existe un criterio unificado ni a nivel global, regional ni local de cómo abordar esta problemática.

El Reino Unido tiene un sistema de salud que es de carácter universal y público: el NHS (National Health Service). En este país el NHS no implementa una política formal de priorización, los hospitales tienen la libertad de gestionar las listas de espera como lo consideren adecuado, lo que da lugar a variaciones locales en las políticas de admisión. En algunos casos, las listas de espera pueden gestionarse según una regla de "primero en llegar, primero en atenderse" o FIFO (First In First Out), mientras que, en otros, la gestión puede tener en cuenta la patología y su gravedad. Por ejemplo, dentro de las metodologías para medir el nivel de urgencia o de dolor de los pacientes en espera para cirugías traumatológicas se crearon cuestionarios que evalúan el dolor y la función articular, clasificando a los pacientes de acuerdo con los puntajes obtenidos en 11 bandas de prioridad. (Gutacker et al., 2016)

Si bien no hay una directiva de priorización si están sujetos a targets de máximos tiempos de espera y son evaluados regularmente en cuanto a su desempeño, pudiendo ver hasta un 5% de sus ingresos retenidos si no cumplen con los objetivos (Gutacker et al., 2016).

En términos de limitaciones de acceso a la atención médica especializada en Noruega y Australia se demostró que los hombres con mayores ingresos y las mujeres con mayor nivel educativo esperan menos tiempo para una cirugía de reemplazo de cadera que aquellos de bajos ingresos (Kaarboe & Carlsen, 2014).

El Royal Adelaide Hospital de Australia en el año 2020 se vio ante la necesidad de desarrollar estrategias de priorización dado el aumento en la lista de espera de pacientes para colonoscopias. Se plantearon tres intervenciones claves: i) ordenar y limpiar la lista revisando los criterios de indicación según las guías de prácticas clínicas. Los pacientes que se consideraban tenían indicaciones inadecuados eran retirados; ii) asignación de un encargado de realizar seguimiento pre-intervención para disminuir la tasa de no show de los pacientes y evitar la pérdida de turnos endoscópicos; iii) aumento en la financiación que permitió agregar un día endoscópico para aumentar el volumen de estudios realizados (Sammour et al., 2021).

En Canadá, en el año 2006, se llevó a cabo un trabajo que buscó definir los tiempos de espera máximos entre la aparición de un síntoma gastrointestinal y la consulta con el especialista y/o el procedimiento endoscópico correspondiente. Este

trabajo surge del creciente tiempo de espera que debían atravesar los pacientes atendidos en el sector público: el 32% de los pacientes consideraba que los tiempos de espera para atención eran muy largos (media de 8.5 semanas), mientras que el tiempo entre indicación de un estudio y su realización era inaceptable (media de 17.5 semanas). El grupo de trabajo definió 6 grandes grupos de patologías: 1) sangrados digestivos, 2) sospecha de patología neoplásica, prevención y seguimiento, 3) sospecha de enfermedad inflamatoria intestinal, 4) dolor abdominal y alteración del tránsito intestinal, 5) disfagia y dispepsia, 6) enfermedades bilio-pancreáticas, y se estimaron cuatro tiempos de espera máximos: 1) Dentro de 24hs, 2) Dentro de las dos semanas, 3) dentro de los dos meses, 4) dentro de los seis meses. A cada patología se le asignó un tiempo de espera máximo de acuerdo con el apoyo de revisiones bibliográficas y la opinión de expertos (Ver tabla 1) (Paterson et al., 2006).

El trabajo mencionado nace de la necesidad de crear benchmarks de tiempos de espera que permitan unificar las conductas, sin consistir estos en una guía de práctica clínica sino en una recomendación para la gestión (Paterson et al., 2006).

En el 2024, en USA, se llevó a cabo un trabajo que evaluó el impacto en los tiempos de espera para colonoscopias de dos factores: la pandemia por COVID-19 y la presencia o no de cobertura en salud de los pacientes. Tomaron como referencia de tiempos de espera máximos los expresados en el trabajo canadiense previamente mencionado aquí. Se evidencio que en el periodo post pandemia los tiempos de espera fueron superiores al periodo prepandémico y que excedían los límites máximos. Los pacientes sin cobertura esperaban más que aquellos con algún tipo de seguro de salud. Por otro lado, desglosando las variables que impactan en los tiempos de atención encontraron que la falta de personal tanto medico como técnico era insuficiente para el aumento de la demanda y que los pacientes sin seguro de salud constituían grupos con barreras financieras y lingüísticas, por lo que era de utilidad la inversión en programas de navegación para asistir a estos pacientes (Shim et al., 2024)

Tabla 1. Tiempos recomendados de atención. Adaptado al español de “*Canadian consensus on medically acceptable wait times for digestive health care*”.

Tiempo máximo recomendado	Condiciones clínicas
Dentro de 24 horas	Sangrado gastrointestinal agudo
Dentro de 24 horas	Obstrucción por bolo alimenticio o cuerpo extraño en esófago
Dentro de 24 horas	Signos clínicos de colangitis ascendente
Dentro de 24 horas	Pancreatitis aguda grave (colangiopancreatografía retrógrada endoscópica dentro de 72 h, si está indicada)
Dentro de 24 horas	Enfermedad hepática descompensada grave
Dentro de 24 horas	Hepatitis aguda grave
Dentro de dos semanas	Alta probabilidad de cáncer según imagenología o examen físico
Dentro de dos semanas	Ictericia obstructiva aguda e indolora
Dentro de dos semanas	Disfagia u odinofagia grave y/o de progresión rápida
Dentro de dos semanas	Signos clínicos sugestivos de enfermedad inflamatoria intestinal activa
Dentro de dos meses	Sangrado rectal con sangre roja brillante
Dentro de dos meses	Anemia por deficiencia de hierro documentada
Dentro de dos meses	Uno o más test de sangre oculta en heces positivos
Dentro de dos meses	Hepatitis viral crónica
Dentro de dos meses	Disfagia estable (no grave)
Dentro de dos meses	Reflujo o dispepsia mal controlados
Dentro de dos meses	Estreñimiento o diarrea crónica
Dentro de dos meses	Cambio reciente en el hábito intestinal
Dentro de dos meses	Dolor abdominal crónico inexplicado
Dentro de dos meses	Confirmación de diagnóstico de enfermedad celíaca (test de anticuerpos)
Dentro de seis meses	Enfermedad por reflujo gastroesofágico crónica para cribado por endoscopia
Dentro de seis meses	Colonoscopia de cribado
Dentro de seis meses	Pruebas hepáticas anormales persistentes (más de seis meses) sin causa aparente

1.2 Servicio de Gastroenterología en el Hospital J.M. Ramos Mejía

Para el desarrollo de esta Tesis se contó con información proveniente de la Unidad de Gastroenterología del Hospital General de Agudos José María Ramos Mejía, perteneciente al Gobierno de la Ciudad Autónoma de Buenos Aires (Argentina). Dicho hospital se fundó en el año 1883 y recibe alrededor de 40.000 pacientes por mes.

La Unidad de Gastroenterología cuenta actualmente con 8 médicos de planta, un jefe de Sección de Endoscopia y 6 residentes, todos ellos dirigidos por un jefe de Unidad. El mismo funciona de Lunes a Viernes, de 7:30 a 17:00 horas, brindando atención ambulatoria (en consultorios externos) y en internación.

Además de la atención clínica realiza estudios endoscópicos, tanto diagnósticos como terapéuticos, en el quirófano mediante anestesia local o sedación. Estos estudios permiten la detección y evaluación de numerosas patologías.

Existen dos tipos fundamentales de estudios endoscópicos: Video Endoscopia Digestiva Alta (VEDA), que estudia el tracto gastrointestinal superior, y Video colonoscopia (VCC), que estudia el tracto gastrointestinal inferior. Dada la alta demanda de estos estudios y los escasos recursos hospitalarios, es común que se tenga una lista de espera como se mencionó en la sección anterior.

Múltiples factores determinan la capacidad de realización de estudios endoscópicos: el recurso humano (tanto endoscopistas como anestesistas), el recurso instrumental (torres y cañas de endoscopia) y la disponibilidad de una sala de procedimientos (en el caso del hospital: quirófano).

Un equipo de endoscopia se conforma de un procesador de imagen, una fuente de energía y una pantalla, esto se conoce como “torre”. Por otro lado, la caña es la parte más flexible y móvil del endoscopio, generalmente el extremo distal. Esta contiene el canal de aspiración o trabajo, así como la fibra óptica, las lentes y los chips que transmiten la luz y la imagen al procesador (Ver Anexo) (Cotton & Williams, 2003).

En los últimos tres años el servicio ha ido reforzando su recurso humano con mayor contratación de médicos especialistas en endoscopia y ha adquirido más cantidad de equipos. Actualmente de los ocho médicos de staff cuatro son endoscopistas, más el jefe de sección y de unidad. En términos de equipos cuentan con dos torres y cinco cañas.

La limitante actual en la capacidad productiva se encuentra fundamentalmente en la disponibilidad de quirófanos y de anestesistas. Este fenómeno fue originándose en los últimos años, posterior a la pandemia, lo que llevó al servicio a sufrir una baja en la capacidad productiva, en términos de estudios, resultando en la formación paulatina de una lista de espera.

Desde 2022 hasta diciembre de 2024 la lista de espera fue acumulando 800 pacientes. A esta demanda se suman los pacientes internados que en el año 2024

representaron el 43% de la demanda atendida. De los 795 estudios endoscópicos realizados, 242 correspondieron a pacientes internados.

1.3 Demanda del servicio: Complicaciones para priorizar.

El equipo de residentes es el encargado de realizar las listas quirúrgicas y contactar a los pacientes. Todas las semanas se arma un plan de quirófano, se contacta a los pacientes para indicarles el día y horario del estudio y cuáles son los pasos para la preparación. En la actualidad el servicio cuenta con dos turnos quirúrgicos de 4 hs aproximadamente los días Martes y Viernes de 7:30 hs - 12:00 hs.

Debido al gran volumen y variedad de casos, resulta difícil asignar las prioridades de atención en los pacientes en función de la urgencia del estudio y el tiempo de espera que llevan. Adicionalmente, esta lista podría estar sobredimensionada debido a la existencia de i) pacientes anotados de forma duplicada, ii) pacientes que ya se hayan realizado los estudios en otro centro, iii) pacientes que se hayan realizado el estudio en el mismo centro pero que, dada la ausencia de digitalización previa de esta lista y de los registros endoscópicos, no se han removido de forma correcta de la misma, iv) pacientes cuyos estudios estén indicados de forma inadecuada.

Estos factores aumentan el volumen de pacientes de forma imprecisa por lo que resulta necesaria una actualización de la lista original previa a la priorización de turnos.

A la fecha, el ordenamiento de los pacientes se realiza de acuerdo con la urgencia de estos entendiendo que existe un cuello de botella agravado por la pandemia como se expuso anteriormente y sin tener en claro cuál es el servicio mínimo requerido para regular la relación oferta y demanda de manera eficiente para la óptima atención de los pacientes. De aquí es que surge la necesidad tanto de encontrar un método para priorizar, como un plan para poder satisfacer la demanda.

1.4 Objetivos

El objetivo general de este trabajo de tesis es desarrollar una herramienta que permita la gestión de los recursos del Hospital General de Agudos José María Ramos Mejía, para la correcta priorización de los pacientes en espera de estudios endoscópicos. Además, evaluar los recursos necesarios que debieran asignarse para que dichos estudios se realicen en forma sostenible (sin la aparición de colas) y apropiada según la patología de los pacientes.

Para cumplir con este objetivo se establecen los siguientes objetivos específicos:

- Desarrollar un algoritmo de ML que contribuya priorizar los pacientes y desarrollar un ranking de atención.
- Utilizar la “teoría de colas” para realizar simulaciones que permitan determinar los recursos mínimos necesarios para la sostenibilidad del servicio, contribuyendo así a la gestión eficiente de los recursos hospitalarios.

En este objetivo se tendrá como parámetro a maximizar al nivel de servicio que se corresponde con el porcentaje de pacientes que se atenderán dentro de los tiempos recomendados por los profesionales, sobre el total de pacientes de la lista.

1.5 Desarrollo general

El servicio de gastroenterología del Hospital General de Agudos José María Ramos Mejía ha disponibilizado información sobre los pacientes que se han anotado en la lista de espera en los últimos años.

Inicialmente realizaremos un análisis descriptivo de la información brindada y analizaremos las series de tiempo para entender la demanda del servicio con respecto a endoscopías ambulatorias.

Posteriormente procederemos a definir un método para la priorización. Para ello nos basaremos en los tiempos de espera máximo mencionados en la sección anterior y en herramientas de Machine Learning.

Una vez abogado eso, retomaremos el análisis exploratorio de las series temporales de la demanda para definir un modelo de Teoría de Colas apropiado para el hospital, definiremos los parámetros de entrada y salida, calcularemos los recursos necesarios para sostenerlo en el tiempo y realizaremos una simulación para entender cómo se comporta la lista de espera a través del tiempo.

Realizada la simulación (la cual registrará momento de entrada y salida del paciente) podremos medir la calidad de nuestro servicio. Definiremos esto bajo la métrica “%On Time”, la cual será resultado del cociente entre cantidad de pacientes atendidos dentro de los tiempos recomendados y el total de los pacientes.

2. Datos

2.1 Dataset principal

Durante los últimos años, el equipo del servicio de residentes de Gastroenterología liderado por la Dra. Sandra Arisio ha estado trabajando en el registro digital de los pacientes para facilitar la gestión.

Este registro, realizado en Excel primero y luego en Google Sheets, registra a aquellos pacientes que requieren algún estudio endoscópico. Los tipos de estudio refieren principalmente a Video Endoscopia Digestiva Alta (VEDA), video colonoscopia (VCC) o Doble cuando el paciente necesita realizarse ambos.

La información suministrada por el Hospital describe edad de los pacientes, los motivos por los cuales se requiere el estudio solicitado y los estudios complementarios requeridos previos a la intervención.

La población registrada es diversa en cuanto a las edades y patologías. Se dispone de esta información desde enero de 2022 hasta diciembre del año 2024. Como se mencionó en el capítulo anterior, esta lista podría estar sobredimensionada y necesita actualización (Ver Sección 1.2). Sin embargo, a los fines del desarrollo de esta Tesis, la misma será considerada sin modificaciones.

Un aspecto importante acerca del trabajo con datos clínicos es la confidencialidad de estos por parte de la institución. Por tanto, para preservar y cuidar la identidad de los pacientes, los datos compartidos por el hospital han sido 100% anónimos. El dataset usa como referencia un “ID” incremental que nos sirve para identificar al paciente, en caso de necesitar información adicional.

Los descriptores relevados para la posterior clasificación de los pacientes fueron:

- Fecha inscripción: Fecha en la que se anotó a la lista.
- ID: Identificador único del paciente.
- EDAD: Edad del paciente
- ESTUDIO: Tipo de estudio a realizarse (VEDA, VCC o DOBLE).
- MOTIVO 1: Motivo principal por el cual se le solicita el estudio.
- MOTIVO 2: Motivo secundario por el cual se le solicita el estudio.
- DOMICILIO: Localidad de residencia.
- OBRA SOCIAL: Dato de si posee o no cobertura privada/ obra social.
- SOLICITANTE: Profesional que solicita el estudio.
- FECHA PEDIDO: Fecha en la que se pidió.
- ASA: Evaluación del anestesista para determinar el riesgo de la intervención según el paciente (I, II, III) siendo III el riesgo más alto.
- CARDIOLOGICO: Evaluación del cardiólogo para determinar el riesgo de la intervención.

2.2 Dataset secundarios

2.2.1 Estandarización de los diagnósticos

En primera instancia estandarizamos los diagnósticos más relevantes para el pedido del estudio: columna “MOTIVO 1” del dataset. Alineamos este motivo en palabras claves que le permitan al cuerpo médico asociar rápidamente el principal diagnóstico, eliminando duplicados o errores de tipo, ya que este es un campo de texto libre.

Posteriormente agrupamos los diagnósticos de las columnas “MOTIVO 1” en las 23 patologías principales descritas, asignando un código numérico (Ver Tabla 2) y sus tiempos de espera máximos recomendados en el estudio canadiense mencionado previamente (Ver Sección 1.1)

Este último punto relacionado a los tiempos es fundamental para el desarrollo de esta Tesis, dado que lo utilizaremos para obtener el ranking de pacientes más adelante.

Tabla 2. Códigos diagnósticos asociados a tiempos recomendados de espera.

Código	Descripción de la Patología	Tiempo recomendado [días]
1	Sangrado gastrointestinal agudo	1
2	Obstrucción esofágica por bolo alimenticio o cuerpo extraño	1
3	Signos clínicos de colangitis ascendente	1
4	Pancreatitis aguda grave (colangiopancreatografía retrógrada endoscópica dentro de las 72 h, si está indicada)	1
5	Enfermedad hepática descompensada grave	1
6	Hepatitis aguda grave	1
7	Alta probabilidad de cáncer según imágenes o examen físico	14
8	Ictericia obstructiva aguda sin dolor	14
9	Disfagia u odinofagia grave y/o de rápida progresión	14
10	Signos clínicos sugestivos de enfermedad inflamatoria intestinal activa	14
11	Sangrado rectal rojo brillante	60
12	Anemia por deficiencia de hierro documentada	60
13	Uno o más test de sangre oculta en materia fecal positivos	60
14	Hepatitis viral crónica	60
15	Disfagia estable (no grave)	60
16	Reflujo/dispepsia mal controlados	60
17	Estreñimiento o diarrea crónica	60
18	Cambio reciente en el hábito intestinal	60
19	Dolor abdominal crónico sin causa aparente	60
20	Confirmación diagnóstica de enfermedad celíaca (test de anticuerpos)	60
21	Enfermedad por reflujo gastroesofágico crónica para endoscopia de cribado	180
22	Colonoscopia de cribado	180
23	Enzimas hepáticas anormales persistentes (más de seis meses) sin explicación	180

2.2.2 Score

El score es una variable cuantitativa que ha construido el servicio de Gastroenterología con el objetivo de dimensionar, mediante el criterio de los profesionales, cuál es la urgencia del paciente a ser atendido.

Tomando cada patología como un Cluster de pacientes, el objetivo del Score es poder discriminar la urgencia Intra-Cluster.

En la Tabla 3 se presentan algunos ejemplos para el caso de Disfagia severa (código 9).

Tabla 3. Score de código 9.

id	edad	estudio	motivo 1	motivo 2	código paper	score
122	56	DOBLE	Disfagia severa	SOMF+	9	2
286	29	DOBLE	Disfagia severa	APF CCR	9	1
340	41	VEDA	Disfagia severa		9	2
344	43	DOBLE	Disfagia severa	Constipación	9	2
519	65	VEDA	Disfagia severa		9	3
767	67	DOBLE	Disfagia severa	Enfermedad diverticular	9	3

Para la construcción del score se tomaron en cuenta criterios de edad y “motivo 2” que influenciaran el pronóstico o grado de alarma del diagnóstico principal.

2.3 Análisis Exploratorio

Se realizó un análisis detallado de cada una de las variables del dataset con el objetivo de comprender la distribución de los datos disponibles.

El conjunto de datos incluye variables categóricas, como la información demográfica, de estudios y evaluaciones prequirúrgicas. También incluye variables numéricas, como la edad y el score.

Adicionalmente, se evaluaron series temporales asociadas a la demanda de estudios, con el fin de entender su comportamiento a lo largo del tiempo y generar estimaciones que permitan proyectar su evolución futura (Bruce et al., 2020; Mckinney, n.d.)

2.3.1 Variables Categóricas

- **ESTUDIO:** La distribución de estudios requeridos se presentan en la Figura 1 donde se observa que más del 50% de la demanda de estudios solicitados son dobles. Esto implica más carga horaria de quirófano.
- **MOTIVO 1:** Se define como el motivo principal por el cual se indica el estudio. Aquí se encuentra una variedad de diagnósticos con un patrón común asociado similar a la Ley de Pareto¹, donde las últimas 5 patologías representan más del 60% de la muestra (Ver Figura 2).
- **MOTIVO 2:** Es un campo opcional que puede añadir más información sobre la sospecha diagnóstica de los pacientes (Ver Figura 3).
- **OBRA SOCIAL:** Más del 60% de la población no cuenta con servicio de obra social para poder realizar el estudio por fuera de la salud pública. Esto es un indicador del grado de vulnerabilidad de las personas integrantes de la lista (Ver Figura 4).
- **ASA (Riesgo Anestésico):** El riesgo anestésico es bajo para la mayor parte de la población de esta muestra. Más del 90% de los pacientes que requieren estas intervenciones están contemplados entre los riesgos I y II (Ver Figura 5)
- **CARDIOLÓGICO:** El cardiológico y anestésico se encuentran directamente relacionados. Por lo tanto, es de esperar que el riesgo encontrado en este aspecto sea bajo para la mayor parte de la población (Ver Figura 6) ("Consenso Argentino de evaluación de riesgo cardiovascular en cirugía no cardíaca," 2016)

¹ El 80% de los resultados proviene del 20% de las causas o esfuerzos.

Figura 1. Distribución porcentual de estudios.

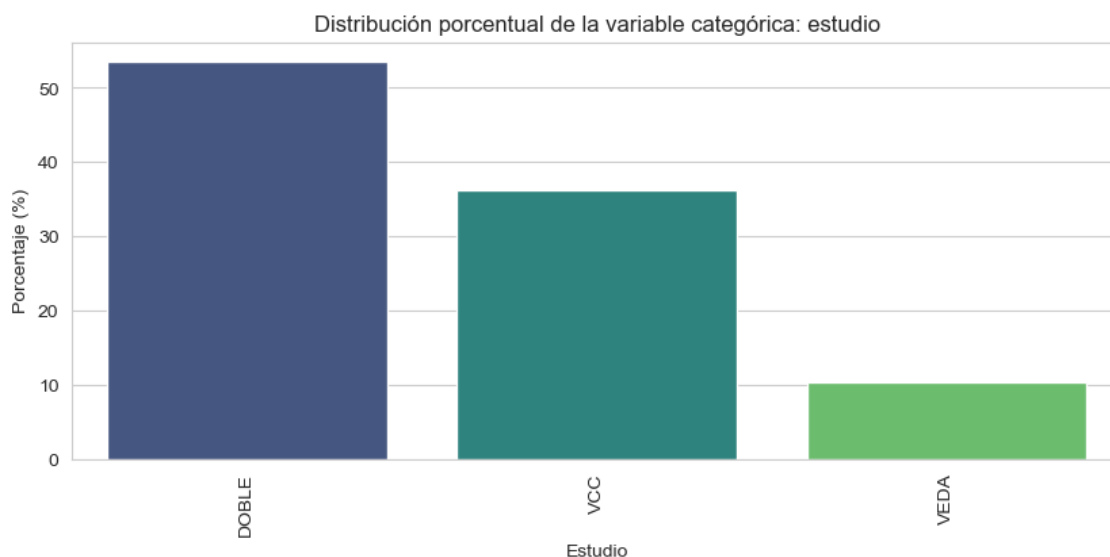


Figura 2. Distribución porcentual de “Motivo 1”.

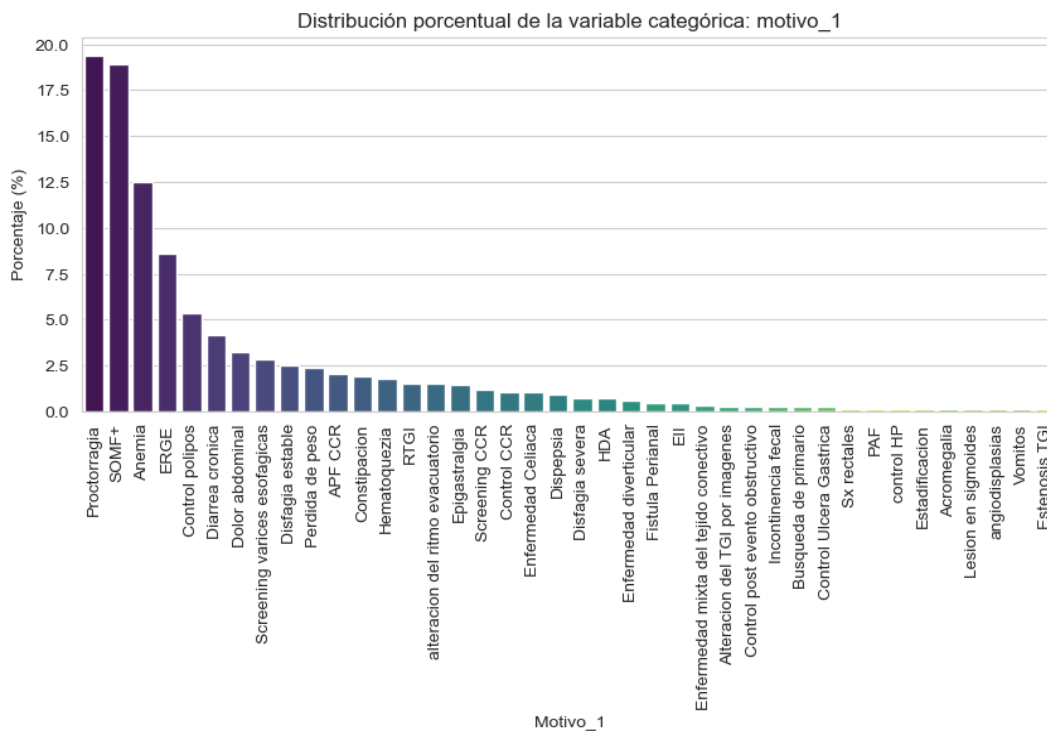


Figura 3. Distribución porcentual de “Motivo 2”.

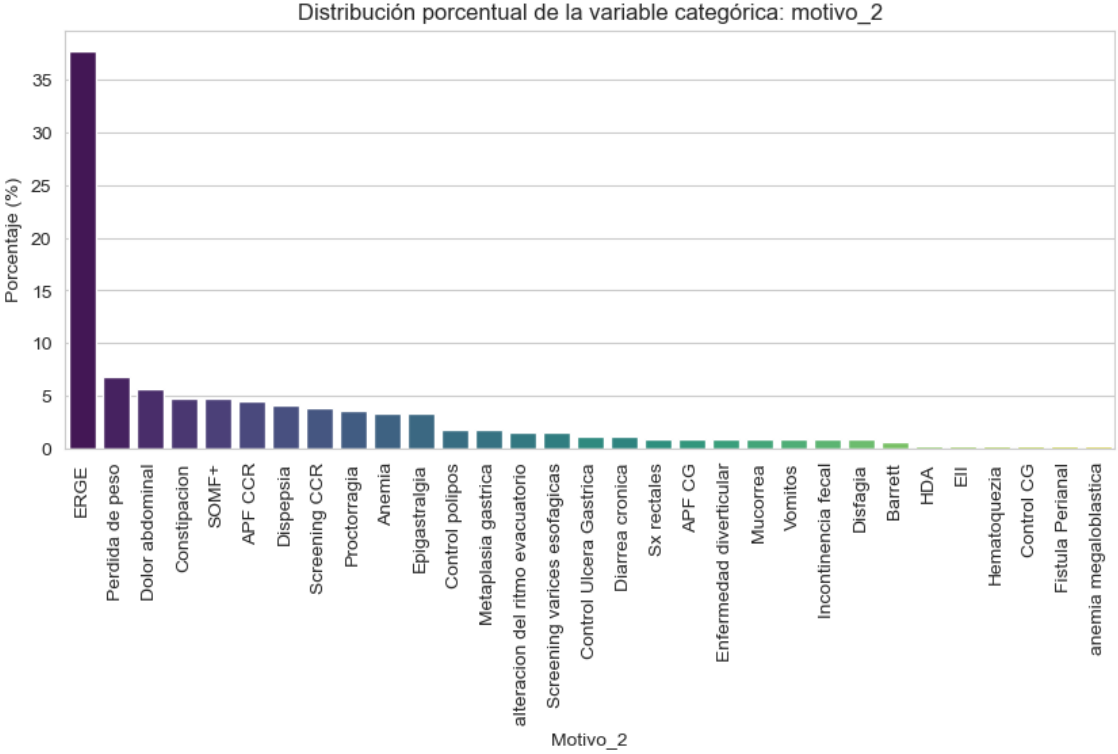


Figura 4. Distribución porcentual de obras sociales.

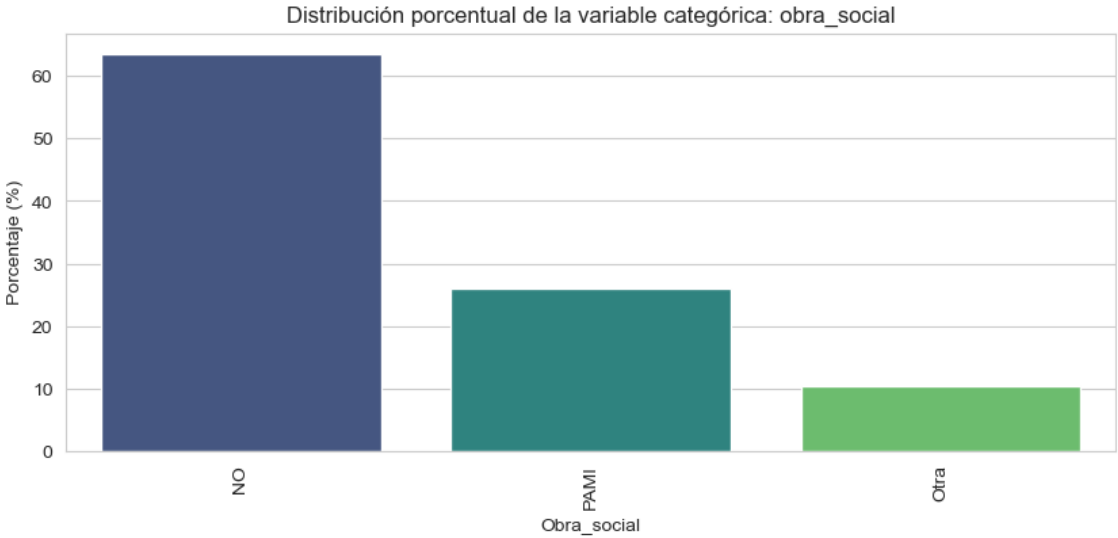


Figura 5. Distribución porcentual de “ASA”.

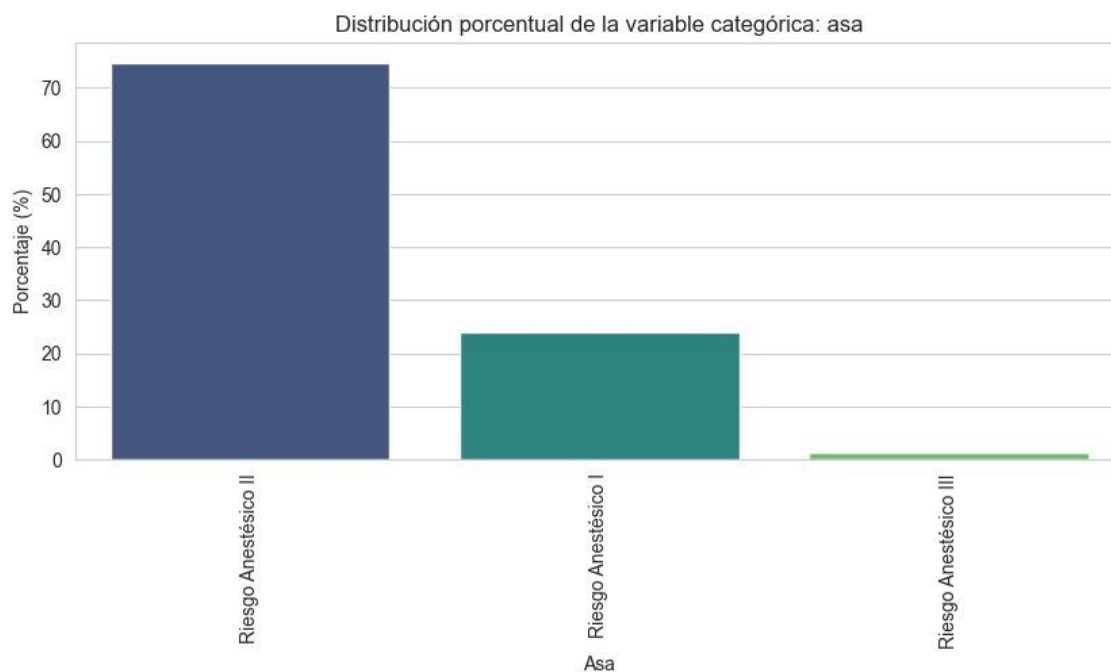
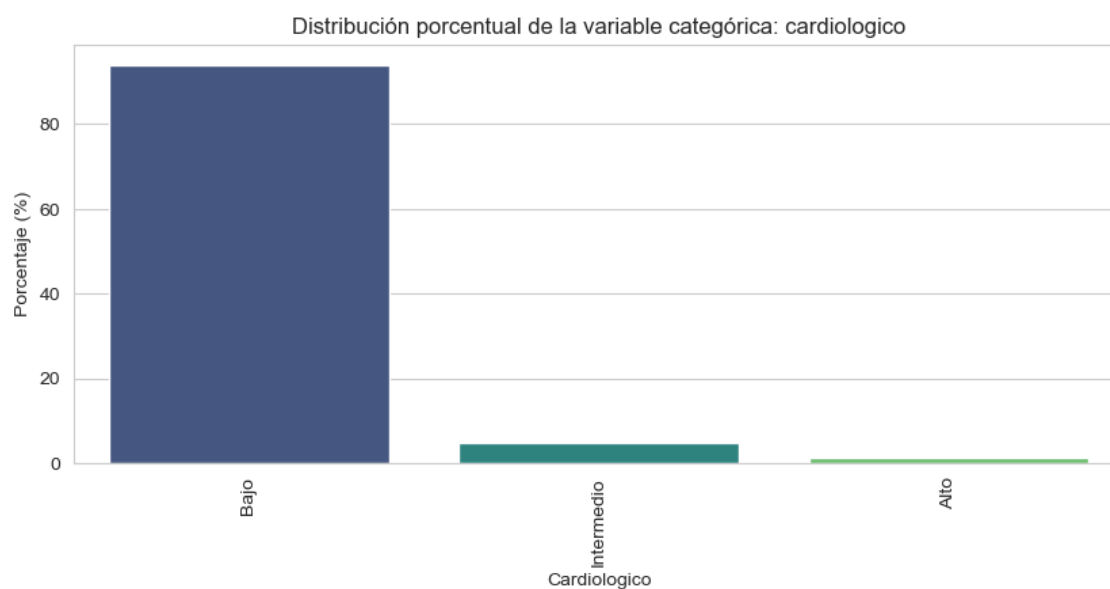


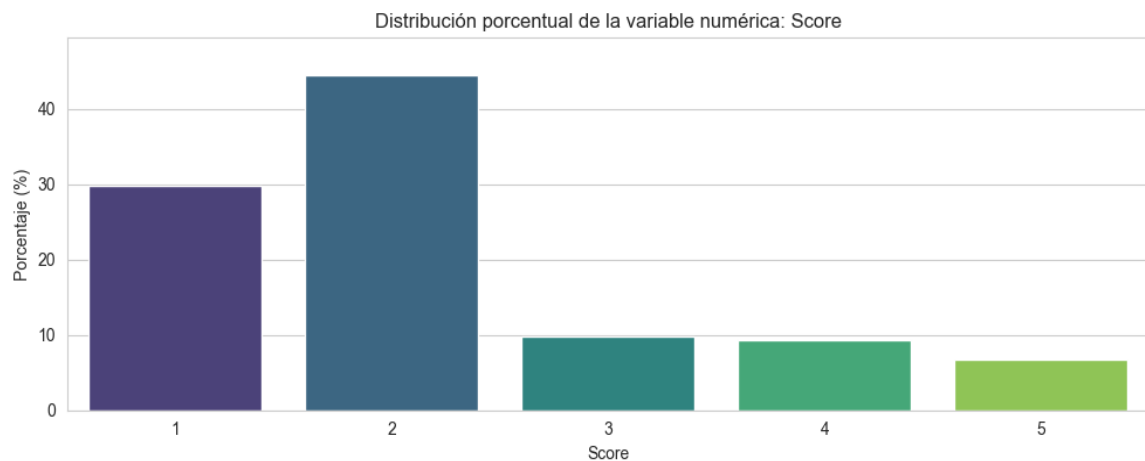
Figura 6. Distribución porcentual de “cardiológico”.



2.3.2 Variables Numéricas

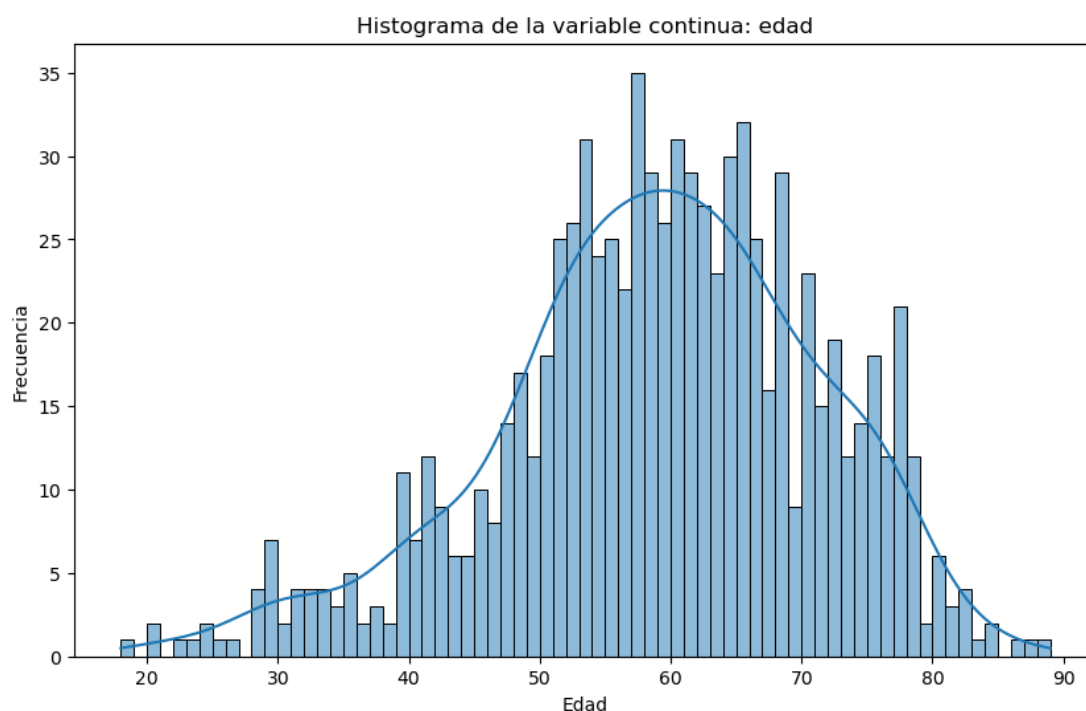
- **SCORE:** El Score es el valor asignado que marca la urgencia de atención dentro de una misma patología. Esta variable es una variable categórica ordinal siendo 1 la menor prioridad y 5 la prioridad máxima (Ver Figura 7).

Figura 7. Distribución porcentual “Score”.



- **EDAD:** Edad del paciente al momento en que se anota en la lista (Ver Figura 8). La distribución etaria presenta una ligera asimetría hacia la izquierda (left-skewed), lo que sugiere que hay una mayor concentración de pacientes en el rango de 50 a 70 años, con una menor cantidad de pacientes más jóvenes o muy mayores. En este sentido, se observa un pico de frecuencia alrededor de los 60 años. La densidad de probabilidad suavizada (línea azul) sugiere que la edad más común en la muestra está entre 55 y 65 años.

Figura 8. Histograma edad.



2.3.3 Series Temporales

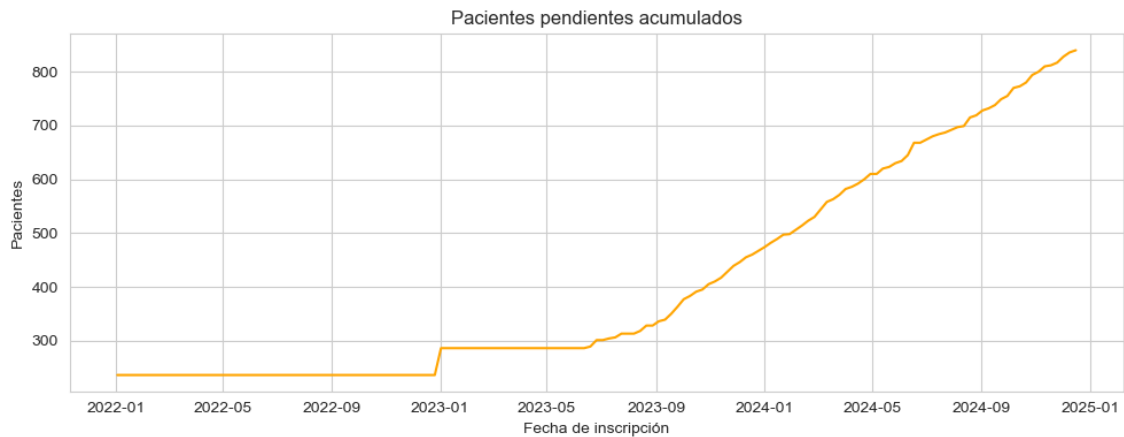
Utilizamos series temporales como recurso para entender el comportamiento de la demanda tanto de pacientes como estudios. Si bien existe una correlación entre ambos, buscamos analizar cada uno por separado para entender su media y picos de demanda.

La lista de espera se creó en 2022 con un registro manual, que en principio no contemplaba una fecha de inscripción. Esta comenzó a registrarse a partir de septiembre del 2023 con la digitalización de la lista, por lo que los pacientes ingresados previos a esta fecha figuran únicamente con el año.

2.3.3.1 Pacientes

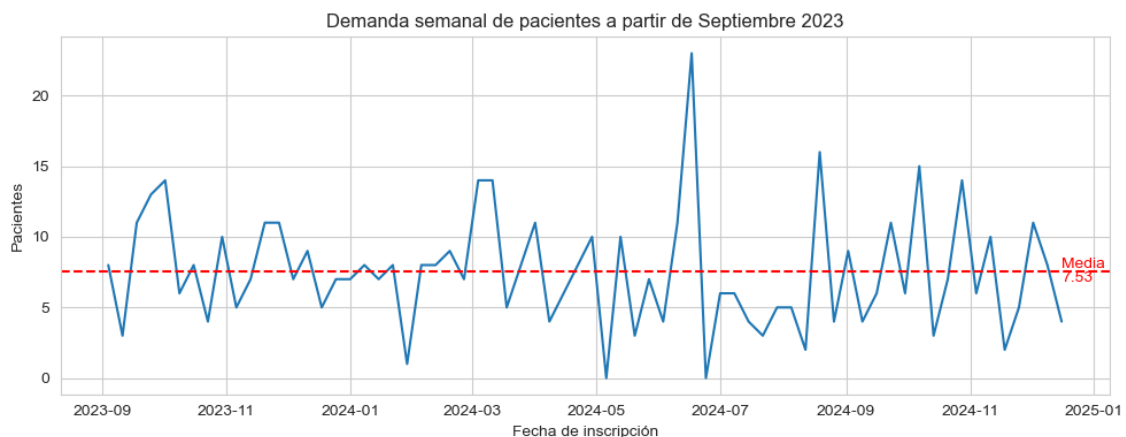
En la Figura 9 podemos observar un crecimiento constante a partir de 2022, con un aumento notable a partir de Septiembre de 2023, alcanzando un acumulado de 800 pacientes a finales de 2024.

Figura 9. Pacientes pendientes acumulados.



La Figura 10 presenta la cantidad de pacientes anotados en función de la fecha de inscripción. Como se observa, hay fluctuaciones notables en el número de inscripciones mensuales, con picos que alcanzan hasta 20 pacientes en ciertos meses, mientras que en otros se registra un número mucho más bajo. El promedio de pacientes por semana es de 7.53 ~ 8 pacientes, lo que sugiere que, aunque hay picos altos de inscripciones, la cantidad tiende a estar generalmente por debajo de este promedio en varias ocasiones.

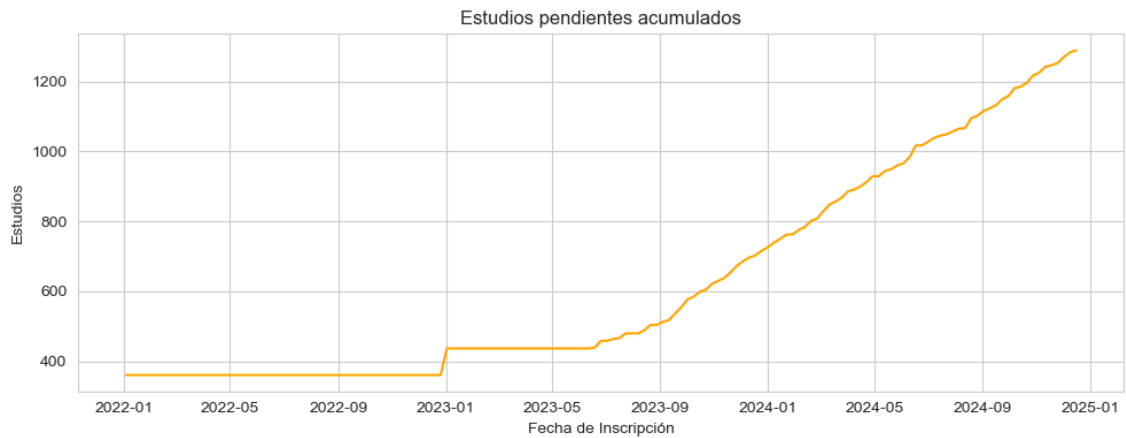
Figura 10. Demanda semanal de pacientes a partir de Septiembre 2023.



2.3.3.2 Estudios

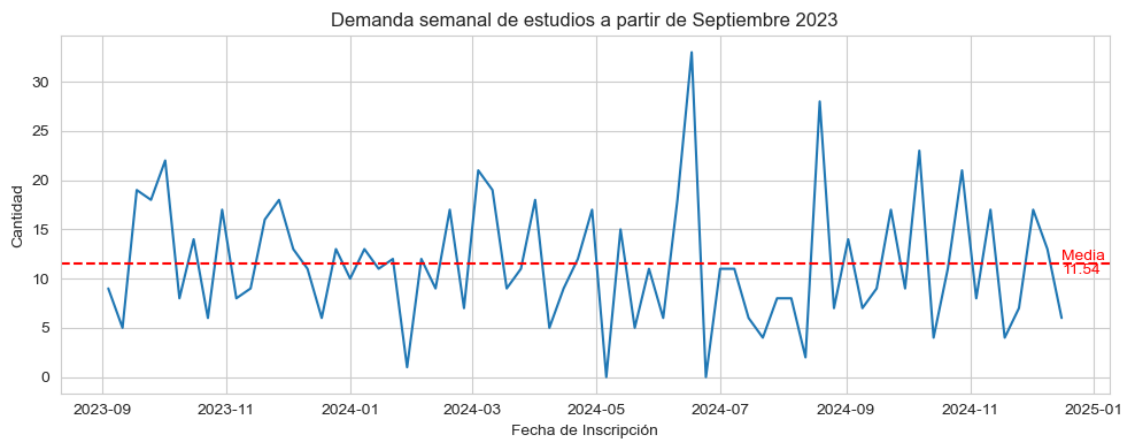
La cantidad de estudios por pacientes está directamente relacionada a la cantidad de pacientes inscritos. Se estima que por cada paciente se realiza 1.5 estudios en promedio. Es por eso, que la línea muestra un crecimiento más pronunciado, alcanzando un máximo de 1200 estudios a finales de 2024.

Figura 11. Estudios pendientes acumulados.



En la Figura 12 observamos la demanda semanal de estudios desde septiembre de 2023, tendiendo en promedio de 11.54 ~ 12 estudios por semana. Para nuestro ejercicio de simulación tendremos esto como input para la tasa de llegadas semanal.

Figura 12. Demanda semanal de estudios a partir de Septiembre 2023.



2.3.3.3 Estacionalidad de la serie

El test de Dickey-Fuller aumentado (ADF) es una prueba estadística utilizada en series temporales para determinar si una serie es estacionaria, es decir, si sus propiedades estadísticas (como la media y la varianza) se mantienen constantes en el tiempo. Este test evalúa la presencia de una raíz unitaria que indicaría que la serie sigue un comportamiento no estacionario y tiende a alejarse indefinidamente de su media. La hipótesis nula del test plantea que la serie no es estacionaria, mientras que la alternativa afirma que sí lo es. Si el valor p del test es menor a un nivel de significancia (por ejemplo, 0.05), se rechaza la hipótesis nula y se concluye que la serie es estacionaria. (Dickey, D. A., & Fuller, W. A. (1979).)

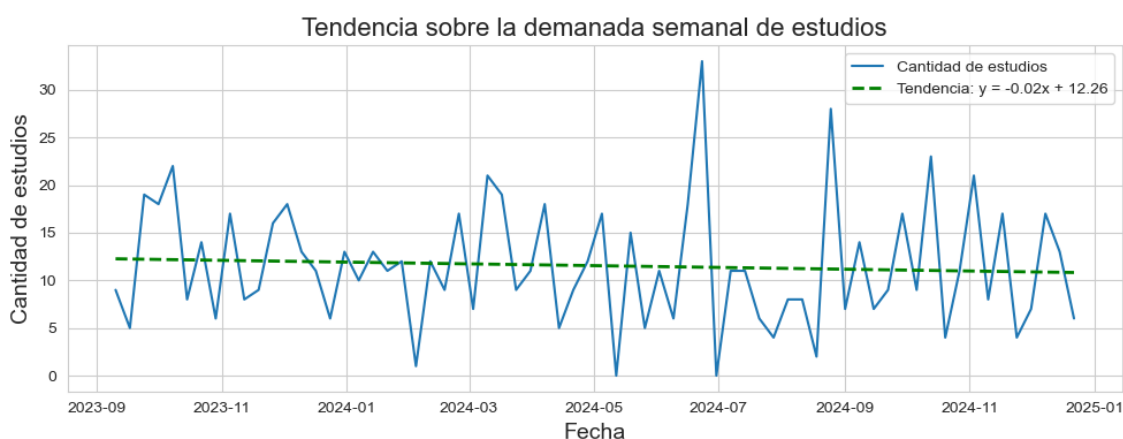
Tomando valores entre Septiembre de 2023 y Diciembre 2024 (donde tenemos una granularidad semanal de la llegada de los pacientes) los resultados del test de Dickey-Fuller aumentan la confianza en que la serie es estacionaria, dado que el estadístico ADF es -10.33 y el valor p es extremadamente bajo ($2.84e-18$), rechazando la hipótesis nula. Esto sugiere que es adecuada para modelos de pronóstico.

Este comportamiento estacionario podría relacionarse a que la cantidad de horas de atención en consultorios es estable y por lo consiguiente lo es la cantidad de estudios endoscópicos solicitados.

A partir de lo mencionado tomamos como supuesto que la demanda del servicio se mantendrá constante en los próximos años, si no hay cambios en su funcionamiento.

En la Figura 13 podemos ver gráficamente como la demanda semanal de los pacientes no varía significativamente observando que la pendiente de la recta de tendencia tiende a cero.

Figura 13. Tendencia demanda semanal de estudios.



3. Metodología

3.1 División del problema

El problema abordado en este trabajo de tesis se divide en dos partes:

1. Lista de Espera: Gestión de la lista de espera de 800 pacientes (que representan alrededor de 1200 estudios):
2. Nuevos Demanda: Estudios a partir de Enero 2025.

3.1.1 Lista de Espera

Destinamos una capacidad extra de turnos de quirófano por semana en función de mitigar la lista, siendo un recurso adicional a la demanda habitual.

Dado que no se agregarán más pacientes utilizaremos los parámetros previamente mencionados para generar por única vez el ranking de prioridad de atención.

La lista contiene pacientes que datan de anotarse en 2022 y 2023. Es de esperar que muchas de estas personas hayan optado por realizar sus estudios en otro centro o que ya no lo requieran. Para contemplar este fenómeno, evaluaremos 3 escenarios: Un escenario pesimista, un escenario moderado y un escenario optimista.

No mediremos el nivel de servicio (Ver Sección 1.5) dado que más del 80% de los pacientes se encuentran excedidos de los tiempos recomendados previamente mencionados (Ver Sección 1.1).

3.1.2 Nueva Demanda

A partir de 2025 los pacientes que conformen la nueva demanda serán incorporados a una lista en la que evaluaremos distintos aspectos clave: la cantidad de pacientes que ingresan semanalmente al sistema, la prioridad asignada a cada uno de ellos, y el comportamiento dinámico de la lista de espera mediante una simulación basada en patrones observados en años anteriores.

Esta simulación nos permitirá proyectar distintos escenarios y evaluar el nivel de servicio alcanzado en función de objetivos establecidos y recursos disponibles.

3.2 Score mediante Machine Learning

3.2.1 Motivación

Si bien el trabajo de tesis tiene foco en la simulación identificamos que uno de los cuellos de botella operativos es la necesidad de puntuar manualmente a cada paciente, una tarea que demanda tiempo, y es difícil de sostener en la práctica diaria. Para aliviar esta carga hemos aprovechado los datos disponibles en la lista para entrenar un modelo de Machine Learning que permita asignar automáticamente un score de prioridad a los nuevos pacientes que ingresen en el sistema a partir de 2025.

Cuando ingresa un nuevo paciente se lo incorpora a la lista digital completando únicamente las variables que fueron descriptas en el dataset principal (Ver Sección 2.1). Una vez ingresados los datos, se ejecuta el modelo previamente entrenado, que procesa la información y asigna automáticamente un score de prioridad entre 1 y 5. Este resultado se incorpora directamente en la lista, sin necesidad de realizar un análisis individual para cada paciente.

Entendemos que este volumen de datos de entrenamiento aún es limitado y que los modelos de este tipo suelen requerir grandes cantidades de información para alcanzar niveles altos de robustez y generalización. Sin embargo, consideramos que representa un punto de partida válido: para su construcción, utilizamos tres años de datos históricos, abarcando un espectro amplio de situaciones clínicas y escenarios reales. Esto nos permitió capturar patrones relevantes para automatizar la asignación de prioridad, y creemos que su desempeño mejorará a medida que se incorporen nuevos datos y acumule más experiencia en su aplicación.

3.2.2 Variable Objetivo

Dado que el objetivo del modelo es predecir un score de 1 a 5, se trata de un problema de clasificación ordinal, ya que las clases tienen un orden inherente (una prioridad de 5 es mayor que una de 1). Sin embargo, en la implementación con Auto Gluon, el problema se trata como una clasificación multiclase.

Para garantizar que el modelo tenga representatividad en todas las clases, utilizamos muestreo estratificado al dividir los datos en entrenamiento y prueba, asegurando que todas las patologías pertenecientes a la columna referencia “*código paper*” estuvieran presentes en ambos subconjuntos. Esta técnica es fundamental para evitar sesgos en la distribución de las clases durante el entrenamiento.

Además, dentro del análisis exploratorio vimos un desbalance en las clases de mayor urgencia. Esto nos hace sentido, ya que si la urgencia es muy alta se le asigna un turno sin entrar en la lista de espera. Empleamos Auto Gluon para realizar un balanceo de clases previo creando muestras sintéticas para poder equilibrar el dataset previo al entrenamiento del predictor.

3.2.3 Entrenamiento del modelo con AutoML

El AutoML (Automated Machine Learning) es una tecnología que automatiza el proceso de entrenamiento, optimización y selección de modelos de Machine Learning. Su objetivo es facilitar la implementación de modelos sin necesidad de una configuración manual extensa, permitiendo que incluso usuarios con menos experiencia en ML obtengan modelos de alto rendimiento.

Para nuestro caso, utilizaremos Auto Gluon, una librería de Auto ML desarrollada por Amazon. El mismo automatiza la selección de algoritmos, la ingeniería de atributos y el ajuste de hiper parámetros, con gran precisión en sus predicciones y bajo esfuerzo de configuración.

Para evaluar el rendimiento de los modelos utilizamos validación cruzada (cross-validation). Este proceso consiste en dividir el conjunto de datos en “k subconjuntos” o folds. El modelo se entrena en k-1 partes y se evalúa en la parte restante, repitiendo este proceso varias veces con diferentes combinaciones. Esto ayuda a obtener una estimación más estable y robusta del rendimiento del modelo, reduciendo el riesgo de sobreajuste (overfitting) (VanderPlas, n.d.).

3.2.4 Ranking de los pacientes

Una vez que establecemos el nivel de urgencia de cada paciente con el score (Ver Sección 2.2.2) evaluamos el tiempo de espera en la lista.

Es importante tener en cuenta el tiempo que el paciente está en la lista para no exceder los tiempos recomendados de espera (Ver sección 1.1)

Entonces, emplearemos el:

- Tiempo Recomendado
- Tiempo en la lista: El cual se calcula como la diferencia entre la fecha de inscripción y la fecha actual en días.

La diferencia entre ambos tiempos será el delta de tiempo de atención como se expresa en la siguiente ecuación:

$$\Delta \text{ tiempo atención} = \text{Tiempo recomendado} - \text{Tiempo en la lista}$$

Ajustaremos el Score, llamándolo Score de Atención Ajustado (SAA), dependiendo del valor de $\Delta \text{ tiempo atención}$ de la siguiente manera:

- Para los casos donde $\Delta \text{ tiempo atención} < 0$

$$\text{Score de Atención Ajustado} = \Delta \text{ tiempo atención} * \text{Score}$$

- Para los casos donde $\Delta \text{ tiempo atención} \geq 0$

$$\text{Score de Atención Ajustado} = \Delta \text{ tiempo atención} / \text{Score}$$

El objetivo de este ajuste es que el Score de Atención Ajustado (SAA) refleje correctamente la prioridad clínica y temporal del paciente, de modo que valores más bajos de SAA indiquen mayor prioridad.

La razón por la cual aplicamos una fórmula distinta según el signo de Δ tiempo atención es la siguiente: Si simplemente dividiéramos Δ tiempo atención por el score en todos los casos, los pacientes con retraso en la atención (es decir, con $\Delta < 0$) podrían obtener valores de SAA cercanos a cero al tener un score elevado. Esto iría en contra del criterio deseado, ya que estos pacientes deberían tener mayor prioridad (es decir, un SAA más negativo).

Cada paciente dentro de la lista de espera tendrá un SAA. Para determinar un ranking, ordenaremos este índice de manera ascendente, es decir, los pacientes que tengan menor SAA irán al principio de la fila.

Es importante remarcar que, si bien el Score no toma el 100% de la decisión del lugar en el ranking, será de vital importancia para asignar la posición a aquellos pacientes que comparten tiempos recomendados y tiempos en lista similares.

3.3 Gestión de la nueva demanda

3.3.1 Teoría de Colas

La teoría de colas es un campo de estudio que se centra en el análisis de sistemas que gestionan la espera de elementos (como personas o paquetes) para recibir un servicio. En estos sistemas, los "clientes" llegan de acuerdo con un proceso estocástico y requieren atención en uno o más "servidores". Los modelos de teoría de colas permiten predecir comportamientos como el tiempo de espera, el tamaño de la cola y la utilización de los recursos, con el fin de optimizar la eficiencia de los servicios y la experiencia del cliente. (Hillier & Lieberman, 2010)

Para simular la atención en el quirófano de un hospital, se ha optado por emplear un modelo de Teoría de Colas. Utilizaremos un modelo M/M/1, en el que la llegada de pacientes se modela según el número de pacientes semanales y el servicio se determina por la cantidad de estudios que se pueden realizar cada semana.

Un paciente puede requerir más de un estudio, por lo cual la demanda de estudios variará semanalmente.

Todo el sistema se modelará bajo la premisa de un solo servidor para simplificar el nivel de cálculos.

3.3.2 Simulación

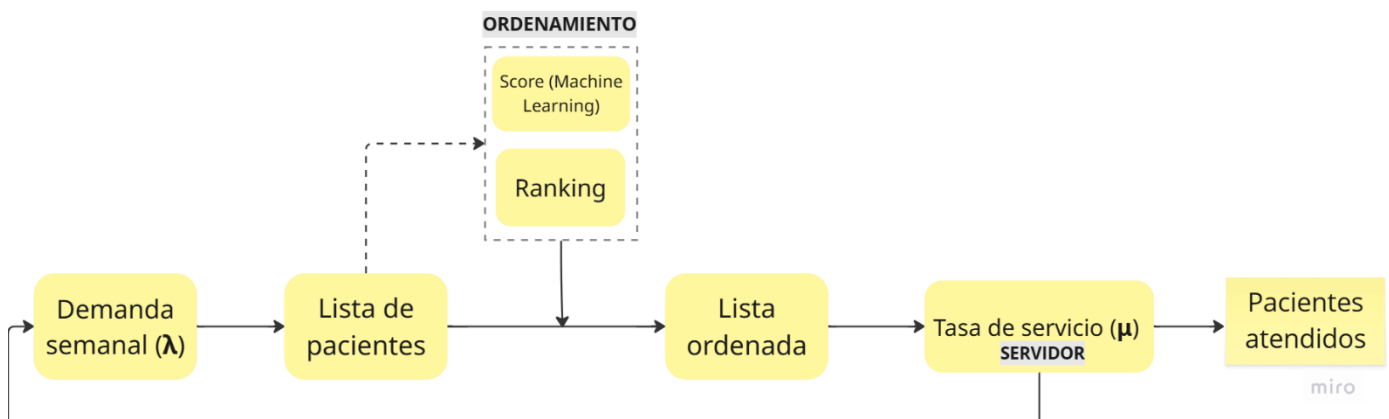
La simulación será desarrollada en Python bajo la librería Simpy. Esta librería nos permite modelar sistemas basándose en demanda, recursos y capacidad. A su vez

nos permite correr tiempo en un espacio virtual y capturar distintas métricas para el posterior análisis de resultados. Para modelar la llegada semanal de los pacientes hemos utilizado como ejemplo el dataset completo (800 personas).

La dinámica de la simulación será (Ver Figura 14):

1. Semana a semana, aleatoriamente y con reposición se extraerá de la lista histórica una cantidad de pacientes. Las llegadas semanales estarán determinadas por una tasa λ que sigue una distribución Poisson, lo que permite incorporar variabilidad y realismo al modelo.
2. Se sumarán estos pacientes (con la información correspondiente a su situación y patología) a una nueva lista que partirá en blanco y se les asignará como fecha de llegada la semana correspondiente a la simulación.
3. Se realizará el scoring para entender cuál es la prioridad que tienen según la patología y con esto se definirá el ranking. Una vez generado el ranking se los llamará para atenderse en el servidor. Si bien el Score se realizará una sola vez porque no depende del tiempo en lista, el ordenamiento se realizará semana a semana.
4. Dentro del servidor la tasa de servicio (μ) estará expresada en estudios por semana y la transformación de cantidad de personas a cantidad de estudios dependerá de los pacientes que sean elegidos aleatoriamente. A medida que se realicen estudios bajará la disponibilidad del servidor. En el caso de que la próxima persona en la lista no coincida con la cantidad de estudios disponibles, se elegirá el próximo paciente que pueda utilizar ese recurso. Por ejemplo: Si solo queda lugar para realizar un estudio, y la próxima persona en la lista requiere dos estudios, se seguirá bajando en la lista hasta encontrar la primera persona que requiera un estudio.
5. Se procederá a atender a los pacientes, asignarles una fecha de salida y sacarlos de la lista de espera.

Figura 14. Dinámica de la simulación.



3.3.3 Nivel de Servicio

Con la salida de los datos en la simulación, contaremos con la fecha de entrada a la simulación y una fecha simulada de salida y los tiempos recomendados para realizar el estudio a nivel paciente.

Con esta información podemos entender si el paciente se realizó los estudios en tiempo y forma. Utilizaremos el tiempo “On Time” en el caso positivo, y “No On Time” en los casos donde los pacientes hayan tardado más de los tiempos recomendados en atenderse.

Mediremos el nivel de servicio bajo la métrica:

$$\% \text{ On Time} = \frac{\sum \text{Pacientes On Time}}{\sum \text{Pacientes}} * 100$$

4. Resultados

4.1 Lista de espera

Como se comentaba en el capítulo anterior, en la Sección 3.1, separaremos en dos la problemática. Por un lado, trataremos la lista pendiente de 800 pacientes (sin el agregado de nuevos pacientes) y por otro lado iremos gestionando la llegada de los nuevos.

En esta sección abordaremos la lista pendiente. Esta lista tiene alrededor de 800 pacientes que representan 1200 estudios.

Lo primero que haremos es utilizar el algoritmo (que explicaremos en el capítulo de priorización) para armar un ranking general. Dado que la lista en cantidad de pacientes podemos alcanzarla con un solo ordenamiento.

La mayoría de los pacientes en esta lista han superado el tiempo recomendado debido a la falta de recursos, no mediremos % On Time de estos pacientes. También, teniendo en cuenta lo anterior, es probable que los pacientes hayan optado por realizar sus estudios en otros centros debido a la impaciencia, o que ya no los requieran. Por lo antes expuesto, definimos 3 escenarios:

- ❖ **Pesimista:** La demanda actualizada sea un 100% de los estudios pendientes
- ❖ **Optimista:** A la demanda pesimista le descontaremos los estudios de pacientes que tengan otra obra social, o que por la urgencia del estudio hayan decidido hacerlo en otro centro (Para esto último consideramos a aquellos con Score 5).
- ❖ **Moderado:** Este será el escenario intermedio a los anteriores, y por lo tanto el esperado.

Para cada uno de los escenarios construimos distintas curvas de demanda en el tiempo, variando la cantidad de estudios por semana para mitigar la lista.

La cantidad de estudios varía en función de turnos de quirófano disponibles:

- ❖ **7 estudios:** Equivalente a 1 quirófano corto (4hs) por semana
- ❖ **14 estudios:** Equivalente a 1 quirófano largo (8hs) por semana
- ❖ **21 estudios:** Equivalente a 1 quirófano corto (4hs) más 1 quirófano largo (8hs) por semana
- ❖ **28 estudios:** Equivalente a 2 quirófanos largos (16hs) por semana

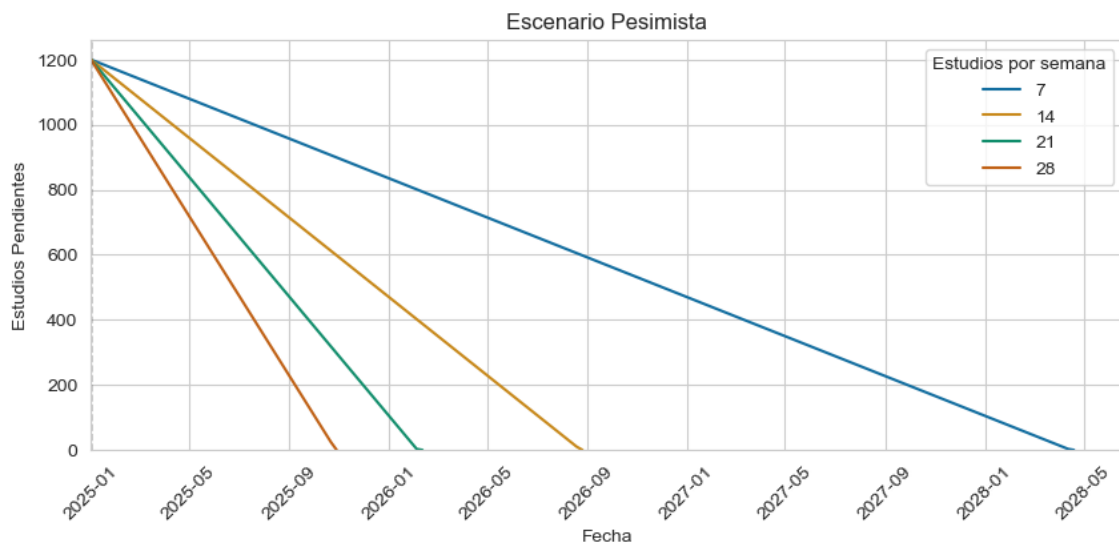
A continuación, están las curvas para cada uno de los escenarios.

El escenario pesimista, contempla el 100% de los estudios, por lo tanto 1200 estudios. A continuación, en la Figura 15 vemos la evolución de la lista a través del

tiempo en caso de hacer una determinada cantidad de estudios por semana en forma sostenida.

Para el caso de siete estudios (quirófano corto), vemos como se hace inviable debido a al tiempo que llevaría licuar esta lista. Las demás opciones, si bien reducen visiblemente los tiempos, siguen estando en rangos elevados para controlar la lista en el corto plazo.

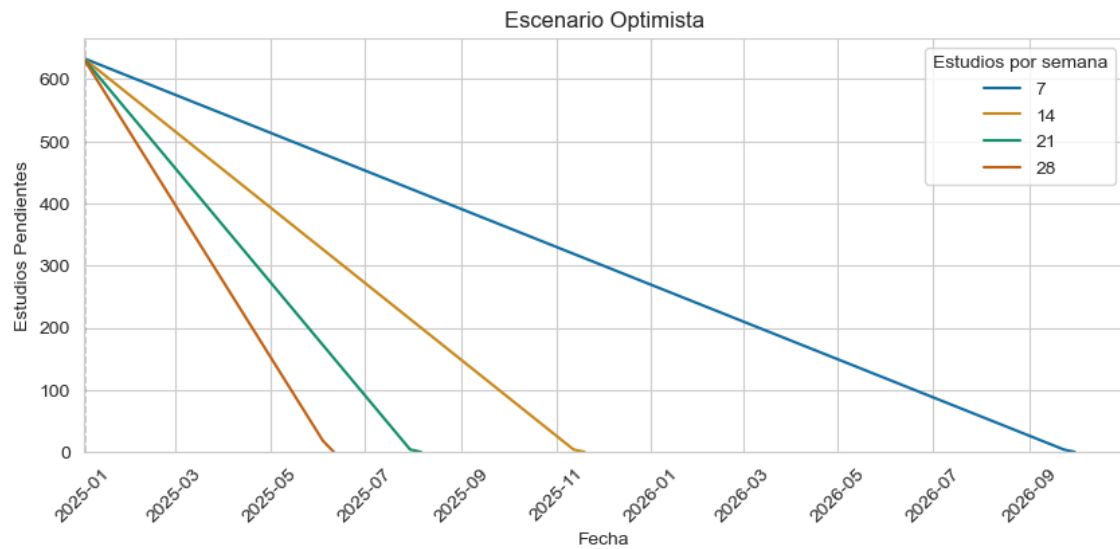
Figura 15. Escenario pesimista.



En el escenario optimista (total de estudios pendientes menos aquellos que hayan requerido personas con obra social, o pacientes con Score nivel 5) tiene un total de 634 estudios. La Figura 16 vemos la evolución de la lista a través del tiempo en caso de hacer una determinada cantidad de estudios por semana en forma sostenida.

En este punto, vemos que los márgenes para controlar la lista oscilan desde los 5 meses y 1 año y 9 meses. En este punto prácticamente descartamos la viabilidad de licuar la lista con 7 estudios, dado que nos interesa poder lograrlo antes del año y medio. La opción de realizar 14 por semana sería la mejor opción para poder lograrlo.

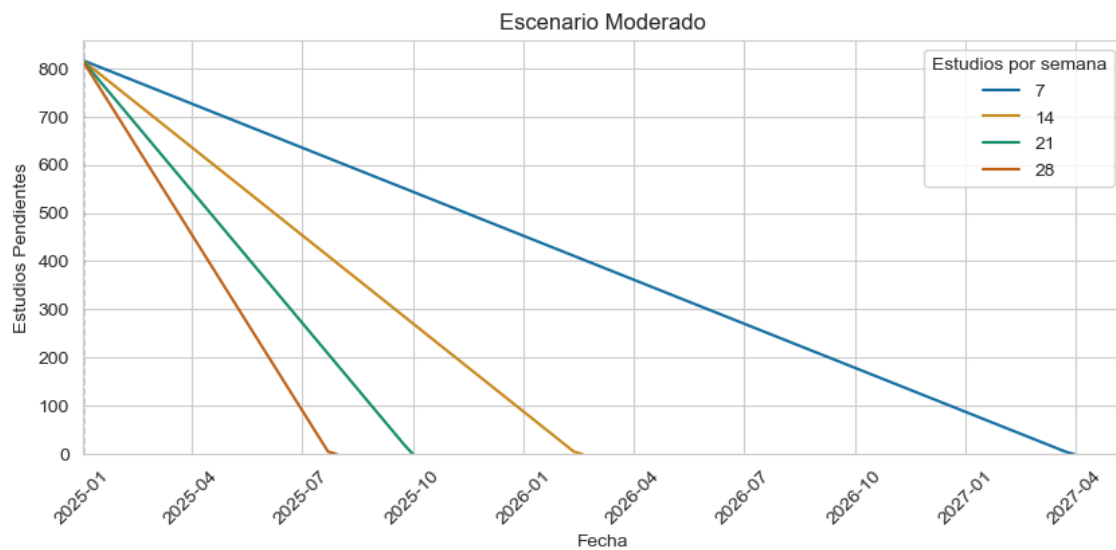
Figura 16. Escenario optimista.



Por último, el escenario moderado, el cual está en un punto medio entre los escenarios anteriores. Arranca en 817 estudios y confirmaremos en este como escenario más probable para la evacuación de la lista. La Figura 17 vemos la evolución de la lista a través del tiempo en caso de hacer una determinada cantidad de estudios por semana en forma sostenida.

Descartada la opción de 7 estudios como mencionamos en el escenario anterior, la siguiente opción de un quirófano largo (14 estudios) es la indicada. Si bien vemos que está por encima del año, sigue respetando nuestro target de año y medio. 14 estudios por semana será nuestra guía para tratar la lista y licuarla dentro del primer trimestre del próximo año.

Figura 17. Escenario moderado.



4.1 Priorización

4.1.1 Score: Selección de modelos y evaluación

El modelo fue entrenado con Auto Gluon, una biblioteca de AutoML que automatiza la prueba de distintos algoritmos y elige el más adecuado según una métrica de rendimiento. Entre los modelos que evaluamos se incluyen: XG Boost, LightGBM, Random Forest, redes neuronales y Cat Boost.

También implementamos modelos ensamblados que combinan predicciones de varios modelos base para mejorar la precisión final.

Se compararon los modelos según su índice de accuracy en validación y prueba. En la práctica, la selección del mejor modelo se realiza en base al score en datos de test (*"score_test"*) que miden el rendimiento del modelo en datos nuevos.

Tabla 4. Comparación de modelos.

ranking	model	score_test	score_val	eval_metric	pred_time_test	pred_time_val	fit_time	fit_order
1	WeightedEnsemble_L2	0.923	0.892	accuracy	0.06	0.01	1.79	14
2	ExtraTreesGini	0.917	0.850	accuracy	0.07	0.05	0.64	9
3	ExtraTreesEntr	0.911	0.833	accuracy	0.07	0.04	0.70	10
4	LightGBM	0.905	0.867	accuracy	0.01	0.01	0.99	5
5	XGBoost	0.905	0.883	accuracy	0.04	0.00	0.72	11
6	RandomForestEntr	0.905	0.817	accuracy	0.07	0.05	0.80	7
7	RandomForestGini	0.905	0.842	accuracy	0.09	0.05	0.91	6
8	NeuralNetTorch	0.881	0.858	accuracy	0.01	0.00	7.73	12
9	CatBoost	0.863	0.858	accuracy	0.00	0.00	29.17	8
10	LightGBMLarge	0.857	0.800	accuracy	0.03	0.01	2.97	13
11	LightGBMXT	0.827	0.800	accuracy	0.02	0.01	1.19	4
12	NeuralNetFastAI	0.631	0.633	accuracy	0.01	0.00	0.49	3
13	KNeighborsUnif	0.482	0.358	accuracy	0.00	0.00	0.00	1
14	KNeighborsDist	0.476	0.350	accuracy	0.00	0.00	0.01	2

En los resultados obtenidos, se observa que el *score_test* es ligeramente superior al *score_val* en todos los modelos evaluados.

Esta diferencia no se debe a data leakage: La diferencia entre ambos puntajes puede explicarse por el tamaño reducido del conjunto de test, que cuenta con 168 filas, lo cual lo hace más sensible a variaciones aleatorias. En cambio, el conjunto de validación proviene de una división interna del entrenamiento (672 filas), donde el modelo se entrena con aproximadamente 537 filas y se valida con otras 135, incluyendo casos de mayor complejidad. Por eso, el score de validación tiende a ser más bajo, pero representa una estimación más conservadora y robusta del rendimiento.

El modelo seleccionado fue *WeightedEnsemble_L2*, que obtuvo el mayor puntaje en validación y test. Se trata de un modelo ensamblado que combina las predicciones de otros modelos base para mejorar el rendimiento.

En la primera capa (L1) Auto Gluon entrena varios modelos base de diferentes tipos (árboles de decisión, redes neuronales, regresión lineal, etc.). Luego, en L2, *WeightedEnsemble_L2* toma sus predicciones y aprende a combinarlas de manera efectiva.

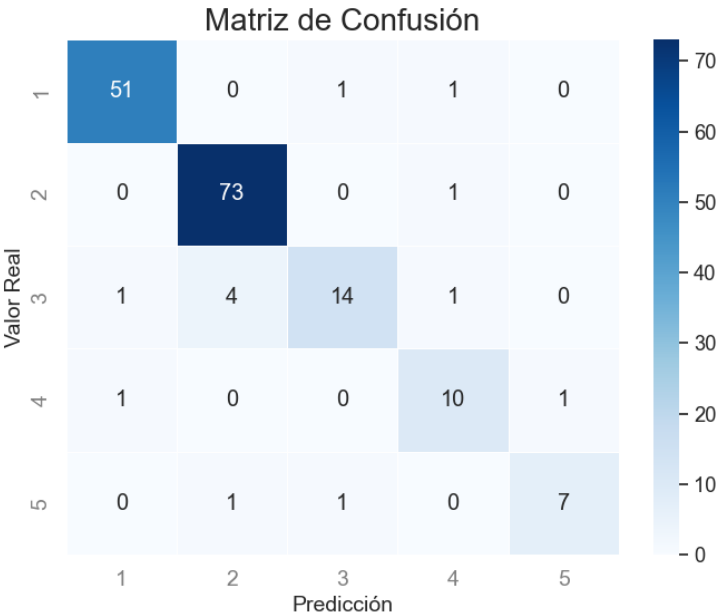
Las principales ventajas de esta técnica son mejorar el rendimiento general del modelo, al aprovechar la diversidad de los modelos base, y reducir la varianza y el sesgo en comparación con los modelos individuales.

4.1.2 Matriz de Confusión

Para evaluar el desempeño del modelo, se construyó una matriz de confusión que permite la cantidad de predicciones correctas e incorrectas por clase. Este análisis es fundamental para identificar la naturaleza de los errores, en particular aquellos casos en los que el modelo clasifica erróneamente a un paciente con alta urgencia como si tuviera baja prioridad. Este falso negativo representa el mayor riesgo, ya que podría retrasar la atención de personas que requieren intervención inmediata.

También es importante controlar los errores en sentido contrario, es decir, cuando el modelo asigna alta prioridad a pacientes que en realidad no la necesitan con urgencia. Aunque clínicamente menos grave, este error podría derivar en un uso ineficiente de los recursos disponibles y en la postergación de otros pacientes más críticos. Por lo tanto, si bien el objetivo principal es minimizar los falsos negativos en casos urgentes, también es necesario mantener bajo control los falsos positivos para no comprometer el sistema de priorización.

Figura 18. Matriz de confusión.



En la Fig. 18 se muestra cómo se relacionan las predicciones del modelo con los valores reales asignados por profesionales médicos. Los valores ubicados sobre la diagonal representan los casos correctamente clasificados, es decir, donde la predicción del modelo coincide con la gravedad asignada por el especialista. A partir de esta diagonal se puede calcular fácilmente la precisión (accuracy) del modelo como la suma de los aciertos dividida por el total de casos evaluados.

$$\% Accuracy = \frac{\sum Predicciones\ correctas}{\sum Predicciones} * 100$$

$$\% Accuracy = \frac{155}{168} * 100$$

$$\% Accuracy = 92,26\%$$

En cuanto a los errores de clasificación (valores fuera de la diagonal), se observa que son relativamente pocos en comparación con los aciertos. Además, una proporción de estos errores se da entre categorías adyacentes, lo que sugiere que, cuando el modelo falla, en varios casos lo hace prediciendo una gravedad cercana al valor real. A su vez, se destaca que no se presentan errores extremos (por ejemplo, confundir una gravedad 1 con una 5), lo cual es indicio de una buena calibración general del modelo.

4.1.3 Ranking

A partir de la fecha actual, tomando como input la fecha de inscripción y el tiempo recomendado, usamos la fórmula del Score de Atención Ajustado SAA (Ver Sección 3.2.4). Una vez calculado, se ordena ascendentemente (mientras menor sea el SAA, más prioridad tendrá). Los valores de este Score Ajustado resultarán negativos para aquellos casos donde el tiempo de espera lista supere a los tiempos recomendados y positivos para aquellos que tengan posibilidad de atenderse a tiempo.

En la Tabla 4 tenemos un extracto de la lista luego de ordenar por Ranking a los pacientes. Es de esperar que empiece por todos aquellos pacientes que se hayan anotado en 2022 y tengan prioridad 5 de atención. Si bien estos casos se conforman en la lista, es probable que se hayan atendido en otros centros debido a que tenían cobertura médica o no hayan podido esperar a atenderse.

Tabla 5. Extracto de ranking pacientes.

id	fecha inscripción	edad	estudio	código	score	tiempo recomendado	SAA	ranking
11	2022-01-01	61	DOBLE	7	5	14	-6015	1
93	2022-01-01	64	VCC	7	5	14	-6015	2
39	2022-01-01	53	VCC	10	5	14	-6015	3
200	2022-01-01	51	VCC	10	5	14	-6015	4
20	2022-01-01	69	DOBLE	11	5	60	-5785	5
60	2022-01-01	63	DOBLE	11	5	60	-5785	6
65	2022-01-01	63	VCC	11	5	60	-5785	7
82	2022-01-01	50	VCC	11	5	60	-5785	8
149	2022-01-01	43	DOBLE	11	5	60	-5785	9
171	2022-01-01	70	DOBLE	11	5	60	-5785	10

4.2 Demanda: Teoría de Colas

4.2.1 Variables

- **Llegada de pacientes (λ):** Se modela con una distribución Poisson con una media de 8 pacientes por semana.
- **Cantidad de estudios requeridos ($\lambda_{\text{estudios}}$):** No todos los pacientes requieren un solo estudio algunos pueden necesitar dos. Definimos un coeficiente de estudios para modelar esta variabilidad (C).
- **Turnos de quirófano (servidor):** Consideramos un solo servidor, modelando la demanda en forma semanal en un recurso que pueda realizar los estudios. Utilizamos un modelo M/M/1
- **Tasa de servicio (μ):** Dado el promedio de estudios evaluamos que el servicio debería realizar 13 estudios por semana para no acumular fila.

Tomamos como coeficiente de estudios a la cantidad promedio de estudios dividido los pacientes promedio:

$$C = 11.54 / 7.54 \approx 1.53$$

Aplicando este coeficiente a la llegada de pacientes semanal obtenemos la tasa de llegada de estudios:

$$\lambda_{\text{estudios}} = 8 \times 1.53 = 12.24$$

Concluimos que en promedio 8 pacientes requieren 12.24 estudios ≈ 13 estudios.

4.2.2 Factor de Ocupación

El factor de ocupación del sistema, definido como la tasa de llegada dividida por la tasa de salida, se expresa como:

$$\rho = \lambda_{\text{estudios}} / \mu = 12.24 / 13 \approx 0.94$$

Este valor indica que el sistema está ligeramente subcargado, lo que sugiere que la capacidad de atención es adecuada y no se generará una acumulación excesiva de estudios pendientes.

4.2.3 Cálculo del Tamaño de la Cola Estimado

El número esperado de estudios en la cola en un sistema M/M/1 se calcula como:

$$Lq = (\rho^2) / (1 - \rho)$$

Dado que el valor de Lq es aproximadamente 14.72 estudios en espera, lo que indica que el sistema puede manejar la demanda sin generar una cola extensa.

Entonces, para encontrar cuántos pacientes equivalen a la cola de 14.72 estudios en espera, usamos la siguiente fórmula:

$$\text{Pacientes en cola} = Lq / C$$

Donde:

- $Lq = 14.72 \text{ estudios}$
- $C = 1.53 \text{ estudios por paciente}$

Realizando el cálculo:

$$\text{Pacientes en cola} = 14.72 / 1.53 \approx 9.62 \text{ pacientes}$$

Concluimos que el valor promedio en cola será de 14.72 estudios en espera equivale a aproximadamente 10 pacientes en espera.

La tasa de llegada de estudios es ligeramente inferior a la capacidad del quirófano, lo que implica que el sistema puede operar de manera estable.

Se recomienda monitorear la demanda y, si es necesario, ajustar la capacidad para mantener una operación eficiente sin generar demoras excesivas.

4.3 Simulación

Validamos empíricamente lo anteriormente descrito con un modelo de Simulación que muestra la evolución de la cola en el tiempo.

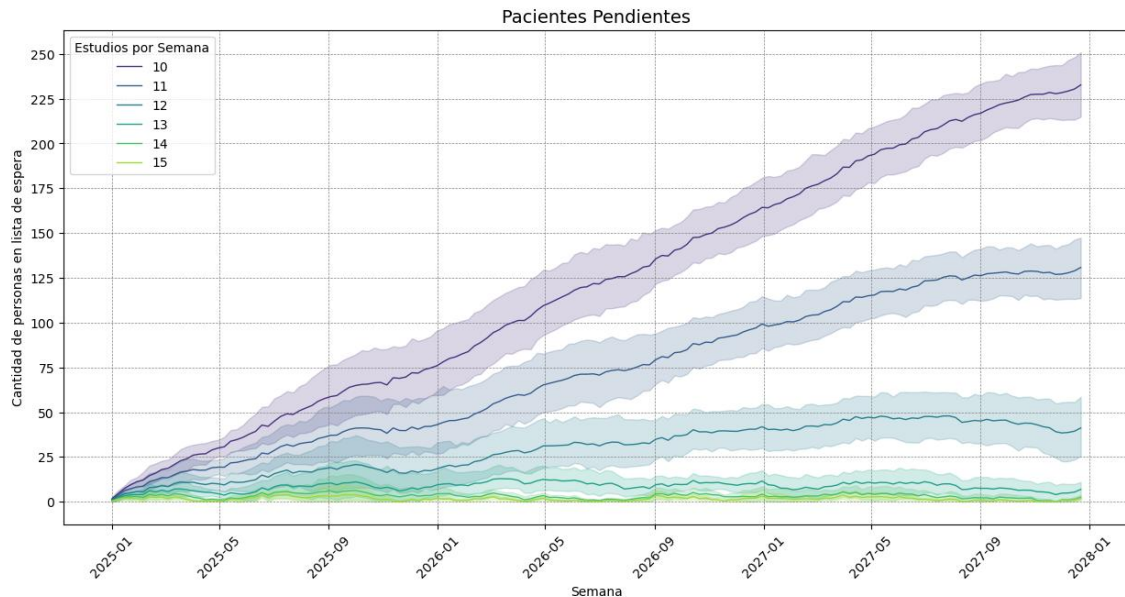
Como se observa en la Figura 19, hemos corrido simulaciones en las que la llegada de pacientes se modela mediante una distribución de Poisson con media igual a 8. Empleamos distintas cantidades de estudios realizados por semana (μ) para entender cómo afecta al sistema en cuanto a acumulación de personas en lista de espera durante tres años.

Para cada escenario hemos corrido 10 simulaciones para poder dar un intervalo de confianza del 90% de los valores expuestos.

Los primeros meses notamos que los intervalos de confianza entre los distintos niveles de servicios se solapan (mostrando que no hay evidencia significativa entre algunos). Sin embargo, con el correr del tiempo se separan quedando marcados los impactos en la lista de espera para las distintas tasas de servicio.

Podemos notar que a partir de 13 estudios las líneas se solapan entre sí. En este punto, se ha alcanzado un régimen estable donde la fila no aumenta en el tiempo y se mantiene dentro de los valores promedios L_q calculados previamente. Vemos además que incrementar recursos para la atención no mejora en forma la situación en significativa, por lo cual será el óptimo para esta demanda de 8 personas por semana.

Figura 19. Pacientes pendientes proyectados en el tiempo.



4.4 Nivel de Servicio

En la Figura 20 podemos ver la evolución del Nivel de Servicio para una demanda de 8 pacientes semanales variando la cantidad de estudios que se realizan.

El mismo consta en el eje x de distintas tasas de servicio (μ) para la misma demanda, y en el eje y del porcentaje de personas atendidas en tiempo (% On Time)

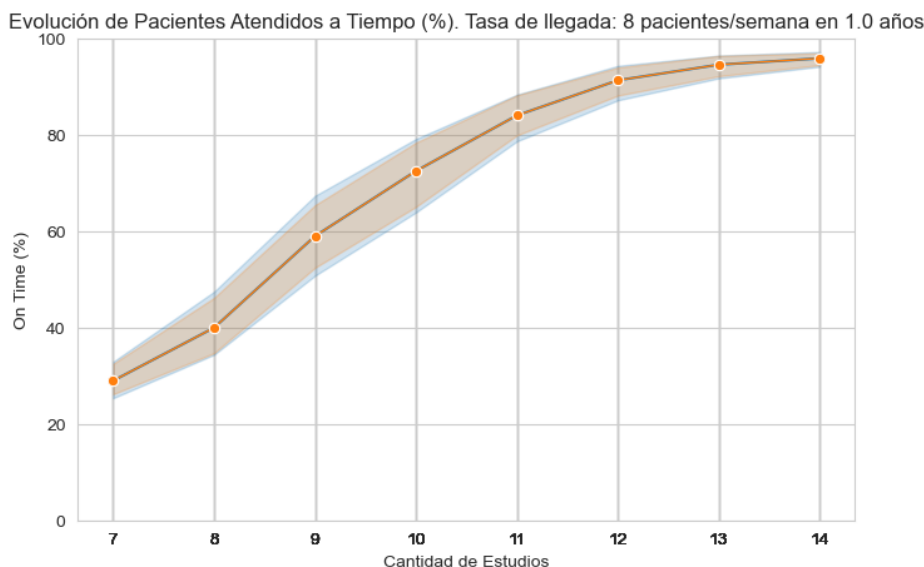
Este mismo está realizado para el plazo de 1 año. Si aumentamos el plazo en años para este mismo gráfico el % On Time caerá para aquellos puntos del sistema que acumulen filas de espera en poco tiempo.

Observamos que para el punto de estabilización de la lista (Cantidad de estudios mayor o igual a 13). En el caso de querer trabajar por debajo del sistema estable y acumular fila, podemos movernos a la izquierda de este punto en la curva y definiendo los targets de nivel de servicio esperado podemos encontrar la tasa de servicio (μ) necesario para alcanzarla.

Es necesario en este punto aclarar que, si trabajamos en un régimen que genera colas, es importante tener un plan para estabilizar y llevarla a cero al menos una vez al año en este caso para confiar en los niveles de servicio.

Por otro lado, para dimensionar la cola generada, podemos apoyarnos en la Figura anterior 19, posicionarnos del siguiente año y tomar el valor esperado como referencia y un intervalo de confianza.

Figura 20. Nivel de servicio.

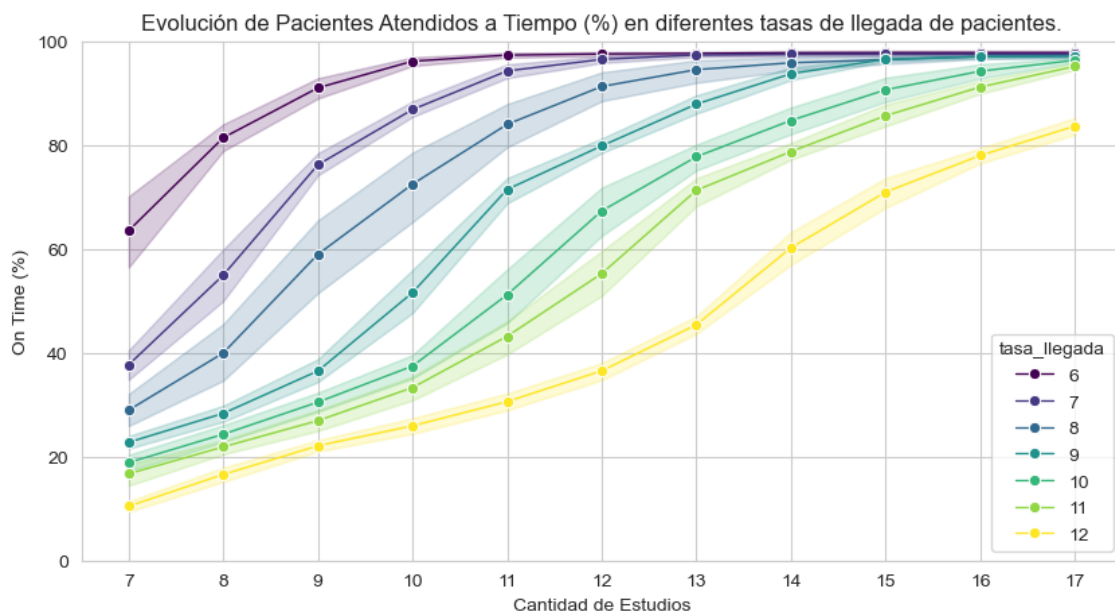


Con el objetivo de entender cómo podría variar el nivel de servicio con distintos tipos de demanda, en la Figura 21 vemos cómo afectaría esto a los distintos niveles de servicio. En el eje x se encuentran las distintas tasas de servicio (μ) y en el eje y del porcentaje de personas atendidas en tiempo (%On Time). Cada línea representa una tasa de llegada (λ) que va de 6 a 12 pacientes por semana.

Esperamos que el comportamiento funcione similar en los distintos niveles de demanda. Sin embargo, observamos que el nivel de servicio para una tasa de llegada de 12 pacientes tiene más inercia a alcanzar buenos niveles de servicio con el aumento de las tasas de servicio.

Notamos además que un quirófano largo (14 estudios) es suficiente para mantener un nivel de servicio por encima del 80%, pero se debe hacer un trade-off frente a cuál es la cola que podría generar en un año, dependiendo la tasa de llegada y poder licuar la en el corto plazo.

Figura 21. Nivel de servicio según diferentes tasas de llegada.



5. Discusión

A continuación, a partir de la información expuesta en los resultados, discutiremos un plan de acción para el Hospital J.M. Ramos Mejía.

El servicio actualmente cuenta 2 quirófanos cortos que equivalen a un quirófano largo (es decir 14 estudios por semana) para mitigar la lista y atender la nueva demanda originada.

Si bien pudimos establecer un ranking que se actualiza semana a semana con el ingreso de nuevos pacientes, hemos expuesto anteriormente que para cubrir la nueva demanda necesitamos una tasa de servicio de 13 estudios por semana. ¿Cuáles serían los recursos necesarios para poder mitigar la lista?

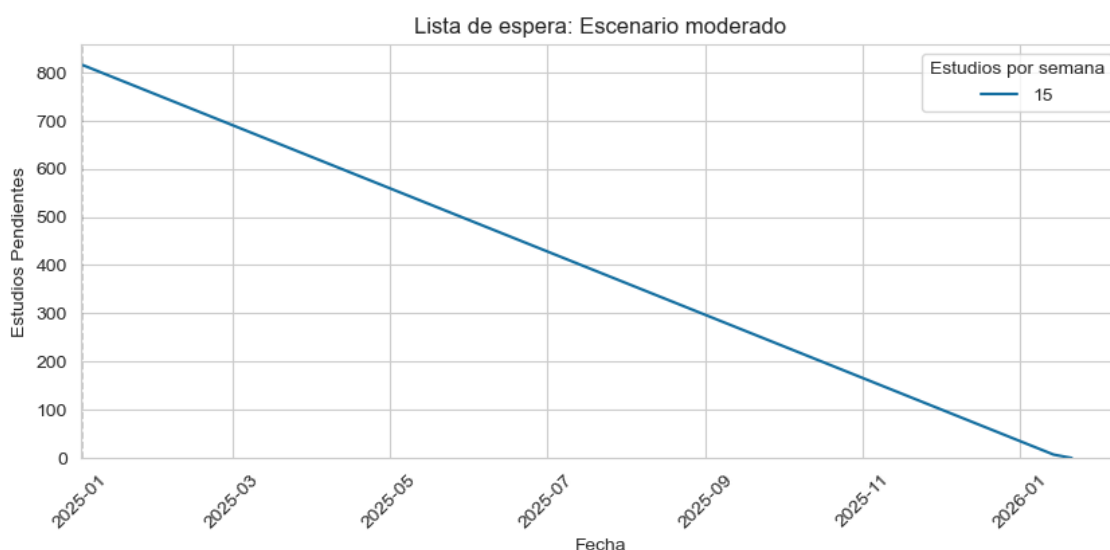
Las tasas de servicio con las que deberíamos cumplir para poder mitigar la lista de espera dentro del año y medio serían:

- Lista de espera: 14 estudios por semana
- Nueva demanda: 13 estudios por semana

5.1 Primer escenario: Quirófano largo

Como escenario óptimo, sumar otro quirófano largo a la demanda actual (28 estudios por semana) permitiría cumplir con la nueva demanda y utilizar la capacidad remanente (15 estudios) terminar con la lista de espera en el período de 1 año aproximadamente para el escenario moderado como se puede ver en la siguiente Figura 22.

Figura 22. Evolución de estudios pendientes en el primer escenario tomando el número de estudios del escenario moderado.



Con esta solución mantenemos un nivel de servicio mayor al 95% como se puede ver en la Figura 20 de la Sección 4.4.

5.2 Segundo escenario: Quirófano corto

En caso de no poder sumar un quirófano largo (14 estudios) al servicio actual, revisamos cómo podríamos dar solución sumando un quirófano corto (7 estudios).

. Las tasas de servicio necesarias seguirán siendo las mismas:

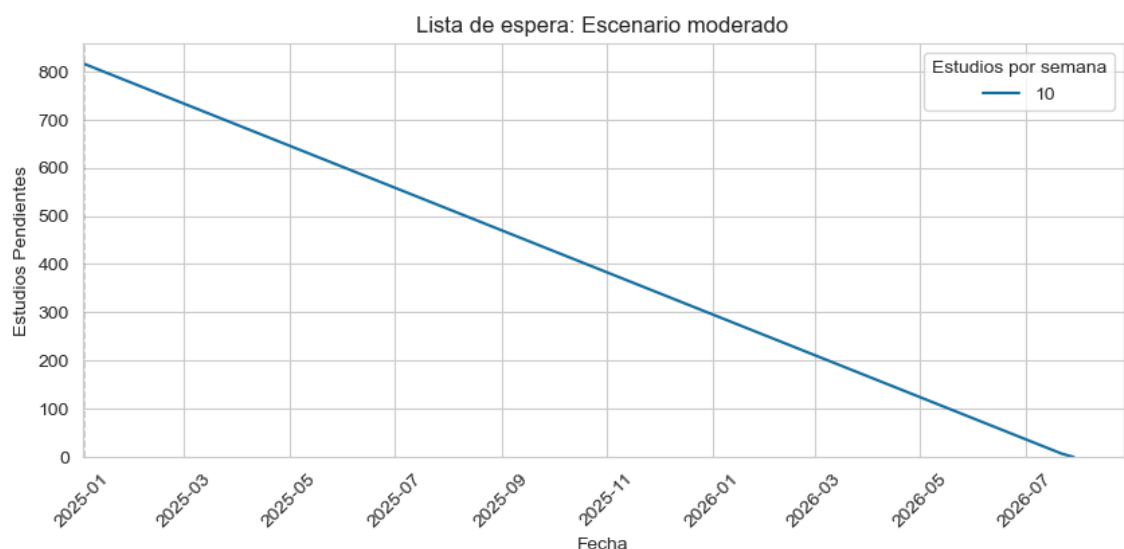
- Lista de espera: 14 estudios por semana
- Nueva demanda: 13 estudios por semana

Con el agregado de este quirófano corto tendríamos 21 estudios. Como no podremos cumplir con ambas demandas, destinamos los recursos de la siguiente manera para equilibrar:

- Lista de espera: 10 estudios por semana
- Nueva demanda: 11 estudios por semana

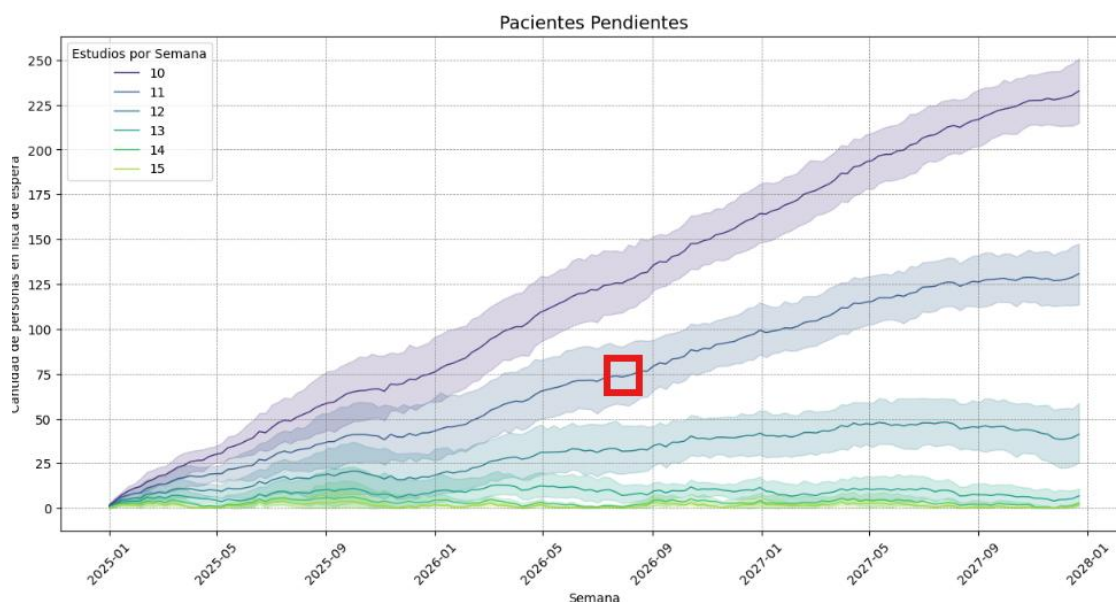
Esta primera etapa genera dos efectos: Por un lado, la lista de espera tardará más en desaparecer. Como podemos observar en la Figura 23, necesitaremos de 1 año y 8 meses para mitigar la lista pendiente.

Figura 23. Evolución de lista de espera en segundo escenario según número de estudios del escenario moderado.



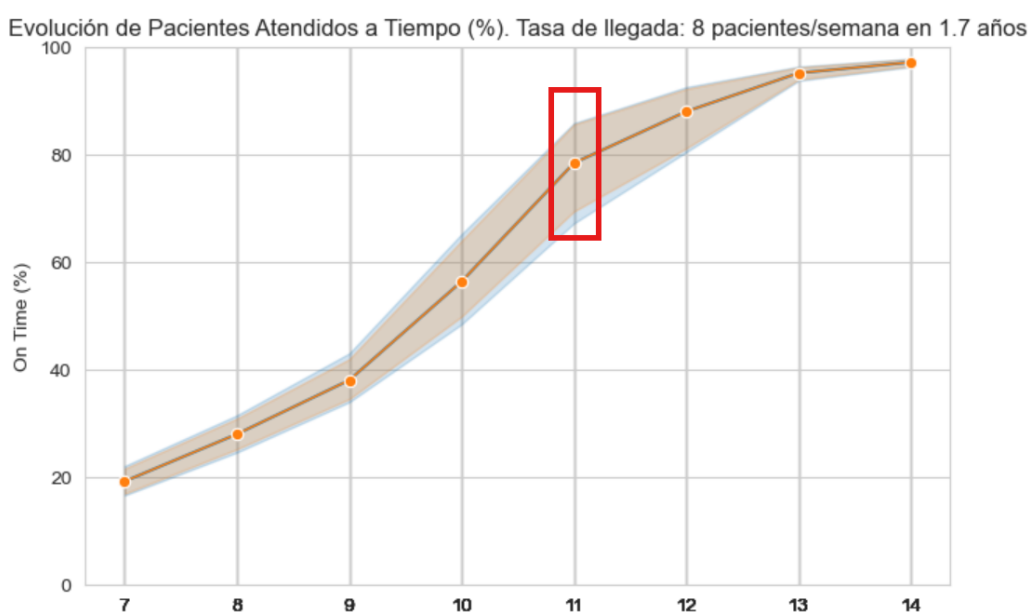
En ese tiempo, la nueva demanda habrá generado una cola esperada de 75 personas, como se puede ver a continuación en la Figura 24, donde tenemos en cuenta 11 estudios por semana por 1 año y 8 meses.

Figura 24. Pacientes pendientes en segundo escenario generados por nueva demanda.



El nivel de servicio para la nueva demanda será del 80% con un límite inferior del 70% y uno superior de 85% como se puede ver en la Figura 25. Consideramos un nivel aceptable para este período dada la limitación en los recursos.

Figura 25. Nivel de servicio en segundo escenario para nueva demanda.

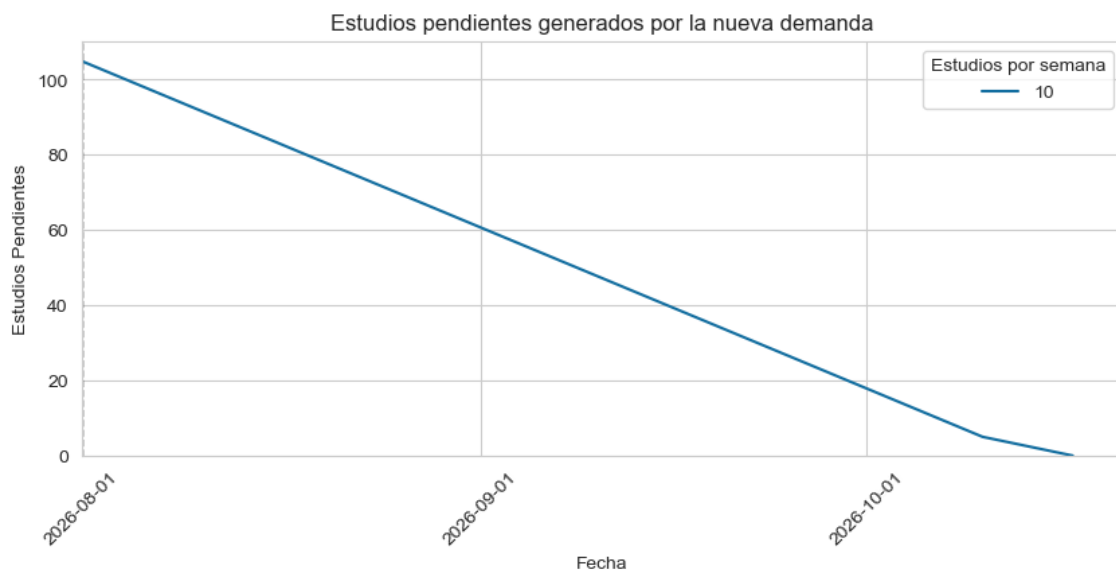


En una segunda etapa, debemos licuar esa nueva lista generada, redefiniendo nuevamente la tasa de servicio como:

- Lista de espera: 8 estudios por semana
- Nueva demanda: 13 estudios por semana

Con el objetivo de no generar más listas de espera utilizaremos 13 estudios por semana para la atención de la nueva demanda y 8 para la lista de espera generada. Las 75 personas pendientes representan aproximadamente 105 estudios. Como vemos en la Figura 26 licuar esta demanda nos llevará casi 4 meses, terminando los remanentes a finales de Octubre de 2026.

Figura 26. Estudios pendientes generados por la nueva demanda en el segundo escenario.



6. Conclusiones

En esta tesis logramos modelar el funcionamiento de un hospital mediante una simulación que requirió varias iteraciones en Python, pero que demostró ser robusta frente a nuevos escenarios. Se identificó una forma de reducir la carga operativa del personal mediante la incorporación de técnicas de Machine Learning para asistir en la priorización de pacientes, y se analizaron distintas estrategias para afrontar la demanda esperada.

Los resultados fueron satisfactorios: se observó que, con la incorporación de pocos turnos adicionales de quirófano, sería posible alcanzar en el corto o mediano plazo un nivel de servicio que podría constituir un caso de estudio para otros hospitales, e incluso para otras áreas que enfrentan problemáticas similares.

A continuación, se presentan algunas reflexiones y recomendaciones que resumen el trabajo realizado

6.1 Cuando la IA Decide: La Importancia de la Ética en Salud

Al desarrollar un modelo de ML para asignar prioridades en la atención médica una de las decisiones clave es qué variables incluir en el entrenamiento. Un aspecto que podría llamar la atención es que en nuestro modelo no hemos considerado:

1. El género de los pacientes: esto no es un descuido, sino una decisión informada basada en la orientación de los profesionales de la salud, quienes nos explicaron que el sexo biológico no es un factor determinante en la priorización de estos estudios.

2. Cobertura en salud: como hemos mencionado (Ver Sección 2.3.1) aproximadamente el 40% de nuestros pacientes tienen alguna cobertura en salud distinta al subsistema público que les podría permitir acceder a los procedimientos en otros centros. Si bien no hay una directiva institucional de excluir a estos pacientes y el sistema público tiene la obligación de atender a todos aquellos que lo soliciten, este factor puede ser un sesgo a la hora de elegir pacientes de la lista.

3. Evaluación de riesgo anestesiológico y cardiológico: habitualmente aquellos pacientes que presentan riesgos elevados requieren condiciones especiales para la realización de los procedimientos, por ejemplo: disponibilidad de cama en unidad cerrada posterior al procedimiento, disponibilidad de hemoderivados, suspensión de medicación habitual, etc. Estos factores acarrearán una planificación más extensa y aumentan la probabilidad de suspensión del estudio. Por lo expuesto estos pacientes también tienen mayor riesgo de sesgo de elección.

Para garantizar mayor equidad en la asignación de prioridades decidimos excluir estas variables del modelo.

Este enfoque nos lleva a reflexionar sobre un tema fundamental en la aplicación de ML en medicina: la ética y la responsabilidad en el uso de modelos predictivos. Si bien estas herramientas pueden optimizar la toma de decisiones y mejorar la eficiencia en la atención, también es crucial asegurarse de que sean justas y libres de sesgos.

Los sesgos pueden surgir de muchas formas, pero una de las más comunes es a través de los datos con los que entrenamos los modelos. Si el conjunto de datos refleja desigualdades previas el modelo puede aprender y perpetuar esas diferencias. Por ello, es fundamental evaluar de manera crítica qué información utilizamos, asegurándonos de que las decisiones automatizadas sean imparciales y alineadas con principios de equidad y justicia.

En este proyecto, hemos sido especialmente cuidadosos en diseñar un modelo que respete estos principios, priorizando únicamente las variables médicamente relevantes y evitando aquellos factores que podrían introducir sesgos innecesarios. De esta manera, buscamos que la tecnología sea un apoyo para la medicina, sin comprometer la equidad en la atención de los pacientes.

6.2 Recomendaciones

El sistema de salud pública enfrenta retos de gestión para poder cumplir con la demanda esperada. La falta de digitalización y estandarización en los procesos es un factor agravante y muchas veces complica el foco y el orden de cómo debería atenderse a los pacientes. En esta tesis hemos desarrollado algunas ideas para ayudar a tomar decisiones de gestión.

Notamos como herramienta fundamental para la gestión la estandarización y priorización de los pacientes de forma sistemática para ahorrar tiempo y esfuerzos, además una métrica para medir el nivel de atención del servicio en el área. Esto brinda mayor visibilidad y una mejora de la sinergia para el trabajo en equipo.

Creemos que la opción de 28 estudios semanales sería el escenario ideal que permitiría evacuar la lista de espera en 1 año, sin generar nuevas colas. Como segunda opción 21 estudios semanales permite en una primera etapa licuar la lista de espera en 16 meses, generando una cola de 75 pacientes (equivalentes a 105 estudios) que podría atenderse en una segunda etapa en los siguientes 4 meses, mitigando las colas en 1 año y 10 meses.

Con lo expuesto concluimos que con las estrategias planteadas el servicio de gastroenterología podría regularizar su atención y mantener un nivel de servicio del 95% para el primer escenario y un 80% para el segundo.

La generación de un ranking, además de ayudar en la priorización médica, podría dar una mejor visibilidad a los pacientes sobre su lugar en la lista de espera (aclarando que el mismo puede moverse dependiendo del nivel de urgencia) acotando el intervalo de incertidumbre.

Como punto a considerar, en esta tesis no desarrollamos el manejo de los pacientes internados. Como ya hemos comentado (Ver Sección 1.2), estos representan un volumen considerable de la demanda habitual. Si los evaluamos mensualmente representan 10 estudios, lo que equivale a 2.5 estudios semanales. Si bien esto puede desarrollarse en un trabajo posterior, proponemos un margen de seguridad de tres estudios más semanales destinados a estos pacientes.

7. Bibliografía

1. Báscolo, E., Houghton, N., & Del Riego, A. (2020). Leveraging household survey data to measure barriers to health services access in the Americas. *Revista Panamericana de Salud Publica/Pan American Journal of Public Health*, 44. <https://doi.org/10.26633/RPSP.2020.100>
2. Bruce, Peter., Bruce, Andrew., & Gedeck, Peter. (2020). *Practical Statistics for Data Scientists, 2nd Edition*. O'Reilly Media, Inc.
3. Cámara Argentina de Especialidades Medicinales. (2024, March 6). *Impacto de la pandemia COVID-19 sobre el sistema de salud Argentino*. <https://www.caeme.org.ar/Impacto-de-La-Pandemia-Covid-19-Sobre-El-Sistema-de-Salud-Argentino/>.
4. Consenso Argentino de evaluación de riesgo cardiovascular en cirugía no cardíaca. (2016). *Revista Argentina de Cardiología*, 84(1). <https://www.sac.org.ar/>
5. Cotton, P. B. ., & Williams, C. B. . (2003). *Practical gastrointestinal endoscopy: the fundamentals*. (5th ed.). Blackwell Pub.
6. Gutacker, N., Siciliani, L., & Cookson, R. (2016). Waiting time prioritisation: Evidence from England. *Social Science and Medicine*, 159, 140–151. <https://doi.org/10.1016/j.socscimed.2016.05.007>
7. Hillier, F., & Lieberman, G. (2010). *INTRODUCCIÓN A LA INVESTIGACIÓN DE OPERACIONES* (E. C. Zuñiga Gutierrez & M. Rocha Martinez, Eds.; J. E. Murrieta & C. R. Cordero Pedraza, Trans.; 9th ed., pp. 708–757). McGRAW-HILL.
8. Kaarboe, O., & Carlsen, F. (2014). Waiting times and socioeconomic status. Evidence from Norway. *Health Economics (United Kingdom)*, 23(1), 93–107. <https://doi.org/10.1002/hec.2904>
9. Kirschbaum, A., & Yonamine, K. (2020). Indicadores de calidad para videocolonoscopia en tamizaje de cáncer colorrectal. In <https://www.argentina.gob.ar/sites/default/files/bancos/2020-10/2020-10-27-indicadores-calidad-para-vcc-en-tamizaje-ccr.pdf>. <https://www.argentina.gob.ar/sites/default/files/bancos/2020-10/2020-10-27-indicadores-calidad-para-vcc-en-tamizaje-ccr.pdf>
10. McKinney, W. (n.d.). *Python for Data Analysis*.
11. Ministerio de Salud de Argentina. (2022, December 30). *Por las consecuencias de la pandemia, el gobierno extendió la emergencia sanitaria por un año*. <https://www.argentina.gob.ar/noticias/por-las-consecuencias-de-la-pandemia-el-gobierno-extendio-la-emergencia-sanitaria-por-un-a%C3%B1o>
12. Organización Mundial de la Salud. (2023, July 11). *Cáncer colorrectal*. https://www.who.int/es/news-room/fact-sheets/detail/colorectal-cancer#_msoanchor_1
13. Organización Mundial de la Salud. (2024, October 1). *Envejecimiento y Salud*. <https://www.who.int/es/news-room/fact-sheets/detail/ageing-and-health>
14. Organización Panamericana de la Salud. (2020, June 17). *La COVID-19 afectó el funcionamiento de los servicios de salud para enfermedades no transmisibles en las Américas*. <https://www.paho.org/es/noticias/17-6-2020-covid-19-afecto-funcionamiento-servicios-salud-para-enfermedades-no>

15. Paterson, W. G., Depew, W. T., Paré, P., Petrunia, D., Switzer, C., Veldhuyzen Van Zanten, S. J., & Daniels Msc, S. (2006). Canadian consensus on medically acceptable wait times for digestive health care. In *Can J Gastroenterol* (Vol. 20, Issue 6).
16. Sammour, T., Schoeman, M., Moore, J. W., Thomas, M. L., Moorcraft, L., & Andrews, J. M. (2021). Clearing a colonoscopy waiting list: how we did it. In *ANZ Journal of Surgery* (Vol. 91, Issues 1–2, pp. 10–12). Blackwell Publishing. <https://doi.org/10.1111/ans.15942>
17. Shim, H. G., Gupta, A., Fu, A., Flores, R., Simmons, R., Steinberg, J., Guerson-Gil, A., Liao, Y., Yang, J., LaComb, J. F., D'Souza, L. S., Monzur, F., Li, E., & Guillaume, A. (2024). A Quality Improvement Study on Colonoscopy Wait Times in Underinsured Patients Following the COVID-19 Pandemic. *Clinical and Translational Gastroenterology*, 15(9), e1. <https://doi.org/10.14309/ctg.0000000000000730>
18. *SimPy Documentation Release 4.0.0 Team SimPy*. (2020).
19. VanderPlas, J. (n.d.). *Python Data Science Handbook*. www.allitebooks.com
20. Yonamine, K., & Kirschbaum, A. (2022). *Recomendaciones para el tamizaje organizado de Cáncer Colorrectal en población de riesgo promedio en Argentina*. (1st ed.).

8. Anexo

Estructura de un endoscopio. Imagen obtenida de “Practical gastrointestinal endoscopy: the fundamentals”.

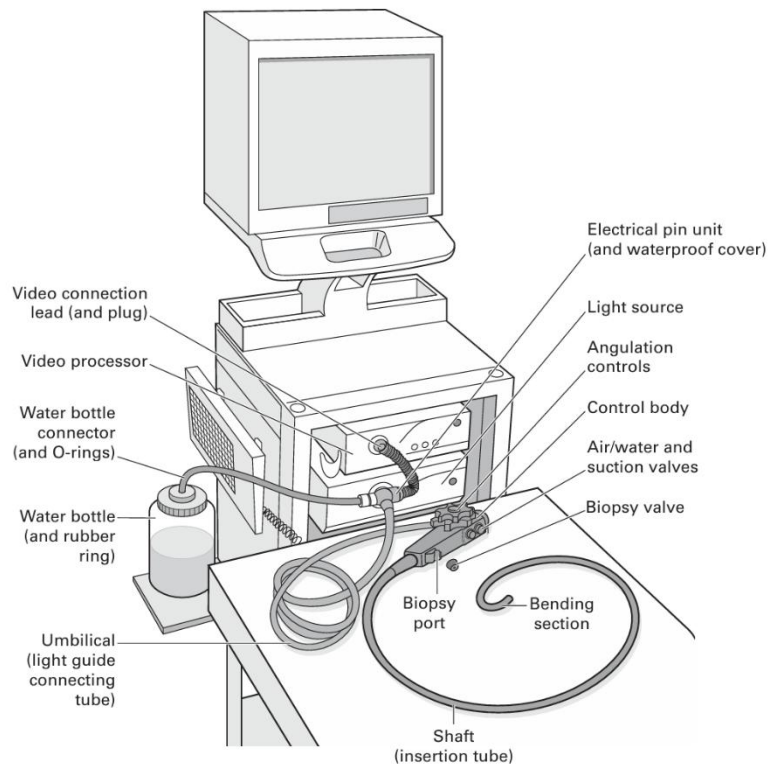


Fig. 2.1 Endoscope system.