



PROYECTO FINAL DE BIOINGENIERÍA

**Aplicación de Inteligencia Artificial (IA) al
servicio de la salud en cáncer: su aplicación
en tumores del tracto gastrointestinal (TGI)
y su valoración costo-efectiva.**

Alumnas:

Andrea Belen, ERBETTI 60348

Olivia, SANGUINETTI 60118

Tutora:

Dra. María Teresa POMBO

Co-Tutor:

Lic. Adrián PÉREZ

Fecha de entrega: 16 de Noviembre de 2023

Dedicatoria

Este proyecto final de carrera es el resultado del apoyo inquebrantable brindado por el ITBA, el invaluable acompañamiento de nuestros queridos profesores y el constante estímulo otorgado por la admirable comunidad científica. Este trabajo está dedicado a cada bioingeniero que compone esta comunidad, cuyo impacto se extiende más allá de estas páginas, influenciando el presente y el futuro del conocimiento.

Agradecimientos

El desarrollo de este trabajo fue un logro alcanzado gracias al apoyo incondicional de nuestros seres queridos. Expresamos nuestra profunda gratitud a nuestras familias y amigos que estuvieron a nuestro lado a lo largo de este proceso, ya que fueron ellos quienes nos ayudaron a sobrepasar los momentos desafiantes, y nos acompañaron en los momentos más movilizantes de la carrera también. Su constante aliento fue fundamental en cada paso. Queremos reconocer a todos nuestros compañeros con los que compartimos éstos últimos increíbles cinco años, especialmente a Ana, Martina, Julieta y Nacho. Sin ustedes ni Esterno, nada hubiese sido posible. Agradecemos a la Licenciada Belén Villegas por su generosa colaboración en la digitalización de las imágenes. También queremos agradecerle a Adrián Etchevarne por su inestimable paciencia y por estar siempre disponible para ayudarnos a resolver todas nuestras inquietudes con respecto a la informática. Asimismo, extendemos nuestro agradecimiento infinito al Licenciado Adrián Pérez, cuya dedicación e inspiración fue imprescindible en este proyecto. Por último, pero no menos importante, agradecemos enormemente a la Doctora María Teresa Pombo, quien fue la fuerza impulsora detrás de este proyecto final. Fue verdaderamente un honor trabajar junto a ella.

Índice

Resumen	7
Objetivos	9
Marco Teórico	10
1. Generalidades	10
1.1. Tracto Gastrointestinal (TGI)	10
1.2. Cáncer Colorrectal y Cáncer Gástrico	11
1.3. Factores de riesgo	12
1.4. Incidencia	12
1.5. Clasificación Histológica	13
1.6. Clasificación Molecular	15
1.7. Diagnóstico	16
1.8. Generalidades de la Inestabilidad Microsatelital (IMS)	16
1.8.1. ¿Qué son los Microsatélites?	16
1.8.2. Inestabilidad Microsatelital (IMS)	17
1.8.3. Mecanismo de Reparación del ADN (MMR)	17
1.8.4. Técnicas de Detección	19
1.8.4.1. Inmunohistoquímica (IHQ)	20
1.8.4.2. Reacción en Cadena de la Polimerasa (PCR)	21
1.8.4.3. Secuenciación de Nueva Generación (NGS)	22
1.9. Laboratorios y centros de salud en países en vías de desarrollo	23
1.10. Tratamientos	23
1.10.1. Inmunoterapia	24

1.10.1.1. Vías de Regulación del Sistema Inmunológico	24
1.10.1.2. PD-1/PD-L1 Inhibidores	26
1.11. Importancia de detección IMS	26
2. Aprendizaje Profundo (Deep Learning)	27
2.1. Redes Neuronales	28
2.2. Clasificación de Capas de Redes Neuronales	29
2.2.1. Capas Densas	29
2.2.2. Capas Convolucionales	29
2.2.3. Capas de Atención	29
2.2.4. Capas Pooling	30
2.2.5. Capas de Regularización	30
2.3. Importancia del Deep Learning en la actualidad	31
3. Visión Computacional (Computer Vision)	32
3.1. Vision Transformers (ViT)	32
3.2. Mecanismo de Atención	33
4. Proceso de Aprendizaje en Modelo de Deep Learning	33
4.1. Inicialización de Pesos	34
4.2. Propagación hacia Adelante	34
4.3. Función de Costo	34
4.4. Retropropagación del Error	35
4.5. Tipos de Aprendizaje	35
4.5.1. Aprendizaje Supervisado	36
4.5.2. Aprendizaje No Supervisado	36
4.5.3. Aprendizaje Débilmente Supervisado	37
5. Evaluación de Modelos	37

5.1. Métricas	38
6. Aprendizaje por Transferencia (Transfer Learning)	39
Materiales y Métodos	41
1. Base de Datos	41
1.1. Descripción de Imágenes Histológicas	41
1.2. Recopilación de Datos	41
1.3. Preprocesamiento de Imágenes	43
2. Líneas de desarrollo del Modelo para la Clasificación de MSI-H/MSS	46
2.1. Modelo ViT	46
2.2. Clustering-constrained Attention Multiple Instance Learning (CLAM)	47
3. Desarrollo de Modelo para la Clasificación MSI-H/MSS	48
3.1. Preprocesamiento de las WSI	48
3.1.1. Segmentación de la WSI	48
3.1.2. Creación de Parches	49
3.1.3. Extracción de Características	51
3.2. Estructura del Modelo para la Clasificación de MSI-H/MSS	51
3.2.1. Compresión de Representaciones de Nivel de Parche	51
3.2.2. Función de Agregación Basada en Atención	52
3.2.3. Instance-Level Clustering	52
3.2.4. Inferencia y Predicción	53
3.3. Entrenamiento y Evaluación del Modelo	53
3.3.1. Parámetros del Entrenamiento del Modelo	56
3.4. Interpretación de la Predicción del Modelo a través del Mapa de Atención	56
4. Interfaz	58
5. Hardware y Software	59

5.1. Servidor	59
5.2. Entorno de Desarrollo	60
Resultados	61
1. Evaluación Interna con WSI de Colon, Recto y Estómago del TCGA	61
2. Evaluación Externa con Colon y Recto	62
3. Ejemplos de Resultados Discordantes	64
Discusión	65
1. Evaluación Interna	65
2. Evaluación Externa	66
3. Limitaciones Generales	68
4. Datos Bibliográficos	68
Conclusión	70
Glosario de Contracciones	71
Referencias	73

Resumen

La detección de células tumorales del tracto gastrointestinal (TGI) en imágenes histológicas es posible por un patólogo entrenado utilizando microscopía de luz. Sin embargo, distinguir entre tumores proficientes (pMMR *Proficient Mismatch Repair*) o deficientes (dMMR *Deficient Mismatch Repair*), ambos biomarcadores subrogantes de tumores con inestabilidad microsatelital (IMS), no es posible por esta metodología y requiere de la utilización de técnicas más sofisticadas como la inmunohistoquímica (IHQ) para su diferenciación. La IHQ es un método de detección antígeno específico, mediante la utilización de anticuerpos mono y policlonales para hacer visible una proteína. Requiere de médicos patólogos capacitados, trabajo metódico y minucioso, y es una técnica costosa, lo que dificulta su realización en los laboratorios de patología en todo el territorio del país. Como contrapartida, identificar a los pacientes con cáncer del TGI con tumores con IMS en función de su estadio clínico, permite diferenciar un subgrupo de pacientes con pronóstico y predicciones distintas.

El objetivo de este Proyecto Final es generar un algoritmo basado en inteligencia artificial (IA) capaz de predecir IMS en TGI a partir de imágenes histológicas en microscopía de luz. La idea es que el mismo pueda ser utilizado como una herramienta accesible y costo-efectiva como soporte a la toma de decisiones. Esto es crucial ya que los pacientes podrían obtener el tratamiento acorde a su subtipo tumoral (Medicina de Precisión), entre otras ventajas que aporta este driver oncogénico.

El algoritmo desarrollado se basa en una arquitectura llamada CLAM (*Clustering-Constrained Attention Multiple Instance Learning*), diseñada para clasificar imágenes de láminas histológicas completas digitalizadas (WSIs, por su sigla en inglés *Whole Slide Images*) utilizando aprendizaje profundo y enfoque en regiones de interés (ROIs). El modelo tiene la capacidad de enfocar su atención en las áreas más relevantes de los WSIs, lo que le permite predecir si la muestra presenta IMS a partir de un corte teñido con hematoxilina-eosina (H&E).

Para el entrenamiento del modelo, se utilizó un conjunto de datos compuesto por 108 WSIs de cáncer colorrectal (CCR), de los cuales 47 eran tumores inestables o MSI-H (del inglés *High Microsatellite Instability*) y 61 tumores estables o MSS (del inglés *Microsatellite Stability*). Estos datos se obtuvieron del repositorio público *The Cancer Genome Atlas* (TCGA). Luego, el modelo fue evaluado en imágenes de este conjunto de datos junto con imágenes del TCGA de cáncer de estómago (CE). Se obtuvieron resultados notables, incluyendo un 76,9 % de sensibilidad, un 77,6 % de exactitud, 78,1 % de especificidad y 76,3 % área bajo la curva ROC (AUC-ROC, del inglés *Area Under the Curve ROC*).

A su vez, también se realizó una validación externa en un conjunto de muestras de archivo, que se consideró como cohorte externa. Esta validación tuvo como objetivo comprobar la capacidad de adaptación del algoritmo en el contexto real en el que se implementaría, en beneficio de los pacientes. Los resultados obtenidos fueron un 84,2 % de sensibilidad, un 56,9 % de exactitud, 67,3 % AUC-ROC y 43,6 % de especificidad. Estos últimos resultados resaltan la necesidad impera-

tiva de optimizar diversos procedimientos de laboratorio, especialmente en lo referente a la etapa preanalítica, con el objetivo de perfeccionar el tratamiento de las muestras y, por ende, fortalecer la eficacia del algoritmo en futuras implementaciones.

Como componente esencial de este proyecto, se desarrolló una interfaz gráfica que integra el modelo y simplifica significativamente la experiencia de profesionales de la salud y patólogos que la utilicen. Permite cargar imágenes y realizar evaluaciones de manera intuitiva, sin requerir acceso al código subyacente; está diseñada para que sea una herramienta rápida, de muy bajo costo, universal, y de fácil utilización. Constituye un primer criterio para subtipificar tumores MSI-H y MSS; se puede visualizar en la Figura 1.

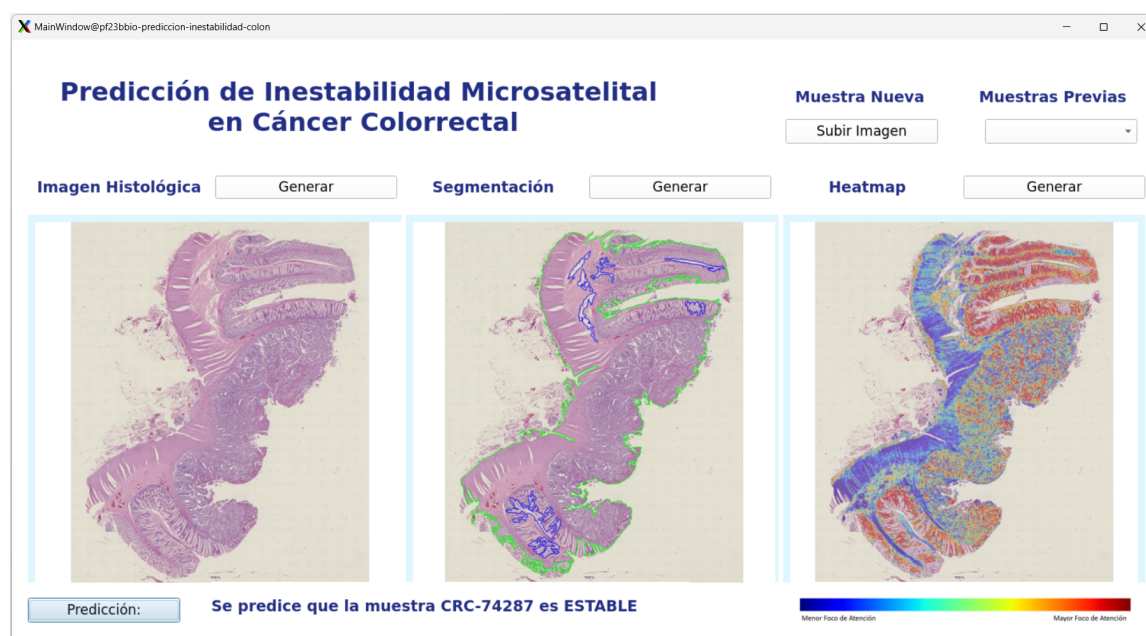


Figura 1: Interfaz gráfica diseñada para la implementación universal del modelo desarrollado.

Debido a la frecuencia de neoplasias del TGI, minimizar la dificultad para la pesquisa de IMS de estos tumores, resulta de gran importancia para los pacientes que los padecen y los profesionales de la salud a cargo. La minuciosa tarea de armonizar IA en el cáncer, representó un desafío que implicará el estudio de mayor número de pacientes para confirmar los resultados obtenidos. En conclusión, este Proyecto Final constituye un avance significativo orientado a acelerar los tiempos de diagnóstico y mejorar la disponibilidad para la detección de la IMS en el cáncer del TGI. Los resultados destacan la promisoría potencialidad de esta tecnología en términos de costo-efectividad para los laboratorios de diagnóstico, lo que inspira a futuros investigadores a impulsar el desarrollo de nuevas plataformas de IA en el cáncer.

Objetivos

El propósito central de esta tesis es el diseño y desarrollo de una herramienta avanzada de patología digital, enfocada específicamente en la identificación de IMS en adenocarcinomas de colon, recto y estómago. El objetivo de máxima es alcanzar una sensibilidad superior al 85 % en la detección, permitiendo así su aplicación práctica en laboratorios de anatomía patológica. Este nivel de sensibilidad no solo asegura la fiabilidad de la herramienta, sino que también la hace adecuada para el uso en entornos clínicos reales. Este trabajo también aspira a generar una herramienta que podría surgir como una alternativa costo-efectiva y rápida a procedimientos de laboratorio más costosos y complejos que se utilizan hoy en día para el diagnóstico. Se espera que sea particularmente beneficiosa en contextos con recursos limitados.

En conclusión, el proyecto busca enriquecer el diagnóstico, impulsando la integración de la IA para una mayor precisión y comprensión en la oncología. Este enfoque innovador abre la puerta a un nuevo paradigma, allanando el camino hacia una práctica médica más rápida y eficiente.

Marco Teórico

1. Generalidades

1.1. Tracto Gastrointestinal (TGI)

El TGI, también conocido como el sistema digestivo, es un sistema anatómico y funcional crucial de la anatomía de un individuo. Sus principales funciones incluyen la digestión de alimentos, la absorción de nutrientes y la eliminación de desechos. Consta de varias partes anatómicas interconectadas, cada una de las cuales desempeña un papel específico en estos procesos y comienza desde la boca hasta el ano.

Este proyecto final, debido a la importancia pronóstica y predictiva de la IMS en CCR, estará basado principalmente en colon y recto. Se hace una breve mención del estómago, debido a que para enriquecer las series de casos del testeo, también se utilizaron muestras de esta localización. Es importante mencionar que el sistema de reparación del ADN (MMR) es idéntico en cualquier órgano de un individuo y se trata de un mecanismo altamente conservado en eucariontes, por lo que la IMS también resulta un mecanismo celular universal.

El colon, también conocido como intestino grueso, es una estructura en forma de tubo que constituye la parte final del sistema digestivo. Como se puede observar en la Figura 2, se divide en varias secciones, que incluyen el ciego, el colon ascendente, el colon transversal, el colon descendente y el colon sigmoide [1].

El recto es la sección final del intestino grueso y se encuentra justo antes del ano. Este área del tracto digestivo actúa como un depósito temporal de heces antes de su eliminación [2].

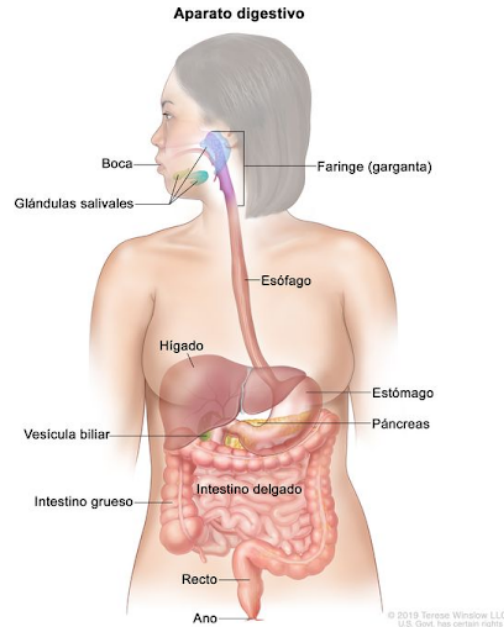


Figura 2: Anatomía del TGI.

El estómago es un órgano en forma de J que se encuentra en el sistema digestivo y desempeña un papel fundamental en la digestión de los alimentos. Se encuentra en la cavidad abdominal, en la parte superior del abdomen, debajo de las costillas. Está situado por debajo del esófago, órgano que lleva los alimentos desde la garganta al estómago y el intestino delgado, por donde los alimentos digeridos pasan al resto del sistema digestivo. Este es un órgano vital para la digestión, donde los alimentos se mezclan con enzimas y jugos gástricos antes de continuar su viaje a través del TGI. Su estructura y función son claves para mantener la absorción adecuada de nutrientes.

1.2. Cáncer Colorrectal y Cáncer Gástrico

El CCR, objeto de estudio de esta tesis, es una de las neoplasias malignas más comunes en el mundo [3]. Se origina en el colon o el recto, generalmente a partir de pólipos que son crecimientos anormales en el revestimiento interno de estas estructuras. A medida que estos pólipos desarrollan distintas variantes genéticas, dan lugar a tumores malignos. La comprensión de la anatomía del colon y el recto es fundamental para abordar el CCR y desarrollar estrategias de diagnóstico y tratamiento efectivas.

El CE, también conocido como cáncer gástrico, es otra de las patologías de gran relevancia en el campo oncológico. Este tipo de cáncer se origina en las células del estómago y su forma más frecuente es el adenocarcinoma. Puede afectar tanto la parte superior (CE proximal) como la parte inferior (CE distal) del órgano. Al igual que en el caso del CCR, la investigación y el desarrollo de métodos de predicción y diagnóstico temprano son esenciales para abordar eficazmente el CE y

mejorar las tasas de supervivencia.

1.3. Factores de riesgo

En la mayoría de los casos (75 % aproximadamente) [4], el CCR se desarrolla en individuos sin antecedentes familiares. Estos casos se consideran esporádicos, lo que significa que las variantes genéticas que causan la enfermedad se desarrollan al azar después del nacimiento. El resto de los CCR se produce en personas con riesgo adicional debido a antecedentes familiares con CCR, presencia de adenomas únicos o múltiples, enfermedad inflamatoria intestinal u otras patologías hereditarias [5].

Lo más frecuente es que no se conozcan exactamente las causas que provocan el CCR. Sin embargo, existen factores que pueden aumentar el riesgo de desarrollarlo. Entre ellos se destaca la edad, ya que la mayoría de los casos se presenta en personas mayores de 50 años. A su vez, también juega un rol importante el estilo de vida: el sedentarismo, la obesidad y el tabaquismo.

En cuanto al CE, los factores de riesgo se asocian con mayor frecuencia al sexo masculino y la infección por *Helicobacter Pylorii*. Es mayor la tendencia en individuos de más de 60 años de edad y parece asociarse a ingesta de ahumados y carnes saladas, como también a anemia perniciosa, entre otras causas. También, existen con baja frecuencia, síndromes hereditarios asociados a neoplasias gástricas y poliposis [6].

1.4. Incidencia

La relevancia del CCR radica principalmente en su alta incidencia tanto en Argentina, como en el mundo. A nivel global, esta enfermedad representa una preocupación significativa ya que, según datos de la Organización Mundial de la Salud (OMS) [3], el CCR ocupa el tercer lugar como causa de muerte, afectando tanto a países desarrollados como también aquellos en vías de desarrollo. En el año 2020, se diagnosticaron más de 1,9 millones de nuevos casos, resultando en más de 930.000 fallecidos. Además, las proyecciones indican que hacia el año 2040, la carga de casos nuevos de CCR aumentará significativamente a 3,2 millones por año, representando un incremento del 63 %. Las muertes asociadas a esta enfermedad alcanzarán 1,6 millones anuales, lo cual implica un aumento del 73 %.

De manera similar, el CE también representa una carga significativa a nivel mundial. Según los datos del Observatorio Global del Cáncer (GLOBOCAN), fue el quinto cáncer más frecuente en 2020, con un millón de casos nuevos (1.089.103), representando el 6 % de todos los tumores a nivel mundial y causando 768.793 muertes [7].

A nivel nacional, mientras que el CCR se ubica en segundo lugar en cuanto a su incidencia, el CE también presenta cifras considerables. Según las estimaciones realizadas por el GLOBOCAN, de la Agencia Internacional de Investigación sobre Cáncer (IARC, por sus siglas en inglés *International*

Agency for Research on Cancer), de los 130.878 nuevos casos de cáncer en ambos sexos que se registraron en Argentina en el año 2020, 15.895 casos fueron de CCR, representando un 12.1 % del total. Por otro lado, el CE ocupó el noveno lugar con un 3,1 % de incidencia, equivalente a 4003 casos para ambos sexos [8]. Estas cifras subrayan la importancia de la detección temprana y el tratamiento adecuado de estas enfermedades en el país. Los números mencionados se reflejan en la Figura 3 que se presenta a continuación.

SITIO TUMORAL	AMBOS SEXOS		VARONES		MUJERES	
	Casos	%	Casos	%	Casos	%
Mama	22.024	16,8	-	-	22.024	32,1
Colon-recto	15.895	12,1	8.493	13,6	7.402	10,8
Pulmón	12.110	9,3	7.738	12,4	4.372	6,4
Próstata	11.686	8,9	11.686	18,7	-	-
Riñón	5.093	3,9	3.370	5,4	1.723	2,5
Páncreas	5.026	3,8	2.357	3,8	2.669	3,9
Cervicouterino	4.583	3,5	-	-	4.583	6,7
Tiroides	4.106	3,1	669	1,1	3.437	5,0
Estomago	4.003	3,1	2.549	4,1	1.454	2,1
Vejiga	3.785	2,9	2.955	4,7	830	1,2
Linfoma No-Hodgkin	3.557	2,7	1.988	3,2	1.569	2,3
Leucemia	3.234	2,5	1.797	2,9	1.437	2,1
Cuerpo de útero	2.455	1,9	-	-	2.455	3,6
Hígado	2.437	1,9	1.467	2,4	970	1,4
Ovario	2.199	1,7	-	-	2.199	3,2
Testículo	2.047	1,6	2.047	3,3	-	-
Esófago	1.993	1,5	1.388	2,2	605	0,9
Encéfalo y otros SNC	1.831	1,4	935	1,5	896	1,3
Otros	9.688	7,4	6.001	15,5	3.687	5,4
Total	130.878	100	62.327	100	68.551	100

Figura 3: Distribución absoluta y relativa de casos incidentes de cáncer estimados por la IARC para Argentina en 2020, según localizaciones tumorales más frecuentes y sexo (N=130.878).

1.5. Clasificación Histológica

La gran mayoría de los CCRs son adenocarcinomas. Los mismos se originan a partir de células que producen mucosidad para lubricar el revestimiento interno del colon y el recto. En adición a los adenocarcinomas, existen otro tipos de tumores que, si bien son mucho menos comunes, pueden desarrollarse en estas localizaciones. Entre ellos se encuentran los tumores carcinoides, que se originan en células especializadas que producen hormonas en el intestino, los tumores estromales gastrointestinales (GIST, del inglés *Gastrointestinal Stromal Tumor*), los linfomas y los sarcomas. [9] [2]. Los distintos tipos tumorales se pueden visualizar en la Figura 4 a continuación.

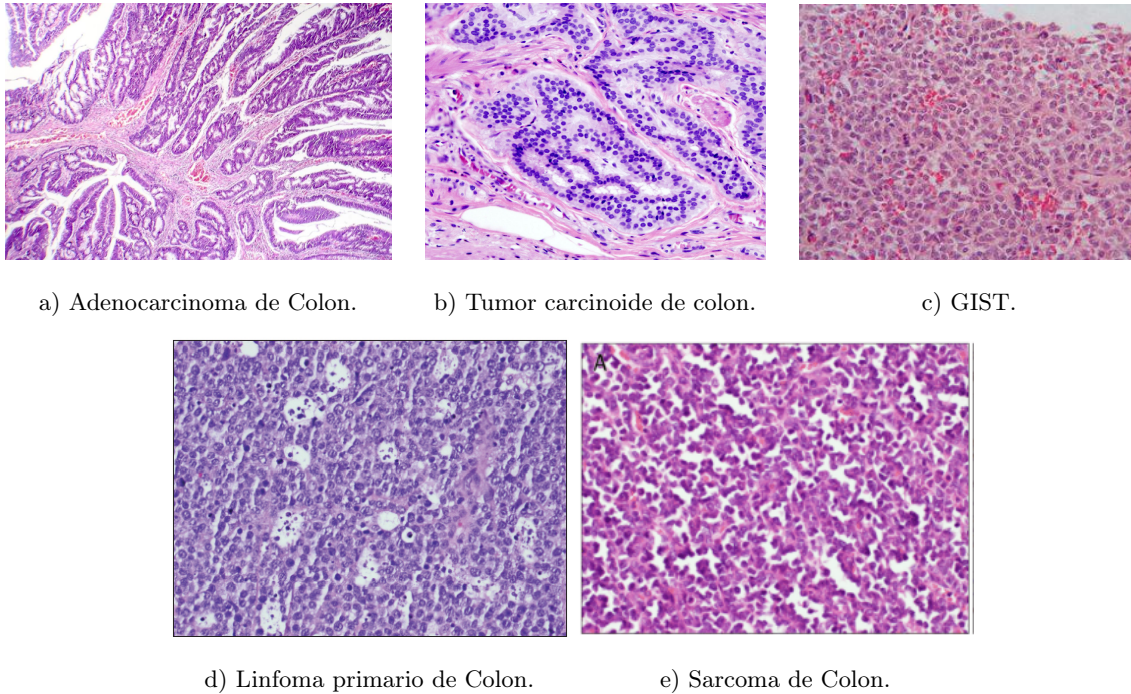
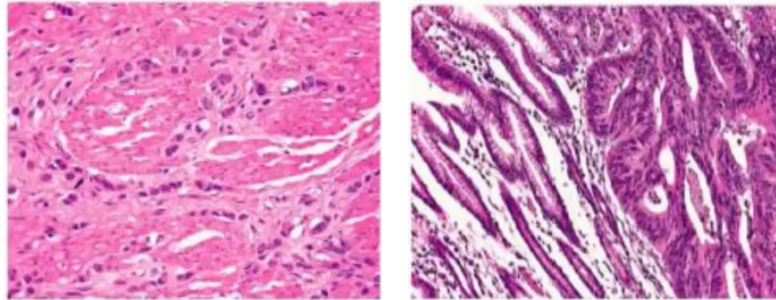


Figura 4: Imágenes histológicas teñidas con H&E obtenidas mediante microscopía de luz.

La clasificación histológica del CE, según la clasificación de Lauren, divide los tumores gástricos en dos categorías principales: el tipo intestinal y el tipo difuso. Los tumores del tipo intestinal, también conocidos como adenocarcinomas, suelen localizarse en el estómago distal y están frecuentemente vinculados a la infección por *Helicobacter pylori*, presentando generalmente un pronóstico más favorable. Por otro lado, los tumores del tipo difuso, conocidos como carcinomas de células en anillo de sello, tienden a no estar localizados, pero suelen ser encontrados más comúnmente en el estómago proximal. Estos se caracterizan por un peor pronóstico y una mayor tendencia a generar metástasis [10]. En la Figura 5 se pueden ver ejemplos de lo mencionado.



a) Variante Difuso o Adenocarcinoma b) Variante Inestinal o Carcinoma

Figura 5: Imágenes histológicas de tumores gástricos teñidas con H&E obtenidas mediante microscopía de luz.

1.6. Clasificación Molecular

El CCR presenta una clasificación molecular que proporciona información crucial para comprender mejor su biología y guiar decisiones terapéuticas. Una de las principales categorías es el CCR con IMS por errores del MMR. La IMS puede estar presente en alrededor del 15 % de los cánceres colorrectales y tiene importantes implicaciones clínicas, ya que estos tumores son más sensibles a la inmunoterapia [11].

Otra categoría es el CCR con inestabilidad cromosómica (CIN, del inglés *Chromosome Instability*), que se caracteriza por variaciones en el número de cromosomas en las células tumorales. Esto puede resultar en una amplia gama de cambios genéticos y cromosómicos, lo que a menudo se asocia con mayor agresividad del tumor. También existen mutaciones específicas en genes como TP53 (p53), BRAF, KRAS y otros, que también desempeñan un papel en la clasificación molecular. Estas mutaciones influyen en el comportamiento del cáncer y son indicativas de la respuesta del tumor a tratamientos específicos. Además de estas categorías principales, algunos cánceres colorrectales se consideran "molecularmente indiferenciados". Esto significa que no encajan claramente en ninguna de las categorías anteriores y requieren un análisis molecular más detallado para comprender su biología y respuesta potencial al tratamiento dirigido, si lo hubiera.

La clasificación molecular es fundamental para la atención personalizada de los pacientes con CCR. Más aún, permite identificar aquellos individuos que podrían beneficiarse con inmunoterapia, algunos de los cuales tendrían respuestas realmente notorias, según las últimas publicaciones en el mundo [5].

En cuanto al CE, también juega un rol importante la IMS, pero existen otros subgrupos moleculares como el asociado al Virus de Epstein-Barr (EBV, del inglés *Epstein-Barr Virus*), el vinculado a CIN y el genómicamente estable (GS, del inglés *Genomically Stable*). En tumores avanzados irresecables, otros biomarcadores moleculares adquieren relevancia clínica como HER2Neu y PDL-1.

1.7. Diagnóstico

El diagnóstico de CCR y de CE es un proceso complejo que implica varias etapas para obtener una evaluación integral de la enfermedad y orientar las decisiones de tratamiento. El proceso de diagnóstico comienza con la obtención de una muestra de tejido tumoral a través de una endoscopia y biopsia. Los patólogos examinan minuciosamente esta muestra bajo el microscopio para determinar el tipo histológico del cáncer. Esta caracterización histológica es fundamental para clasificar el tumor. Paralelamente al examen histológico, se realiza un análisis molecular para detectar mutaciones en genes específicos mediante diferentes técnicas genómicas. Se buscan variantes en genes como TP53, BRAF, KRAS y RAS, en el caso de CCR que intervienen como drivers oncogénicos, y poseen valor pronóstico y predictivo [5]. En CE como ya fue mencionado, es de gran importancia si resulta irresecable, el valor de HER2, PDL-1 y de IMS.

La evaluación de IMS es un paso crítico en el diagnóstico de estos tumores, ya que los pacientes serán divididos en dos subgrupos moleculares que tendrán distinta chance terapéutica, podrán tener un mejor pronóstico y recibirán o no asesoramiento genético. Con la información recopilada de los exámenes histológicos y moleculares, se establece un diagnóstico completo que incluye el tipo histológico específico del cáncer, cualquier mutación genética presente y la categorización de IMS para la toma de decisiones terapéuticas y el seguimiento oncológico.

1.8. Generalidades de la Inestabilidad Microsatelital (IMS)

1.8.1. ¿Qué son los Microsatélites?

Los microsatélites, también conocidos como repeticiones en tándem cortas (STRs, del inglés *Short Tandem Repeats*), son secuencias de ADN repetidas (1-6 pares de bases) dispersas en todo el genoma; tanto en regiones codificantes como no codificantes. Estos pueden ocurrir en regiones intragénicas o intergénicas, y representan aproximadamente el 3 % del genoma humano. Debido a su estructura repetitiva, los microsatélites son propensos a una alta tasa de mutación. A continuación se observa una ilustración de un microsatélite en la Figura 6.



Figura 6: Microsatélites, ejemplo de un di-nucleótido A-C(n).

1.8.2. Inestabilidad Microsatelital (IMS)

IMS se define como la presencia de secuencias de ADN repetitivas de tamaños alternativos, que no están presentes en el ADN germinativo correspondiente, debido a una expansión o contracción de los microsatélites [12]. IMS se encuentra en alrededor del 15% de todos los cánceres de colon, el 3% se encuentra asociado al síndrome de Lynch (patología de condición hereditaria) y el porcentaje restante se debe a apariciones esporádicas. La determinación del estado de IMS en el CCR tiene implicaciones pronósticas y terapéuticas. El IMS se ha asociado con un mejor pronóstico ya que los microsatélites inestables son altamente inmunogénicos, por lo que la terapia que activa el sistema inmunológico puede tener efectos importantes en estos tumores.

Bajo el foco de la IMS, existen tres tipos de categorización tumoral: MSI-H, Inestabilidad Microsatelital Baja (MSI-L, del inglés *Low Microsatellite Instability*) y MSS. Los tumores MSI-H son aquellos que presentan una alta tasa de mutación, una eficiente presentación de neoantígenos y la infiltración de células T efectoras, y es por este motivo que se los denominan tumores “calientes”. Además, atraen a células inmunosupresoras como las células supresoras derivadas de la médula ósea (MDSCs) y las células T reguladoras (Tregs). Estos factores explican las altas tasas de respuesta observadas con los inhibidores de puntos de control inmunológico (ICIs). Notablemente, los tumores MSI-H tienen un mayor potencial angiogénico y una mayor densidad microvascular, lo que puede facilitar el reconocimiento por parte del sistema inmunológico, y se traduce en altas tasas de respuesta y supervivencia con la inmunoterapia. Estos se diferencian notablemente de los tumores MSS o MSI-L, ya que estos presentan una cantidad muy baja de neoantígenos en sus membranas, lo que hace que el sistema inmunológico no pueda reclutar a las células inmunitarias al tumor, lo que es un obstáculo fundamental para la eficacia de la inmunoterapia. A estos tumores se los denomina tumores “fríos”.

1.8.3. Mecanismo de Reparación del ADN (MMR)

El MMR corrige inserciones, deleciones y desajustes de bases erróneas generadas durante la replicación y recombinación del ADN. Esta vía de reparación está altamente conservada desde las bacterias hasta los humanos y protege la integridad del genoma [12]. En condiciones óptimas, MMR tiene la capacidad de identificar un error, eliminar la porción de la cadena recién sintetizada que contiene dicho error y rellenar el espacio resultante con las bases correctas.

MMR cuenta con cuatro proteínas esenciales: MLH1, PMS2, MSH2, y MSH6 que funcionan mediante la formación de dímeros; MLH1 con PMS2, y MSH2 con MSH6. Un tumor se considera dMMR cuando al menos una de las proteínas involucradas no se expresa. Esto se contrasta con pMMR, donde todas las proteínas del sistema de reparación se expresan y funcionan correctamente.

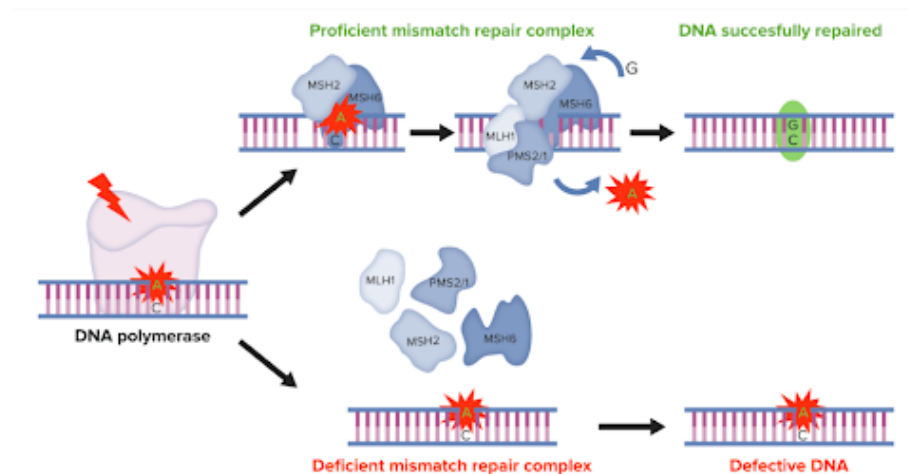
Los defectos en los genes que codifican estas proteínas pueden ocurrir como eventos somáticos o de línea germinal. Los defectos de línea germinal son mutaciones en los genes que codifican para las proteínas de reparación de MMR. La mutación de alguno de estos genes hace que la proteína correspondiente no se produzca, por lo que el sistema de reparación se vuelve deficiente. Este tipo

de variaciones genéticas son hereditarias, la más conocida es el Síndrome de Lynch, que representa aproximadamente el 2-3 % de los casos de cáncer de colon.

Por otro lado, los defectos somáticos son aquellos que causan el IMS esporádico, el tipo más frecuente. Este tipo de defecto ocurre por un fenómeno epigenético: la metilación del ADN. La metilación es un mecanismo mediante el cual una célula silencia de manera permanente a los genes. Esto se observa en el caso de la proteína MLH-1, que se inactiva debido a la hipermetilación de la secuencia promotora del gen. Este tipo de evento esporádico es común entre las personas de edad avanzada, probablemente como resultado de factores ambientales, dietéticos y el envejecimiento.

En la Figura 7, se podrán observar dos imágenes que ilustran el MMR en condiciones proficientes y deficientes.

Es fundamental resaltar que IMS y un MMR defectuoso no son lo mismo. Sin embargo, como existe una fuerte correlación entre ambos fenómenos (>95 %) [13], al referirnos a dMMR consideraremos IMS y pMMR será subrogante en nuestro trabajo de MSS, como todo lo referido en la bibliografía.



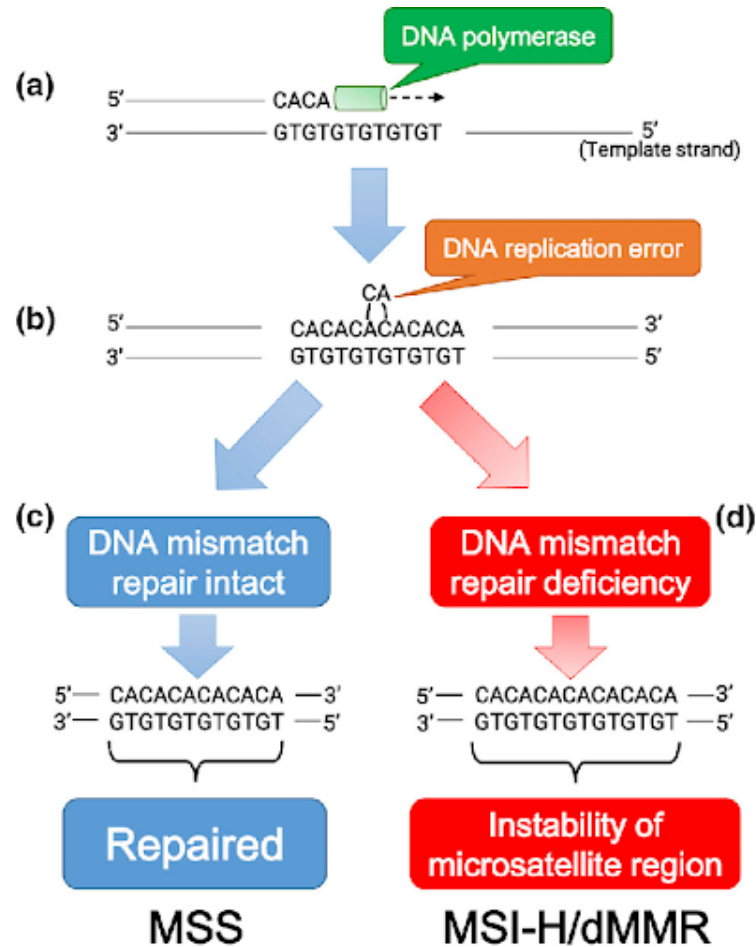


Figura 7: Dos diagramas esquemáticos de las diferencias entre un MMR deficiente y proficient.

1.8.4. Técnicas de Detección

Para el diagnóstico de IMS se utilizan comúnmente tres métodos: la IHQ, la reacción en cadena de la polimerasa (PCR, del inglés *Polymerase Chain Reaction*) seguida de análisis de fragmentos y la secuenciación genómica, exclusivamente a través de la técnica NGS. NGS e IHQ son métodos de alta sensibilidad y especificidad. NGS tiene una sensibilidad del 95,8 % y una especificidad del 99,4 %, con un valor predictivo positivo del 94,5 % y un valor predictivo negativo del 99,2 % en comparación con la PCR. En el caso de la IHQ, la sensibilidad varía del 80,8 % al 100,0 %, y la especificidad del 80,5 % al 91,9 % [14].

1.8.4.1. Inmunohistoquímica (IHQ)

La IHQ es una técnica que se basa en la unión de anticuerpos a sus antígenos objetivo para localizar y detectar proteínas o antígenos específicos en células y tejidos. Se suele realizar en tejido embebido en parafina y fijado en formalina, que tiene la ventaja de ser fácil de almacenar.

Este método requiere de un anticuerpo primario, secundario, y un reactivo que pueda unirse a alguno de ellos para que sirva como agente de detección. En general, el anticuerpo primario suele ser monoclonal ya que son específicos a un único epítipo. Para poder visualizar la interacción antígeno-anticuerpo, el anticuerpo primario o secundario debe estar marcado. En el método directo, el anticuerpo primario es aquel marcado, pero éste método no es comúnmente utilizado debido a la falta de amplificación de la señal. En el método indirecto, el anticuerpo secundario es marcado, lo que permite que la señal sea amplificada con facilidad y se pueda utilizar una variedad más amplia de anticuerpos primarios [15]. Se puede visualizar la esquematización de la técnica en la Figura 8.

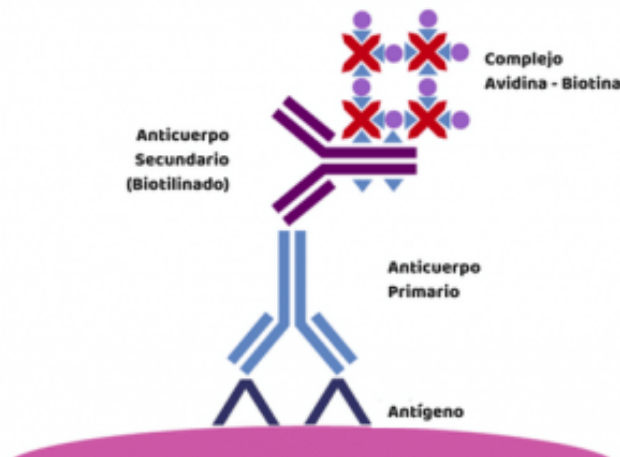


Figura 8: Representación esquemática de la técnica de IHQ.

En este método, los anticuerpos contra las proteínas MMR (MLH1, MSH2, PMS2 y MSH6) proporcionan información sobre la funcionalidad del sistema MMR. El análisis de IHQ con los anticuerpos de PMS2 y MSH6 puede detectar la mayoría de las anomalías en los genes correspondientes, así como las mutaciones en MLH1 y MSH2. Sin embargo, la prueba de IHQ con los anticuerpos MLH1 o MSH2 puede detectar una fracción de las anomalías de MLH1 o MSH2, pero no todas. Por lo tanto, el análisis de IHQ con los anticuerpos MSH6 y PMS2 tiene un mayor potencial diagnóstico que el análisis con los anticuerpos MLH1 y MSH2 [12]. Se presentan las IHQ positivas y negativas para cada una de las proteínas en la Figura 9.

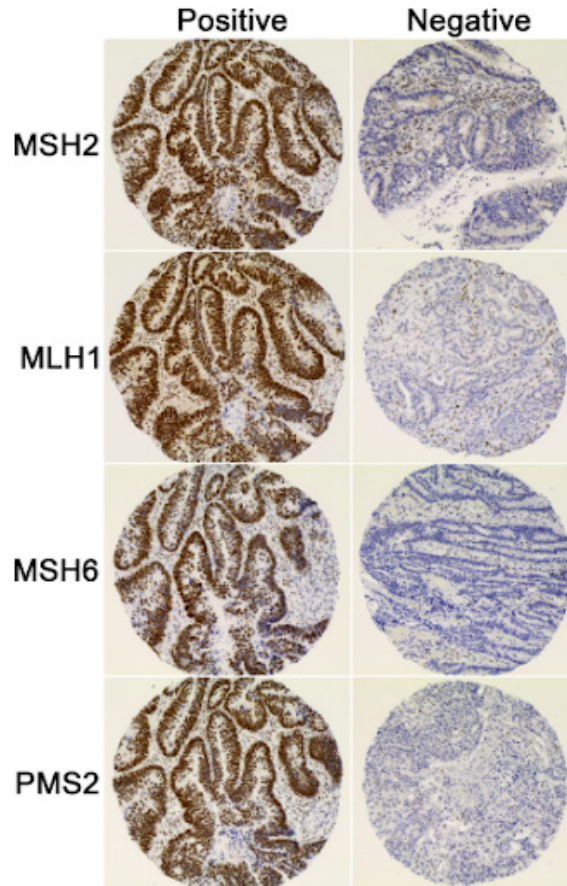


Figura 9: Análisis inmunohistoquímico de las cuatro proteínas de reparación del ADN [16].

1.8.4.2. Reacción en Cadena de la Polimerasa (PCR)

El método de Reacción en Cadena de la Polimerasa (PCR) se emplea para amplificar segmentos específicos del ADN, permitiendo la generación de múltiples copias a partir de una secuencia particular. Posteriormente, el análisis de los fragmentos de ADN amplificados se lleva a cabo mediante la metodología de Sanger. Este proceso, fundado en la técnica de secuenciación Sanger, implica la identificación y caracterización de los fragmentos de ADN amplificados para discernir su composición y secuencia.

El principio subyacente de este método radica en la evaluación de variaciones en la longitud de marcadores de microsatélites particulares presentes en células tumorales, en comparación con las células normales. Esta técnica posibilita la detección de alteraciones en la estructura del ADN, ofreciendo información relevante para el análisis genético y la identificación de posibles anomalías asociadas a procesos patológicos. Se pueden observar los pasos de la técnica en la Figura 10.

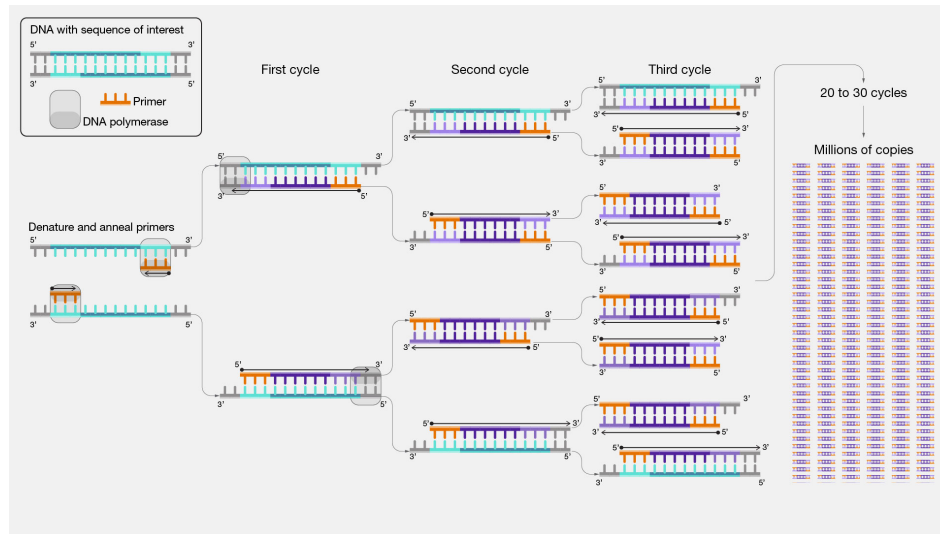


Figura 10: Pasos de la técnica PCR.

1.8.4.3. Secuenciación de Nueva Generación (NGS)

La NGS es una tecnología utilizada para determinar la secuencia de nucleótidos en el ADN o en el ácido ribonucleico (ARN) de una muestra de manera rápida y eficiente. Ha reemplazado los métodos de secuenciación genómica antiguos, debido a su capacidad para secuenciar múltiples fragmentos de ADN o ARN de forma paralela. A diferencia de los métodos anteriores, la NGS permite una visualización a nivel nucleotídico, pudiendo así identificar con precisión si existen bases faltantes o excedentes.

Esta técnica es parte del campo conocido como Biología Molecular, ya que permite el análisis profundo de los nucleótidos presentes en la región del genoma que se está analizando. Al ser una técnica que requiere trabajo riguroso y patólogos con mucha experiencia, como también es elevadamente costosa, es utilizada en casos muy puntuales. En el análisis de IMS, se puede utilizar cuando un caso es no concluyente por IHQ para confirmar que el tumor es verdaderamente MSI-H. También, se utiliza cuando se predice que el MMR es deficiente (dMMR), debido a una falta de tinción de alguna de las cuatro proteínas por IHQ, para la confirmación del diagnóstico. Los pasos de ésta técnica se pueden visualizar en la Figura 11.

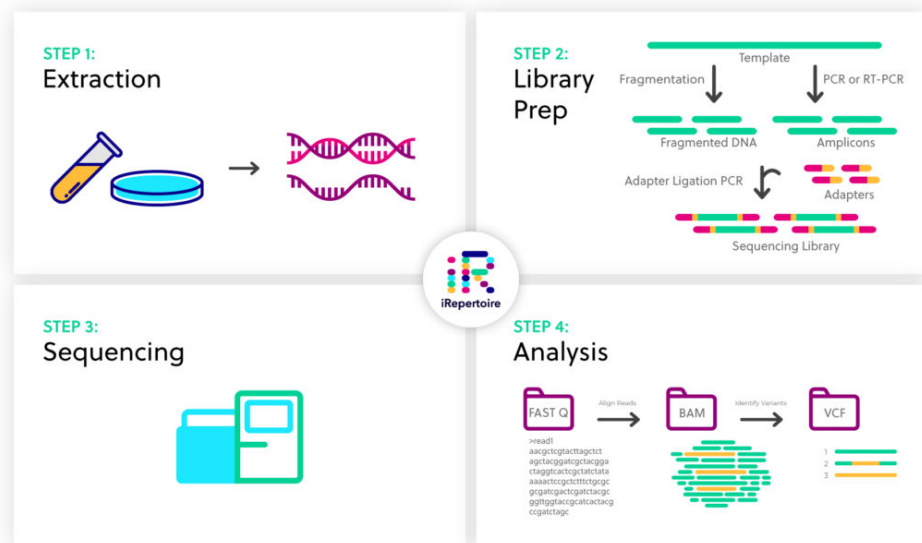


Figura 11: Pasos de la técnica NGS.

1.9. Laboratorios y Centros de Salud en Países en Vías de Desarrollo

No todos los laboratorios tienen la disponibilidad de tener herramientas como IHQ, PCR u otras herramientas de biología molecular para poder realizar medicina de precisión, por lo que resulta dificultoso separar a los pacientes en subgrupos moleculares y entre ellos, los individuos con MSI-H. Vemos en nuestra tecnología, la IA, enorme potencialidad en términos de costo-efectividad para los laboratorios de diagnóstico, ya que el desarrollo de estas plataformas permite con muy bajos costos, pocos recursos humanos y sencilla infraestructura, obtener datos fidedignos y altamente calificados, que permite al equipo de salud, el tratamiento de un paciente acorde a su patología.

1.10. Tratamientos

Hay varios tratamientos disponibles para los pacientes con CCR y CE. En primer lugar, dependiendo del estadio, se puede realizar cirugía. En caso de que el cáncer se encuentre en un estadio temprano, la extirpación del mismo puede ser con márgenes negativos. Esto involucra una extirpación del cáncer y una pequeña cantidad del tejido sano que lo rodea, incluyendo los ganglios linfáticos cercanos para determinar si éstos también se encuentran afectados.

Uno de los tratamientos más comunes para el CCR y CE es la quimioterapia. Este tratamiento utiliza medicamentos conocidos como agentes quimioterapéuticos, que interfieren en el ciclo de división celular, entre otros mecanismos, para detener el crecimiento y proliferación de las células tumorales. Estos medicamentos pueden ser administrados por vía intravenosa, oral o en forma de tabletas o cápsulas. La quimioterapia generalmente se utiliza en conjunto con la radioterapia,

que utiliza rayos X de alta energía u otros tipos de radiación para ayudar la destrucción de las células cancerosas. El gran problema de este tratamiento es que puede tener efectos secundarios incómodos para el paciente, como fatiga, náuseas, vómitos, caída de cabello, supresión del sistema inmunológico, etc. [17].

La inmunoterapia es un tipo de tratamiento desarrollado en la última década, en el intento de clasificar especialmente el CCR en base a su perfil molecular, buscando posibles blancos terapéuticos [18].

1.10.1. Inmunoterapia

1.10.1.1. Vías de Regulación del Sistema Inmunológico

La presencia de linfocitos T en el lecho tumoral del CCR se ha asociado con resultados favorables. Los linfocitos T son células del sistema inmunológico que desempeñan un papel crucial en la respuesta inmunitaria contra las células anómalas, incluidas las células tumorales.

La interacción entre los linfocitos T y las células en general, tanto tumorales como normales, se realiza a través de los receptores de células T (TCR), que reconocen complejos de péptidos presentados por las moléculas del complejo principal de histocompatibilidad (MHC) en la superficie celular. La mera identificación de los complejos péptido-MHC de clase I por el TCR no es suficiente para la activación completa de las células T. Las vías de señalización entre el TCR y el complejo péptido-MHC están reguladas por señales coestimuladoras (que promueven la activación) o co-inhibitorias (que suprimen la activación)[18].

Para lograr la activación completa de las células T y desencadenar una respuesta inmunitaria eficaz, se necesitan señales coestimuladoras. Estas señales son esenciales para potenciar la actividad de las células T y asegurar una respuesta inmunitaria adecuada. La regulación de estas vías de señalización está mediada por puntos de control inmunológico. Estos puntos de control, pueden desempeñar un papel de coestimuladores, promoviendo la activación celular, o de co-inhibidores, suprimiendo la activación excesiva y manteniendo la homeostasis del sistema inmunitario. Están compuestos por proteínas, siendo las más conocidas PD1 y CTLA-4. El delicado equilibrio entre estas señales es fundamental para garantizar una respuesta inmunitaria efectiva sin provocar respuestas autoinmunes no deseadas.

Las células tumorales tienen la capacidad de unirse a una célula T a través de las uniones PD1 y PDL-1, como también B7 y CTLA-4. Estas interacciones tienen como resultado que el sistema inmunitario no reconozca a estas células como amenazas malignas, evadiendo así la respuesta inmunitaria. Esto se puede analizar de manera gráfica en la Figura 12.

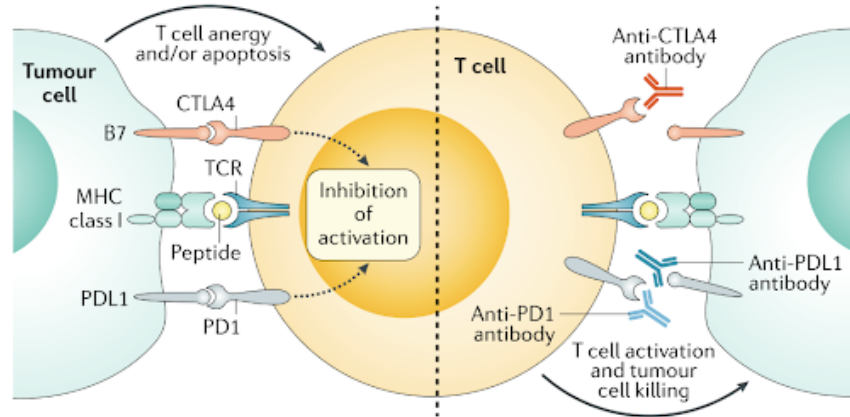


Figura 12: Representación esquemática de los ICI.

La inmunoterapia es una forma de tratamiento que utiliza el sistema inmunitario del individuo para combatir enfermedades, como el cáncer. En el contexto del cáncer, los ICI son una forma de inmunoterapia que bloquea las interacciones, que frena la respuesta inmunitaria. En el caso de tumores dMMR-MSI-H en estadios II y III, donde hay una mayor presentación de péptidos mutados en el complejo principal de histocompatibilidad, se facilita el reconocimiento de los tumores por los ICI. Estos ICI, dirigidos a proteínas como CTLA-4, PD-1 o PD-L1, bloquean las interacciones específicas entre las células del sistema inmunitario y las células tumorales, lo que evita la disfunción y muerte celular, favoreciendo la activación y potenciando la capacidad del sistema inmunitario para destruir las células tumorales. Este tipo de terapia fue aceptado por la Administración de Drogas y Alimentos de los Estados Unidos (FDA, del inglés *Food and Drug Administration*) en 2017. Esto permite desencadenar la activación de las células inmunitarias T CD8+, células T cooperadoras tipo 1 (TH1) CD4+ y la secreción de IFN, y su posterior infiltración en el entorno del tumor [18].

Los tumores MSS o MSI-L se distinguen significativamente de los tumores MSI-H al presentar una cantidad notablemente menor de neoantígenos en sus membranas. Esta característica limita la capacidad del sistema inmunitario para atraer células inmunitarias hacia el tumor, lo que representa un desafío fundamental para que la inmunoterapia sea efectiva en estos casos [14]. Esto puede ser visualizado de manera gráfica en la Figura 13 que se presenta a continuación.

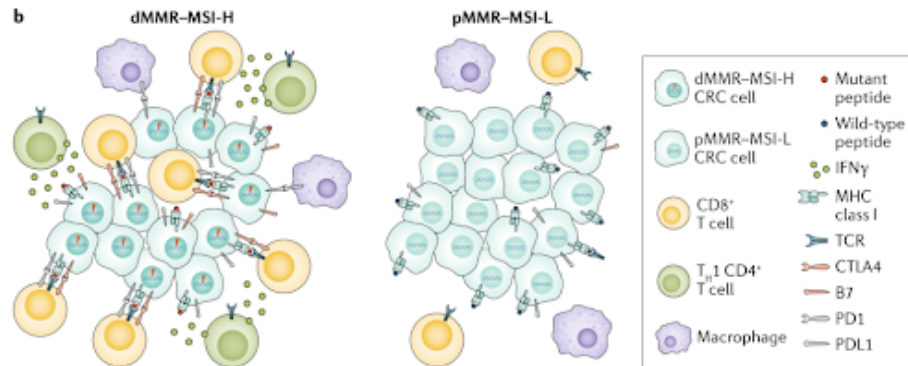


Figura 13: Reconocimiento del sistema inmunitario a los tumores dMMR-MSI-H y pMMR-MSI-L.

1.10.1.2. PD-1/PD-L1 Inhibidores

El PD-1 es un receptor inhibidor en la superficie celular que se expresa en linfocitos T activados, células pro-B y macrófagos. El PD-L1 es el ligando para el PD-1 y se expresa comúnmente en células tumorales, pero también en células supresoras del sistema inmunológico, como los Tregs, los macrófagos tipo M2 y las MDSCs.

En condiciones fisiológicas, la unión de PD-L1 al receptor PD-1 produce un mecanismo de retroalimentación negativa que evita que el sistema inmunológico del huésped “agreda” sus propias células. Sin embargo, las células neoplásicas utilizan este mecanismo, a través de la expresión de PD-L1, para evadir la vigilancia inmunológica. Se sabe que la unión PD-1- PD-L1 desactiva la función efectora de los linfocitos T al disminuir su activación y proliferación. Los inhibidores de PD-1 (como nivolumab o pembrolizumab) y de PD-L1 (como durvalumab, atezolizumab o avelumab), bloquean la interacción inhibitoria, permitiendo así que los linfocitos T citotóxicos ejerzan su efecto antitumoral, especialmente en tumores MSI-H con una alta cantidad de neoantígenos y una respuesta inmunológica mediada por linfocitos T ya establecida. Esto explica por qué los inhibidores de PD-1, han mostrado altas tasas de respuesta y supervivencia en el CCR con MSI-H.

1.11. Importancia de detección IMS

La detección de IMS juega un papel esencial en el manejo del CCR, ya que proporciona información crucial para personalizar las estrategias de tratamiento y pronosticar los resultados de los pacientes con CCR. Como se mencionó previamente, los tumores que presentan IMS, tienen una mayor sensibilidad a la inmunoterapia, lo que representa un avance significativo en el tratamiento del CCR. Según datos bibliográficos, los pacientes con CCR MSI-H tratados con inmunoterapia tuvieron una tasa de supervivencia a 5 años del 83 %, en comparación con el 64 % para los pacientes tratados con quimioterapia [19]. Esto subraya la importancia de la detección precisa de IMS, para brindar a los pacientes el tratamiento más efectivo y mejor tolerado.

Por otro lado, los pacientes con tumores MSS generalmente obtienen beneficios de los regímenes de quimioterapia convencionales. Esta distinción en las estrategias de tratamiento según su diagnóstico en cuanto a la IMS, resalta la importancia de la detección precisa, para orientar la atención al paciente y optimizar los resultados del tratamiento.

Además de las ventajas terapéuticas, el estado de IMS también tiene implicancias pronósticas ya que pacientes MSI-H tienden a evolucionar más favorablemente que aquellos tumores MSS. Estudios han demostrado que la supervivencia de los pacientes con MSI-H es hasta un 15 % más alta que la de aquellos con tumores MSS [20].

En resumen, la detección de IMS ha cambiado el paradigma del tratamiento del CCR, permitiendo a los médicos ofrecer a los pacientes el tratamiento más adecuado para su enfermedad, lo que a su vez mejora las posibilidades de éxito terapéutico y la calidad de vida de los pacientes con CCR. Este nuevo planteamiento para tratar estos pacientes, con valor pronóstico y predictivo, nos desafía a generar herramientas de diagnóstico, sencillas y accesibles para separar grupos de individuos que se beneficiarán con inmunoterapia y es lo que impulsa a este trabajo a generar un algoritmo de IA para pacientes con este tipo de neoplasias, altamente beneficiarios de estas nuevas ventajas terapéuticas.

2. Aprendizaje Profundo (Deep Learning)

La IA representa un campo multidisciplinario que busca desarrollar sistemas capaces de realizar tareas que, cuando son realizadas por seres humanos, requieren de inteligencia. Una de las ramas más destacadas de la IA es el Aprendizaje Profundo, también conocido como *Deep Learning*. La subdivisión de éste último puede visualizarse en la Figura 14 a continuación.

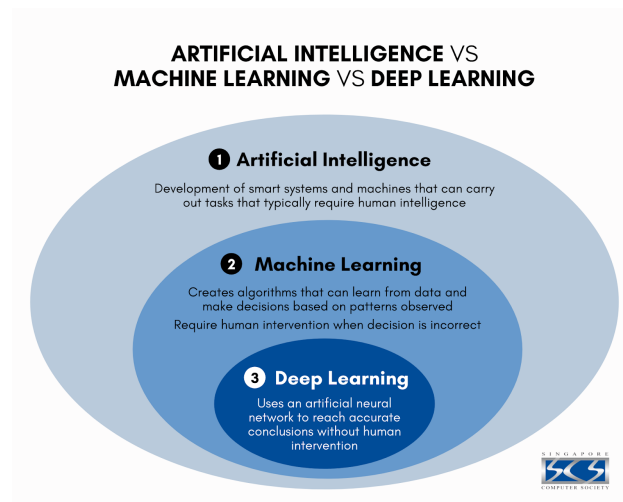


Figura 14: El aprendizaje profundo como capa dentro de la IA.

Esta se caracteriza por el uso de redes neuronales artificiales profundas, que consisten en modelos compuestos por múltiples capas de neuronas interconectadas. Estas capas trabajan de manera conjunta para procesar los datos de entrada, y aprender características y representaciones útiles de los mismos. La detección y extracción de dichas características se realiza de forma automática en el proceso de entrenamiento de la red, que ajusta parámetros de las capas para optimizar el desempeño de la tarea. Cada capa agrega una nueva dimensión de abstracción y refinamiento, por lo que las representaciones extraídas se tornan cada vez más complejas a medida que se avanza a través del modelo. El proceso de interconexión y transformación de datos en estas múltiples capas forma la base de la capacidad del aprendizaje profundo para aprender y comprender patrones de alta complejidad en una variedad de aplicaciones [21][22][23].

2.1. Redes Neuronales

Las neuronas artificiales son las unidades fundamentales de procesamiento en las redes neuronales. Estas neuronas toman un tensor de entrada y aplican operaciones matemáticas a este tensor para producir un tensor de salida. Las operaciones que realizan pueden estar influenciadas por parámetros conocidos como pesos, los cuales se ajustan durante el entrenamiento de la red. Aunque en algunos casos, las operaciones pueden ser fijas y no requerir parámetros específicos, o incluso no tener parámetros [24].

Como se puede ver en el ejemplo de la Figura 15, las redes neuronales se componen de capas de neuronas organizadas en una estructura que consta de tres partes fundamentales: la capa de entrada, las capas intermedias u ocultas y la capa de salida. Los datos se ingresan en la capa de entrada, donde se inicia el procesamiento. Luego, a medida que los datos fluyen a través de las capas intermedias, se aprenden representaciones cada vez más complejas. Puede haber un gran número de capas ocultas en una red neuronal. Finalmente, se obtiene un resultado desde la capa de salida, que proporciona la respuesta o predicción deseada.

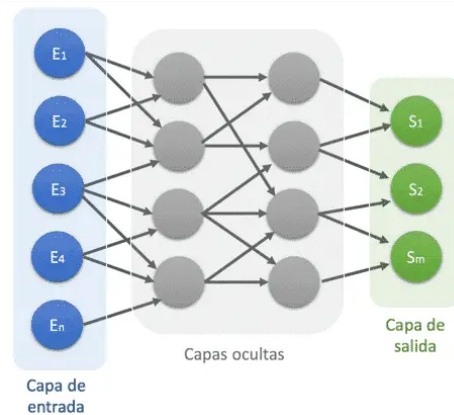


Figura 15: Ejemplo de una red neuronal. Cada círculo representa a una neurona y cada columna a una capa; se distinguen la capa de entrada (en azul), las intermedias (en gris) y la de salida (en verde).

2.2. Clasificación de Capas de Redes Neuronales

Las capas que componen las redes neuronales se pueden clasificar de varias maneras según su función y características. Sin embargo, la clasificación más fundamental y ampliamente utilizada se basa en su función dentro del proceso de procesamiento de datos. Cada tipo de capa está diseñado para abordar desafíos específicos y ofrecer a la red la capacidad de aprender y representar patrones en los datos de manera efectiva. A continuación, se destacan las capas más comunes.

2.2.1. Capas Densas

Las capas densas, también conocidas como capas totalmente conectadas (*fully connected*), son un elemento central en las redes neuronales. En estas capas, cada neurona está conectada con todas las de la capa siguiente, lo que permite capturar relaciones complejas en los datos. Si bien son ampliamente utilizadas en tareas de clasificación, especialmente en problemas binarios, presentan un alto costo computacional y no preservan la estructura espacial de los datos, lo que es crítico en el procesamiento de imágenes. Por esta razón, en tareas de clasificación de imágenes, se prefieren las capas convolucionales, diseñadas específicamente para conservar la correlación espacial de los píxeles en las imágenes.

2.2.2. Capas Convolucionales

Las capas convolucionales son componentes esenciales en las redes neuronales convolucionales (CNN) y desempeñan un papel crucial en el procesamiento de datos estructurados en cuadrícula, como imágenes. Una de las ventajas más notables de estas capas es su capacidad para preservar la estructura espacial de los datos. Mientras que las capas densas conectan cada neurona con todas las de la capa siguiente y tratan cada entrada por separado en un vector unidimensional, las capas convolucionales aplican filtros en regiones locales, que se deslizan por toda la imagen. Esto les permite ser capaces de aprender patrones y características invariantes en toda la imagen de forma eficiente. Cada vez que se aplica un filtro, se produce un mapa de características, que es una representación de las características relevantes encontradas en diferentes ubicaciones de la entrada. Estos mapas de características son fundamentales para el reconocimiento de patrones espaciales, como bordes, texturas y otros detalles. Su capacidad para detectar patrones locales y preservar la estructura espacial las hace ideales para tareas de visión por computadora, como la clasificación de imágenes, la detección de objetos y la segmentación semántica.

2.2.3. Capas de Atención

Las capas de Redes de Atención son la base de arquitecturas de atención como los modelos *Transformers*. Su función principal radica en permitir que la red dirija su “atención” hacia partes particulares y relevantes de los datos de entrada. A diferencia de la tradicional secuenciación de

datos, estas capas pueden considerar simultáneamente las relaciones entre elementos.

En el contexto del procesamiento del lenguaje natural, las capas de atención hacen que un modelo pueda enfocarse en palabras clave dentro de una oración o en secciones vitales de un documento. En el ámbito de la visión por computadora, ayudan a las redes a destacar áreas específicas de una imagen que contienen información de importancia. Esta capacidad ha revolucionado profundamente la manera en que las máquinas comprenden y procesan tanto el lenguaje natural como la información visual. Como resultado, son fundamentales en aplicaciones como la traducción automática, resumen de texto, generación de contenido y reconocimiento de objetos, entre muchas otras.

2.2.4. Capas Pooling

Las capas de *pooling* son fundamentales para el procesamiento de datos con estructura de cuadrícula, como imágenes. Estas capas realizan un submuestreo de los mapas de características generados por las capas convolucionales, lo que significa que reducen la resolución espacial disminuyendo la carga computacional en la red.

El proceso de *pooling* consiste en seleccionar un valor representativo de un grupo de valores en una región local del mapa de activación. Existen diferentes estrategias para hacer esto. Con *max pooling*, por ejemplo, se selecciona el valor máximo en cada región local. *Average pooling*, en cambio, toma el valor promedio. También existe el *global pooling*, que realiza una reducción global de la resolución al calcular una única estadística (por ejemplo, el valor máximo) para todo el mapa de activación.

La reducción de resolución que proporciona el *pooling* es beneficiosa ya que ayuda a prevenir el sobreajuste al reducir la dimensionalidad de los datos. Además, le permite al modelo enfocarse en regiones mayores y características esenciales, como bordes, texturas y patrones importantes. Es importante destacar que, aunque el proceso de *pooling* reduce la resolución espacial, mantiene la capacidad de preservar características relevantes en los datos. Esto lo convierte en una técnica ideal para aplicaciones de visión por computadora, donde se busca mantener la información importante en las imágenes mientras se reduce su tamaño. Las capas de *pooling* son especialmente útiles en la detección de objetos y la clasificación de imágenes, ya que permiten un procesamiento más eficiente sin sacrificar la capacidad de reconocer patrones clave.

2.2.5. Capas de Regularización

Las capas de Regularización, como *Dropout*, son un componente esencial en el entrenamiento de redes neuronales y tienen como objetivo prevenir el sobreajuste. El sobreajuste ocurre cuando una red neuronal se ajusta en exceso a los datos de entrenamiento y, como resultado, no generaliza bien a nuevos datos, lo que disminuye su capacidad predictiva. Para abordar este problema, se utiliza la técnica de *Dropout*.

La capa de *Dropout* funciona apagando aleatoriamente un porcentaje de neuronas durante cada iteración de entrenamiento. Este porcentaje se denomina “*dropout rate*” y se configura previamente. Por ejemplo, un dropout rate del 0.25 indica que el 25 % de las salidas de las neuronas se anulan durante el entrenamiento. Es importante destacar que esta fracción suele ser menor que 0.5 para evitar apagar demasiadas neuronas y deshabilitar efectivamente una parte significativa de la red.

La aleatoriedad en el apagado de neuronas durante el entrenamiento es esencial para prevenir la dependencia excesiva de características específicas en los datos de entrenamiento. Al forzar que diferentes conjuntos de neuronas se apaguen en cada iteración, la red se vuelve más resistente al sobreajuste y aprende a generalizar mejor. En otras palabras, el *Dropout* actúa como un mecanismo de regularización que promueve la robustez y la generalización de la red.

2.3. Importancia del Deep Learning en la actualidad

El *Deep Learning* ha revolucionado la manera en que abordamos los desafíos de la IA y se ha convertido en una herramienta ampliamente difundida en diversas disciplinas. Este crecimiento y aplicación generalizada de la técnica se pueden atribuir a tres factores clave.

En primer lugar, destaca la simplicidad de implementación, ya que el aprendizaje profundo elimina la necesidad de la ingeniería de características y simplifica los flujos de trabajo de aprendizaje automático. Anteriormente, la implementación de dichos modelos requería conocimientos profundos en matemáticas, programación y una comprensión completa de los detalles técnicos. Sin embargo, en la actualidad, con la disponibilidad de marcos de trabajo como *TensorFlow* y bibliotecas como *Keras*, las barreras de entrada se han reducido significativamente, permitiendo a una audiencia más diversa y amplia involucrarse en el desarrollo y aplicación del aprendizaje profundo. Habilidades básicas de scripting en *Python* son suficientes para llevar a cabo investigaciones avanzadas en esta disciplina, lo que hace que esta técnica sea más accesible y eficiente.

En segundo lugar, el aprendizaje profundo se destaca por su capacidad de adaptarse de manera inherente a conjuntos masivos de datos y su habilidad de aprovechar eficazmente la computación en paralelo. Esto se traduce en una escalabilidad y eficiencia significativas en comparación con enfoques anteriores. El acceso a grandes volúmenes de datos y el aprovechamiento de unidades de procesamiento gráfico (GPUs, del inglés *Graphics Processing Units*) para cálculos paralelos son dos factores clave que permiten que los modelos de aprendizaje profundo se beneficien de volúmenes de datos más grandes, lo que a su vez les permite superar muchas de las limitaciones de los enfoques clásicos de aprendizaje automático.

Por último, estos modelos tienen la capacidad de adaptarse y ser re-entrenados con nuevos datos sin necesidad de partir desde cero, lo que los hace ideales para el aprendizaje en línea continuo. Esta versatilidad y reutilización de modelos previamente entrenados en diferentes tareas es una ventaja significativa, ya que evita la necesidad de construir modelos desde cero para cada nueva aplicación. Esto ahorra tiempo y recursos, lo que contribuye a la eficiencia en una variedad de

contextos, desde la investigación hasta la implementación en la industria.

3. Visión Computacional (Computer Vision)

La Visión por Computadora es una disciplina fundamental de la IA que se enfoca en desarrollar algoritmos y modelos capaces de interpretar y comprender información visual a partir de imágenes y videos. Para lograr este objetivo, utiliza técnicas avanzadas, dentro de ellas el *Deep Learning*.

3.1. Vision Transformers (ViT)

Los *Vision Transformers (ViT)* representan una innovadora arquitectura de modelos de Aprendizaje Profundo que ha revolucionado la Visión por Computadora en los últimos años. Estos modelos son variantes de las redes neuronales que han demostrado una capacidad notable para abordar tareas de procesamiento de imágenes de manera altamente eficaz.

A diferencia de las arquitecturas de redes neuronales tradicionales, los ViT adoptan una estructura basada en mecanismos de atención, similar a la que se encuentra en modelos como el *Transformer*. Esta característica permite que los modelos de ViT capturen relaciones globales y locales en las imágenes, lo que resulta fundamental para tareas de visión más avanzadas que involucren información contextual y relaciones globales en las imágenes. La estructura básica de un ViT se puede visualizar en la Figura 16 [25].

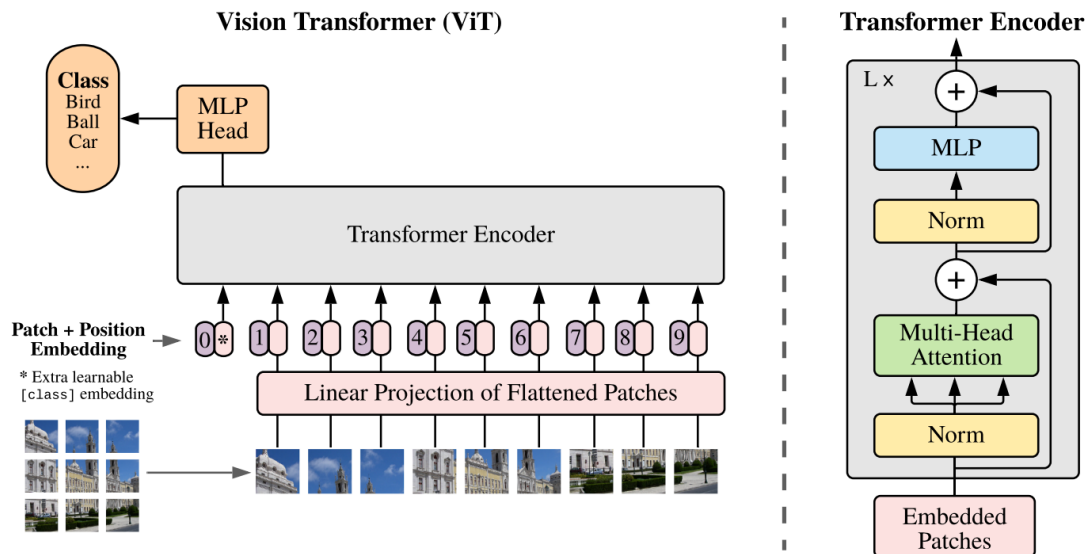


Figura 16: Ilustración de un ViT básico.

3.2. Mecanismo de Atención

El mecanismo de atención en los ViT es una técnica esencial que les permite procesar imágenes dividiéndolas en parches y asignando importancia a cada parche en función de la tarea que se desea realizar. Este proceso implica calcular puntuaciones de atención que determinan la relevancia de un parche con respecto a otros. Los parches más relevantes reciben una mayor ponderación y se utilizan para crear una nueva representación de la imagen. Este enfoque permite a los ViT comprender la estructura y las relaciones espaciales en las imágenes, lo que es crucial para tareas de Visión por Computadora, como detección de objetos y segmentación, donde se requiere una comprensión detallada de la información visual y las relaciones entre elementos en la imagen. Esto podrá ser visualizado en la Figura 17. [26]

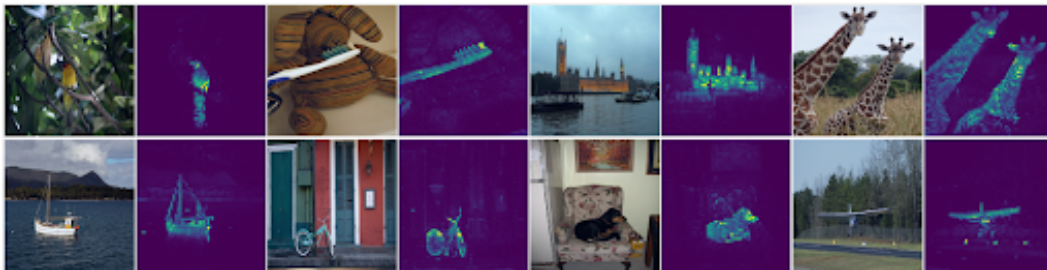


Figura 17: Visualizaciones de atención generadas por un Vision Transformer con parches de 8×8 entrenados sin supervisión [26].

4. Proceso de Aprendizaje en Modelo de Deep Learning

El proceso de aprendizaje en modelos de *deep learning*, ya sean redes neuronales tradicionales o ViT, sigue un enfoque común. El mismo se define como la modificación del comportamiento del modelo a medida que se le presentan datos de entrada o estímulos. En esencia, el conocimiento se refleja en los pesos de las conexiones entre las neuronas o unidades de atención en el caso de los ViT. Estos cambios específicos en los pesos no son aleatorios, sino que se llevan a cabo a través de un proceso de optimización durante el entrenamiento del modelo. Durante este proceso, el modelo se ajusta con el objetivo de minimizar la diferencia entre sus predicciones y los valores reales, lo que es esencial para mejorar su rendimiento y permitir que “aprenda” a realizar tareas específicas.

El proceso de aprendizaje o entrenamiento del modelo se lleva a cabo mediante una secuencia de pasos fundamentales que se describen a continuación. En la Figura 18, se puede observar un esquema que representa este proceso en el caso de una red neuronal.

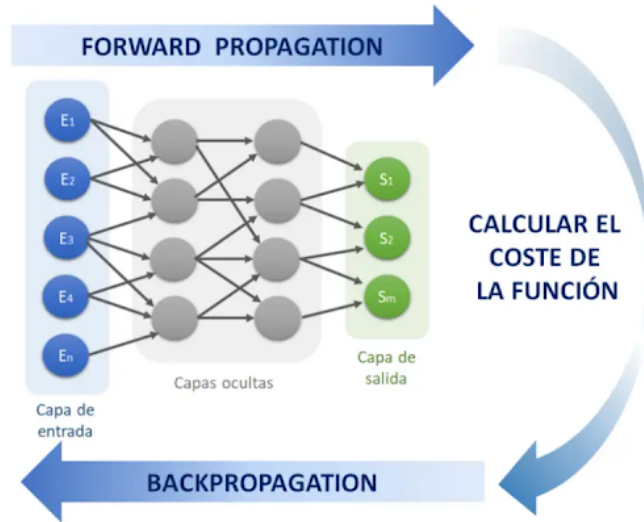


Figura 18: Esquema de entrenamiento de una red neuronal.

4.1. Inicialización de Pesos

En primer lugar, los pesos y sesgos de la red o de las unidades de atención se inicializan con valores pequeños y aleatorios. Inicializar los pesos de manera aleatoria es esencial, ya que permite que el modelo comience sin ningún conocimiento previo. A su vez, aunque no garantiza evitar quedar atrapado en óptimos locales durante el entrenamiento, ésta práctica se recomienda debido a que maximiza las posibilidades de alcanzar el mínimo global, o en su defecto, un mínimo local lo suficientemente óptimo.

4.2. Propagación hacia Adelante

A continuación, se da la fase de propagación hacia adelante (*Forward Propagation*). Durante la misma, los datos de entrenamiento se introducen en la capa de entrada del modelo. Cada uno, ya sea una red neuronal convencional o un ViT, procesa los datos y genera predicciones en función de los valores de sus pesos actuales.

4.3. Función de Costo

Tras generar predicciones, se calcula una función de pérdida o de costo (*Loss Function*) que mide el error, es decir, la diferencia entre la salida predicha por la modelo y la salida deseada, que es el valor real. El objetivo es minimizar esta función de pérdida.

Existen distintas funciones de costo según la tarea que se desea resolver. Para tareas de

clasificación con dos clases se suele emplear la Entropía Cruzada Binaria (*Cross Binary Entropy*) como función de costo. La misma cuantifica la discrepancia entre la distribución de probabilidad de las predicciones del modelo y la distribución de probabilidad real de las etiquetas. También, se utiliza comúnmente en problemas de clasificación la función de pérdida de Bisagra o SVM de Margen Suave (*Hinge Loss*). Esta función de costo ayuda a maximizar el margen entre las clases y penaliza las predicciones incorrectas que están demasiado cerca del límite de decisión.

4.4. Retropropagación del Error

El último paso en el ciclo de entrenamiento implica la retropropagación del error, donde se ajustan los pesos de las conexiones para minimizar el error previamente calculado. Los optimizadores son los encargados de llevar a cabo esta actualización de pesos. Esto lo realizan estimando la contribución parcial de cada uno de los pesos, es decir, calculando la derivada parcial de la función de costo con respecto a cada uno de los parámetros del modelo.

Existen múltiples tipos de optimizadores utilizados en el entrenamiento de redes neuronales. La elección del mismo depende del problema específico. El optimizador más básico es el Descenso de Gradiente Estocástico (SGD, del inglés *Stochastic Gradient Descent*). El mismo actualiza los pesos en función del gradiente de la función de costo para cada muestra de entrenamiento. Si bien puede ser efectivo, también puede ser ruidoso y tomar tiempo en converger. Introduciendo el término de “momentum” al SGD surge el optimizador *Momentum*. El *momentum* es una especie de inercia que ayuda a acelerar la convergencia al reducir las oscilaciones en el espacio de parámetros. Ayuda a superar mínimos locales y a mantener una dirección coherente hacia la convergencia. Otro posible optimizador es RMSprop (*Root Mean Square Propagation*), que ajusta las tasas de aprendizaje de manera adaptativa para cada parámetro. Por último, dentro de los más utilizados se encuentra Adam (*Adaptive Moment Estimation*), que combina las ventajas de los dos anteriores. Ajusta las tasas de aprendizaje de cada parámetro individualmente y utiliza momentos de primer y segundo orden para acelerar la convergencia. Su eficiencia en una amplia variedad de aplicaciones, y su capacidad de adaptación lo hacen una elección popular.

4.5. Tipos de Aprendizaje

Los algoritmos de aprendizaje profundo, como también los de aprendizaje automático, se pueden clasificar principalmente en supervisados y no supervisados. La principal diferencia entre estas dos clases radica en la presencia de etiquetas en el subconjunto de datos de entrenamiento [27] [28] [24].

A continuación en la Figura 19, se presenta un esquema que presenta de manera visual los tipos de aprendizajes y sus técnicas asociadas, que serán explicadas posteriormente.

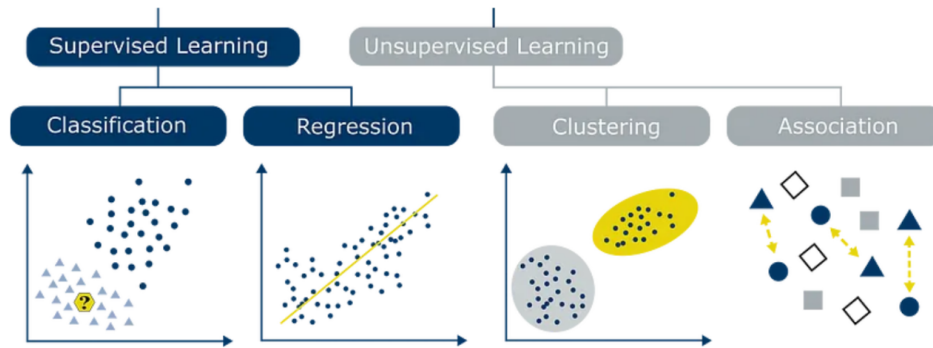


Figura 19: Esquema de los tipos de aprendizaje y sus técnicas asociadas.

4.5.1. Aprendizaje Supervisado

En el aprendizaje supervisado se entrena al modelo utilizando un conjunto de datos etiquetado. En este contexto, “etiquetado” significa que cada ejemplo de entrenamiento del conjunto de datos tiene una respuesta deseada conocida, es decir, se conoce cuál debería ser la salida o la respuesta correcta. En caso de que ésta no coincida con la deseada, se procederá a modificar los pesos de las conexiones, con el fin de conseguir que la salida obtenida se aproxime a la deseada. Una representación gráfica de esto puede ser visualizado en la Figura 20. El objetivo de este proceso es que el modelo pueda aprender a hacer predicciones precisas sobre datos no vistos utilizando el patrón que encuentra en los datos etiquetados. Esta técnica se utiliza comúnmente en problemas de clasificación y regresión.

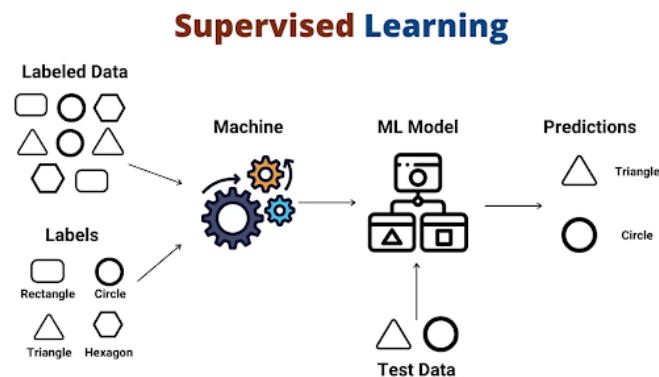


Figura 20: Esquema del Aprendizaje Supervisado.

4.5.2. Aprendizaje No Supervisado

El aprendizaje no supervisado, también conocido como autosupervisado, es un enfoque de entrenamiento de modelos de *deep learning* en el cual el modelo se entrena con datos no etiquetados.

A diferencia del aprendizaje supervisado, aquí el modelo no recibe información sobre si su salida es correcta o incorrecta en respuesta a una entrada específica. En cambio, el modelo debe descubrir por sí mismo las características, regularidades, correlaciones o categorías presentes en los datos de entrada. Una representación gráfica de lo explicado aquí puede ser visualizado en la Figura 21.

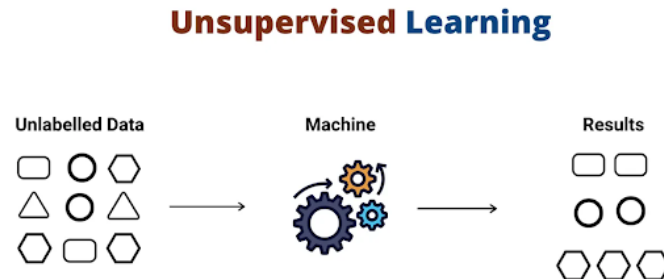


Figura 21: Esquema del Aprendizaje No-Supervisado.

4.5.3. Aprendizaje Débilmente Supervisado

El aprendizaje débilmente supervisado es una categoría menos frecuente que se ubica entre el aprendizaje supervisado y el no supervisado. En este enfoque, las etiquetas proporcionadas a los datos de entrenamiento son menos específicas que en el aprendizaje supervisado. En lugar de etiquetas precisas para cada ejemplo, se utilizan etiquetas más generales o parciales [29].

5. Evaluación de Modelos

La evaluación de modelos es fundamental para determinar el rendimiento de sus predicciones en nuevos datos. Una herramienta esencial en esta evaluación es la matriz de confusión, que nos permite ver cuántos casos de cada categoría se han clasificado correctamente y cuántos incorrectamente. Cuando se trata de clasificación binaria, hay cuatro medidas clave que nos ayudan a entender la comparación entre las predicciones de un modelo y las etiquetas reales: [30] [31]

- **Verdaderos Positivos (VP):** Estos son los casos en los que el modelo predice correctamente que algo pertenece a la categoría positiva.
- **Verdaderos Negativos (VN):** Representan los casos en los que el modelo predice correctamente que algo pertenece a la categoría negativa.
- **Falsos Positivos (FP):** Estos son los casos en los que el modelo predice incorrectamente que algo pertenece a la categoría positiva.
- **Falsos Negativos (FN):** Corresponden a los casos en los que el modelo predice incorrectamente que algo pertenece a la categoría negativa.

		Actual Values	
		Negative (n)	Positive (s)
Predicted Values	Negative (n)	TN	FN
	Positive (s)	FP	TP

Figura 22: Matriz de confusión.

Estas medidas se combinan en la matriz de confusión, como se muestra en la Figura 22. Con los valores presentes en la matriz de confusión, se puede calcular una variedad de métricas clave, que son esenciales para interpretar cómo se está desempeñando el modelo. A continuación, definiremos estas métricas y proporcionaremos sus ecuaciones:

5.1. Métricas

- **Exactitud (*Accuracy*):** Mide la proporción de casos correctamente clasificados con respecto al total de observaciones. Se calcula como:

$$Exactitud = \frac{VP + VN}{VP + VN + FP + FN} \quad (1)$$

- **Precisión (*Precisión*):** Mide la proporción de casos positivos correctamente clasificados en relación con el total de casos clasificados como positivos. Se calcula como:

$$Precisión = \frac{VP}{VP + FP} \quad (2)$$

- **Sensibilidad (*Recall*) o Tasa de VP:** Representa la proporción de casos positivos correctamente clasificados en relación con el total de casos positivos. Se calcula como:

$$Sensibilidad = \frac{VP}{VP + FN} \quad (3)$$

Es la capacidad del modelo para identificar todos los casos positivos.

- **Especificidad o Tasa de VN:** Mide la capacidad del modelo para identificar correctamente los casos negativos. Se calcula como:

$$Especificidad = \frac{VN}{VN + FP} \quad (4)$$

- **Valor-F1:** Combina las medidas de precisión y sensibilidad en una sola métrica, lo que proporciona una evaluación global del desempeño del modelo. Se calcula como:

$$\text{Valor-F1} = \frac{2 * \text{Precisión} * \text{Sensibilidad}}{\text{Precisión} + \text{Sensibilidad}} \quad (5)$$

- **AUC-ROC:** La curva ROC (*Receiver Operating Characteristic*) es una representación gráfica del rendimiento de un modelo de clasificación en diferentes umbrales de decisión. El AUC-ROC mide el área bajo esta curva y proporciona una métrica que evalúa la capacidad del modelo para distinguir entre clases. Un valor de AUC-ROC cercano a 1 indica un mejor rendimiento del modelo. A continuación, en la Figura 23, se puede observar un gráfico de la curva ROC, a modo de ejemplo. La diagonal muestra el desempeño de un clasificador aleatorio, y se muestran tres clasificadores de ejemplo (azul, naranja, verde)



Figura 23: Curva ROC con tasa de falsos positivos y tasa de verdaderos positivos.

Todas estas métricas pueden tener valores entre 0 y 1, y cuanto más cercanas a 1 sean, mejor será el rendimiento del modelo. Cada una de estas métricas tiene su importancia y puede ser relevante según los objetivos del modelo y la aplicación específica. La elección de la métrica adecuada depende de si se da más peso a la minimización de FP, FN o una combinación de ambos en la evaluación del rendimiento del modelo.

6. Aprendizaje por Transferencia (Transfer Learning)

El aprendizaje por transferencia, a menudo referido como *Transfer Learning*, es una estrategia fundamental en el campo de la IA. Consiste en mejorar el rendimiento de un modelo en una nueva tarea al aplicar conocimiento previamente adquirido en tareas relacionadas que ya han sido aprendidas. Similar al proceso de aprendizaje humano, el *Transfer Learning* representa un avance

en la eficiencia de la IA, permitiendo que los modelos aprendan de experiencias pasadas para abordar nuevas tareas [32] [33]. La necesidad de aplicar *Transfer Learning* se manifiesta en áreas menos desarrolladas tecnológicamente, donde la obtención de datos de entrenamiento puede ser difícil y costosa. Esta limitación en la disponibilidad de datos de prueba puede deberse a la rareza de los datos, los altos costos asociados con su recopilación y etiquetado, o la falta de acceso a los mismos. Sin embargo, con la creciente disponibilidad de grandes repositorios de datos, se ha vuelto más factible utilizar conjuntos de datos existentes que están relacionados, aunque no sean idénticos, al dominio de interés, lo que hace que el *Transfer Learning* sea una estrategia atractiva.

Existen tres enfoques principales para la aplicación del aprendizaje por transferencia. La elección de enfoque dependerá de la naturaleza del problema y la disponibilidad de datos [34]. El primer enfoque implica resolver la tarea deseada utilizando como punto de partida un modelo que fue entrenado con abundantes datos para realizar una tarea relacionada. La extensión en la que se reutiliza y reentrena el modelo dependerá del problema en cuestión. Si ambas tareas comparten la misma entrada, se puede hacer predicciones para los nuevos datos sin realizar adaptaciones. En caso contrario, se puede modificar y reentrenar partes del modelo específicas para adaptarlas a la nueva tarea. El segundo enfoque implica aprovechar modelos pre entrenados que están disponibles. Diferentes marcos de trabajo, como *Hugging Face*, *Tensorflow* o *Keras*, ofrecen una variedad de modelos pre entrenados adecuados para tareas de transferencia, y estas elecciones se basan en el problema en cuestión. Este método es ampliamente utilizado en el aprendizaje profundo. Por último, el tercer enfoque consiste en utilizar el aprendizaje profundo para descubrir la mejor representación de tu problema, es decir, encontrar las características más importantes sin necesidad de reentrenar un modelo completo. Este enfoque, conocido como extracción de características, se centra en extraer características relevantes de los datos y puede ofrecer un rendimiento mejorado en comparación con el diseño manual de características.

Materiales y Métodos

1. Base de Datos

1.1. Descripción de Imágenes Histológicas

La base de datos utilizada se compone de imágenes histológicas de TGI (CCR y de CE) que han sido seleccionadas. Estas imágenes provienen de láminas de tejido tumoral obtenidas mediante muestras quirúrgicas de pacientes. Todo el proceso de recopilación y preparación de estas muestras incluye una etapa crítica conocida como la fase preanalítica. Durante esta fase, se llevan a cabo una serie de procedimientos esenciales que aseguran la integridad de las muestras y la calidad de las imágenes resultantes. La fase preanalítica comprende varios pasos, que incluyen la fijación de las muestras en formalina, su inclusión en parafina y el corte en secciones de hasta 5 micrómetros (μ m). Para permitir la visualización de las estructuras celulares presentes en estos cortes, se realiza una tinción con H&E.

Posteriormente, las imágenes histológicas se obtienen capturando fotografías de estas láminas con la ayuda de microscopios digitales de alta resolución o escáneres de láminas enteras. A esta imagen histológica digitalizada se la denomina WSI. Estos escáneres realizan un barrido de toda la lámina a una alta resolución. Es importante destacar que las WSI se almacenan en archivos con formato .svs. Estos archivos son conocidos por su gran tamaño y peso, alcanzando aproximadamente 1 o 2 *gigabytes* por imagen. El manejo de archivos de tal envergadura es un desafío en la patología computacional y requiere recursos significativos en términos de capacidad de almacenamiento y procesamiento de datos.

1.2. Recopilación de Datos

Las imágenes histológicas utilizadas en este proyecto se obtuvieron de dos fuentes distintas. Por un lado, se recopilaron a partir del repositorio público TCGA. El mismo es de acceso público y alberga datos de pacientes principalmente de los Estados Unidos. Todas las imágenes están disponibles públicamente en <https://portal.gdc.cancer.gov>. Las imágenes provienen de tres cohortes principales: TCGA-COAD, TCGA-READ y TCGA-STAD. Dichos cohortes están compuestos por imágenes de láminas de tejido teñidas con H&E de pacientes diagnosticados con adenocarcinoma de colon, recto y estómago. Para determinar el estado de IMS de los pacientes en cuestión, se utilizó secuenciación genómica [35]. Se tuvieron en cuenta únicamente aquellos casos cuyos tumores MSS y MSI-H, excluyendo los tumores con MSI-L. Como el desequilibrio de clases afecta significativamente el rendimiento de los clasificadores basados en aprendizaje profundo, se tomaron cantidades equitativas para cada grupo.

Por otro lado, se obtuvieron muestras histológicas de archivo de adenocarcinomas de colon y recto para realizar una validación externa del modelo. Para todos los pacientes, una lámina

histológica de H&E estaba disponible y se conocía el estatus de IMS/MMR a través de biología molecular o IHQ.

Las muestras se digitalizaron usando el Aperio LV1 IVD, un digitalizador de portaobjetos de microscopio con una resolución de $0.28 \mu\text{m}/\text{px}$. Este escáner permite la visualización en vivo y el escaneo automático de hasta 4 portaobjetos en alta resolución. Además, brinda acceso rápido a imágenes en menos de 15 segundos y control remoto para ajustar aumentos (2.5x, 5x, 20x, 40x). Una imagen del mismo puede ser visualizada en la Figura 24.

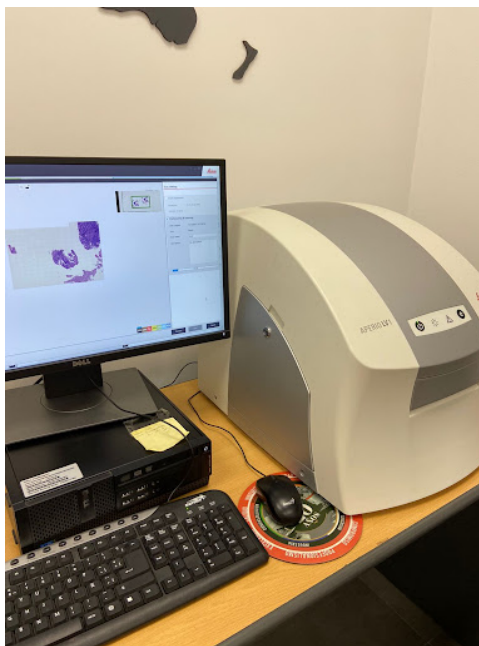


Figura 24: Digitalización de las imágenes histológicas con el digitalizador Aperio LV1 IVD.

A continuación, se presenta el Cuadro 1 que detalla la cantidad de imágenes por cohorte, dividiéndolas en las categorías MSS y MSI-H. Estas categorías se encuentran organizadas según el origen de la obtención de las muestras.

	MSI-H	MSS	Total
TCGA Colon y Recto	51	68	119
TCGA Estómago	22	25	47
Cohorte Externo de Colon y Recto	19	39	58

Cuadro 1: Cantidad de imágenes por cohorte utilizadas para las diferentes fases del modelo.

1.3. Preprocesamiento de Imágenes

Como fue mencionado anteriormente, la fase preanalítica es crucial ya que sienta las bases para el proceso analítico posterior. Cualquier error o insuficiencia en esta etapa puede comprometer la integridad de las muestras y, por lo tanto, influir en la precisión de las pruebas posteriores. La atención meticulosa a los detalles en la fase preanalítica es esencial para garantizar la obtención de resultados clínicos confiables y precisos. Por este motivo, todas las láminas del TCGA y las del cohorte externo fueron revisadas de forma individual por observadores capacitados, bajo la supervisión de patólogos expertos con el propósito de asegurar la presencia de tejido tumoral en la lámina y la calidad diagnóstica de la misma.

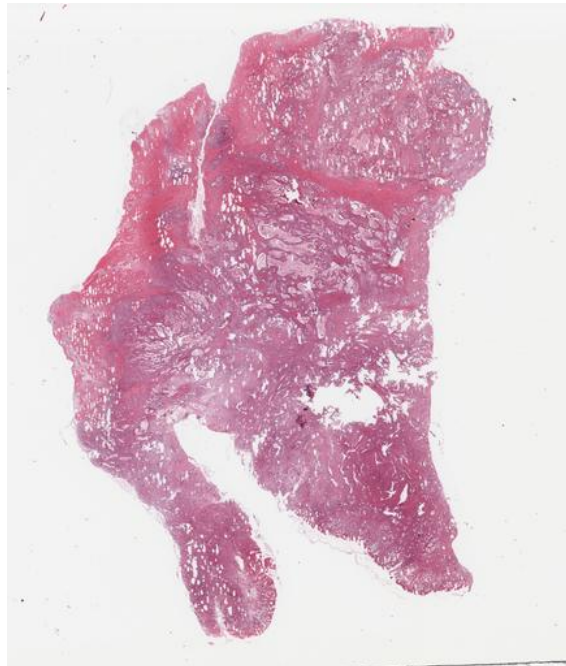
Bajo un minucioso análisis, procedimos a la exclusión de ciertas imágenes de la base de datos por diversas razones. En primer lugar, se identificaron imágenes que presentaban marcas realizadas con marcador indeleble. Esto podría afectar la integridad de los datos y por tanto, estas imágenes se excluyeron del conjunto de muestras. Pueden observarse ejemplos de estos casos eliminados en la Figura 25 a continuación.



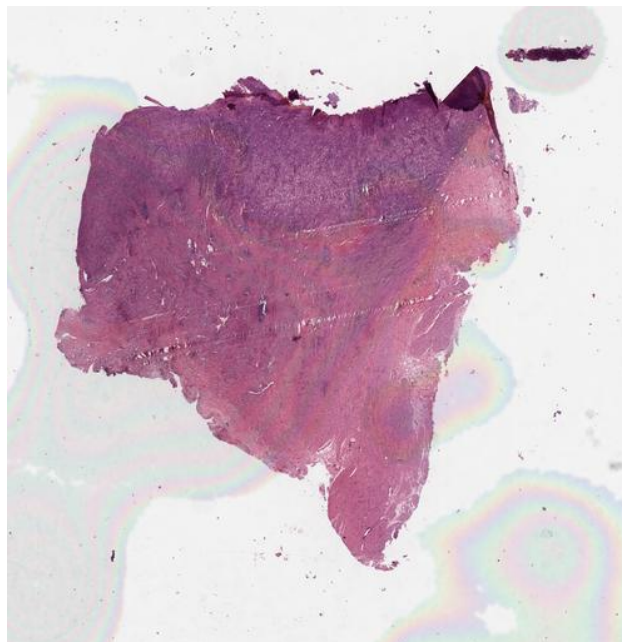
Figura 25: Imágenes histológicas extraídas del TCGA-STAD excluidas de la base de datos por sus anotaciones en marcador.

Asimismo, se excluyeron las imágenes con tinciones de H&E inadecuadas, que se traduce en un exceso o falta de hematoxilina o eosina. Estos problemas dificultan la identificación de las estructuras celulares como las características nucleares, que son el objeto de análisis del modelo. Es por este motivo que la calidad de la tinción es esencial, ya que influye en la capacidad de obtener una interpretación precisa de la morfología de las muestras.

Además, se excluyeron las láminas que presentaban suciedad, manchas, burbujas, como también imágenes con evidencia de fijación inadecuada. Esto último puede degradar las estructuras celulares, resultar en la pérdida de información crucial, generar resultados sesgados y, en algunos casos, hacer que las muestras sean inutilizables para el análisis histológico. Para ilustrar estos casos, a continuación, en la Figura 26, presentamos imágenes con este tipo de errores.



a) WSI con tejido roto.



b) WSI con inadecuada fijación.

Figura 26: WSI excluidos por errores en la fase preanalítica laboratorial.

Dentro de los problemas en los portaobjetos que encontramos para la digitalización, hallamos el exceso de bálsamo, muestras rotas que dejan vacíos, líneas irregulares que confunden el algoritmo de análisis, o pliegues, lo que ocasiona la acumulación de tejido y vuelve ilegibles ciertos detalles. Estas imágenes se pueden diferenciar de las que sí se encontraron en condiciones óptimas para su análisis, como se puede observar en la Figura 27.

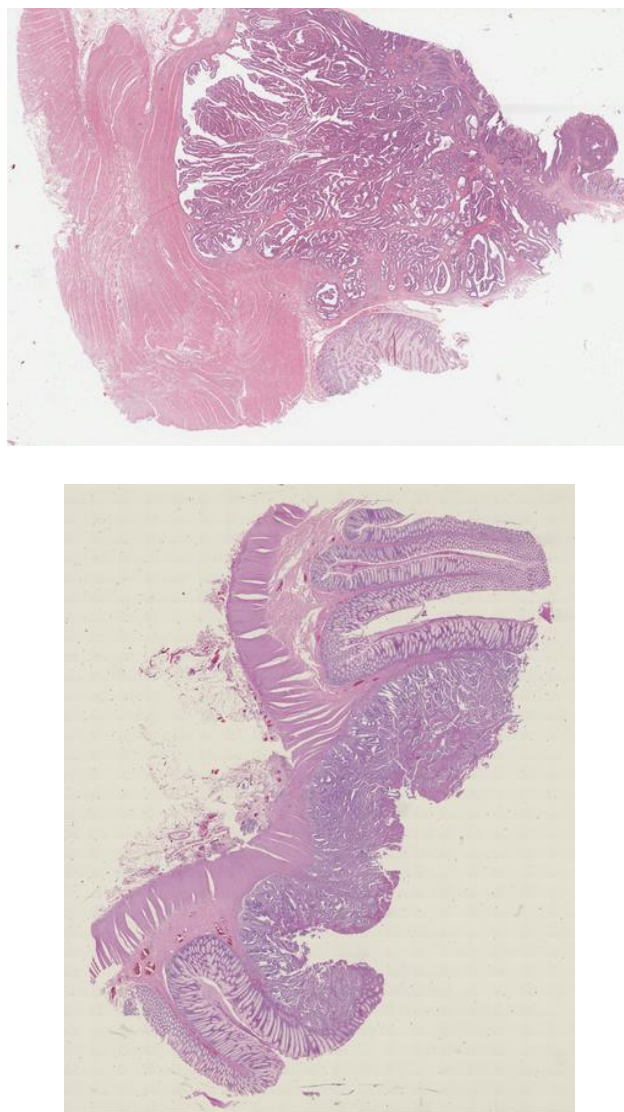


Figura 27: WSIs en condiciones óptimas para su utilización.

Es importante destacar que estos criterios de exclusión se aplicaron de manera rigurosa y meticulosa para preservar la calidad de la base de datos. Esta selección minuciosa garantiza que las imágenes utilizadas en el entrenamiento y evaluación del modelo sean representativas y confiables. Prestar atención a estos detalles durante la fase preanalítica es fundamental para asegurar

resultados clínicos precisos y confiables.

2. Líneas de desarrollo del Modelo para la Clasificación de MSI-H/MSS

En el marco del proyecto final, el objetivo fue desarrollar un modelo de IA capaz de predecir IMS en TGI mediante el análisis de imágenes histológicas obtenidas por microscopía óptica. Para alcanzar el objetivo establecido, se desarrollaron dos modelos en simultáneo. En primera instancia, se desarrolló un modelo basado en ViT, y en segundo lugar, se implementó un modelo que se basó en la metodología propuesta por la herramienta CLAM. A continuación, se describe la secuencia de pasos que se llevaron a cabo para el desarrollo de cada uno de los modelos.

2.1. Modelo ViT

El desarrollo del modelo basado en ViT siguió un proceso estructurado. Inicialmente, las WSI fueron preprocesadas; en este paso se realizó la normalización y la subdivisión en parches de 224x224 píxeles. Estos parches se sometieron a una extracción de características y posteriormente, se convirtieron en tensores de *PyTorch*, ajustando sus valores de píxeles al rango necesario.

El punto de partida fue un modelo de ViT pre-entrenado con las imágenes de ImageNet-21k (14 millones de imágenes, 21,843 clases). Este modelo fue nuevamente entrenado con las imágenes específicas del conjunto de datos del proyecto. La manera de generar los conjuntos de datos de entrenamiento, validación y evaluación fue de forma manual, generando únicamente un solo conjunto de datos para cada una de las fases. El proceso de entrenamiento consistió en 10 epochs, utilizando un optimizador AdamW, una tasa de aprendizaje (LR, del inglés *Learning Rate*) de 1e-4, y la función de pérdida de entropía cruzada.

El modelo se empleó para clasificar cada parche como MSI-H o MSS. Para lograr una clasificación definitiva a nivel del paciente, se sumaron las probabilidades de pertenencia a cada clase de todos los parches dentro de una imagen del tejido completo. Al consolidar estos valores acumulados para ambas categorías, se calculó un promedio para cada paciente, lo que permitió determinar su clasificación final.

Los resultados obtenidos con el modelo ViT no fueron satisfactorios, ya que a pesar de su potencial, no se le dedicaron todos los esfuerzos necesarios para optimizarlo debido a la implementación en paralelo con el modelo CLAM. El enfoque dual reveló que el modelo CLAM arrojaba resultados más prometedores en comparación con el ViT. Esta observación fue crucial para enfocar los esfuerzos hacia el desarrollo y perfeccionamiento del CLAM, dejando el modelo ViT en una fase preliminar con mejoras potenciales pendientes, una de ellas siendo la ejecución de validación cruzada con K-pliegues en lugar de generar conjuntos de datos manualmente. Estos resultados fueron determinantes en la priorización del CLAM como la dirección principal para el análisis, al tiempo que se consideró el modelo ViT como un área de interés para investigaciones y refinamiento a futuro.

2.2. Clustering-constrained Attention Multiple Instance Learning (CLAM)

CLAM es un modelo de aprendizaje profundo de alto rendimiento utilizado en la patología computacional. Es un método flexible, de propósito general y adaptable que puede utilizarse en una variedad de tareas diferentes en patología computacional, tanto en entornos clínicos como de investigación. Se utiliza principalmente en la clasificación de WSI y es capaz de manejar problemas de subtipificación multi-clase [36].

CLAM se caracteriza por utilizar como técnica el aprendizaje basado en atención, lo que le permite al modelo enfocarse en las partes más relevantes de las WSI. Este método se clasifica como supervisión débil, lo que significa que no requiere etiquetas detalladas a nivel de instancia, como extracciones de ROI o anotaciones de parches individuales. El modelo se entrena únicamente con etiquetas globales de WSI, es decir, una sola etiqueta por muestra. Para ello, este enfoque emplea una capa de atención que identifica las regiones de alto valor diagnóstico en las WSI.

La capacidad de CLAM de realizar clasificaciones en las imágenes de láminas completas sin requerir anotaciones de parches o ROIs representa una valiosa contribución a la patología computacional. Esto se justifica por la eliminación del procedimiento meticuloso que conlleva la marcación individual de las imágenes, un proceso que es laborioso y exige la experiencia de patólogos altamente capacitados. Simultáneamente, reduce considerablemente la etapa de preprocesamiento de las imágenes, lo que es especialmente beneficioso dado que las imágenes se encuentran almacenadas en archivos .svs que, como se mencionó anteriormente, son pesados en términos computacionales.

CLAM puede cumplir con dos tareas principales; en primer lugar, puede diferenciar el tejido entre tumoral y no tumoral, que dentro del modelo se lo denomina *“task_1_tumor_vs_normal”*. En segundo lugar, puede hacer subclasificaciones del tejido en varias clases diferentes dependiendo de la cantidad de etiquetas con las que se entrena al modelo, cuya tarea se la denomina *“task_2_tumor_subtyping”*.

Dada la naturaleza de la clasificación requerida en este proyecto final, la tarea más adecuada para nuestros propósitos fue el *“task_2”*, en la cual se proporcionaron directamente las etiquetas de MSS o MSI-H. De esta manera, el mecanismo de atención del modelo no se centra en las regiones tumorales del tejido, sino que se enfoca en las características del WSI que considera reflejar el comportamiento indicado por la etiqueta.

El enfoque de aprendizaje débilmente supervisado suprime la necesidad de efectuar una bifurcación en el proceso de clasificación, como se evidencia en el trabajo de Yamashita et al [37]. En su enfoque previo, se desarrollaba inicialmente un modelo destinado a discernir entre tejido tumoral y sano, y posteriormente, esos resultados se empleaban en un segundo modelo para lograr un diagnóstico de MSS y MSI-H. La fusión de ambos modelos en uno solo, capaz de obviar el primer paso y concentrarse en características imperceptibles para el ojo humano pero descriptivas de la estabilidad o inestabilidad del tejido, confiere a esta herramienta su utilidad, innovación y eficiencia en términos de costos.

Gracias al aprendizaje basado en atención, CLAM también es capaz de generar mapas de

calor (en inglés conocidos como “*heatmaps*”), altamente interpretables que permiten a los patólogos visualizar, para cada diapositiva, la contribución relativa y la importancia de cada región de tejido en las predicciones del modelo. Estos heatmaps no solo ofrecen información valiosa para interpretar las decisiones del modelo, sino que también comparten de dónde está enfocando mayormente su atención, evitando que el algoritmo se convierta en una caja negra. CLAM está disponible públicamente como un paquete Python de fácil uso en GitHub.

3. Desarrollo de Modelo para la Clasificación MSI-H/MSS

3.1. Preprocesamiento de las WSI

Este proceso es esencial para la preparación de las WSI, que actúan como el input inicial del modelo. Desde la segmentación automatizada, la creación de parches y la extracción de características, se transforman las imágenes WSI en vectores de características, que, a su vez, sirven como el input del modelo clasificador para la diferenciación en MSI-H y MSS.

3.1.1. Segmentación de la WSI

El primer paso necesario para poder aplicar el modelo comienza con la segmentación automatizada de las regiones de tejido en una diapositiva digitalizada. Dicho proceso comienza con la carga de la WSI en la memoria a una resolución reducida, menor que la resolución original. Esta reducción se realiza para facilitar el procesamiento computacional subsiguiente. Una vez cargada, la imagen se convierte del espacio de color RGB al espacio de color HSV (del inglés *Hue*, *Saturation*, *Value*). La conversión a HSV puede facilitar la identificación de las regiones de tejido en la imagen.

La segmentación se lleva a cabo aplicando un umbral en el canal de saturación de la imagen HSV, lo que crea una máscara binaria que resalta las regiones de tejido o primer plano, marcándolas en blanco y el resto de la imagen en negro. Para mejorar la precisión de esta máscara, se aplica un suavizado mediano, lo que permite obtener contornos más definidos de las regiones de tejido.

Después de esta etapa, se realiza un cierre morfológico adicional en la máscara. Esto implica llenar pequeños espacios y agujeros en la máscara, asegurando que las regiones de tejido estén continuas y bien definidas. La máscara y los contornos aproximados de las regiones de tejido se almacenan para su procesamiento posterior.

El proceso de segmentación es altamente adaptable a las necesidades específicas de cada imagen, permitiendo un control preciso sobre los resultados. Los parámetros de segmentación se ajustan para garantizar una segmentación precisa que se adapte a las particularidades de cada imagen. Esto es esencial en nuestro dataset, donde las imágenes del TCGA y las imágenes digitalizadas del cohorte externo requieren parámetros diferentes debido a sus diferentes orígenes y procesos de preparación. Los parámetros que se debieron ajustar en particular fueron el umbral de

segmentación, el tamaño del filtro mediano, los umbrales de área para tejido y cavidades, el uso de relleno y la elección del método de Otsu. Estas adaptaciones aseguran una segmentación óptima y adecuada para cada conjunto de imágenes. En la Figura 28 que se presenta a continuación, se aprecia claramente la segmentación y máscara aplicada en imágenes procedentes tanto del TCGA como del cohorte externo. Las áreas de tejido han sido meticulosamente remarcadas con la línea verde, mientras que las áreas de cavidades o agujeros, quedan remarcadas en azul.

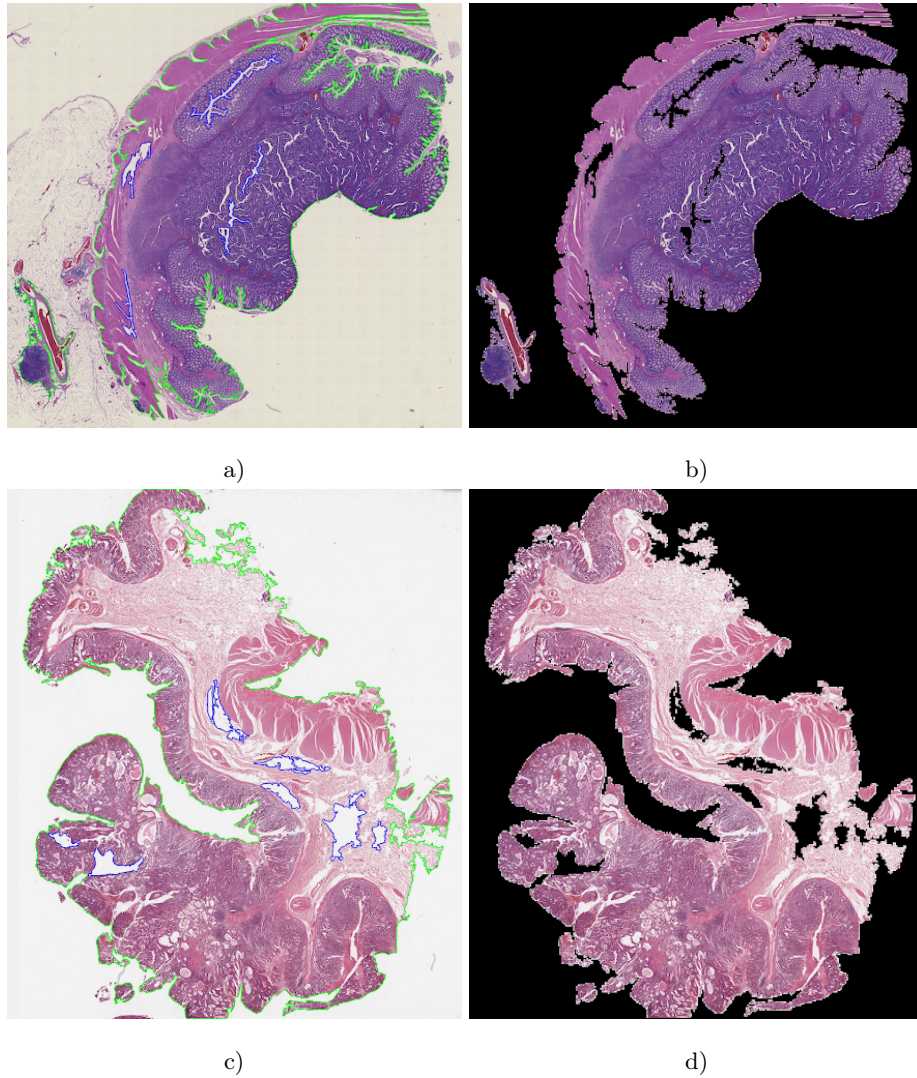


Figura 28: Segmentación y máscaras de WSIs: a) y b) Cohorte Externo; c) y d) Cohorte interno, TCGA.

3.1.2. Creación de Parches

Después de la segmentación, para cada WSI, se recortan parches de 256×256 píxeles desde los contornos de primer plano previamente segmentados, a una magnificación de 40x. Estos parches

se almacenan junto con sus coordenadas y metadatos de la WSI utilizando el formato de datos jerárquicos hdf5. La cantidad de parches extraídos de cada imagen puede variar significativamente dependiendo de su tamaño y la magnificación especificada. La cantidad de parches extraídos de cada WSI puede llegar hasta cientos de miles, especialmente en el caso de las imágenes de resecciones más grandes. En las Figuras 29 y 30 a continuación se puede observar un esquema de este proceso y parches obtenidos del mismo.

La creación de parches es una práctica esencial para simplificar la gestión de imágenes histológicas de alta resolución, como las WSIs. Estas imágenes son voluminosas y altamente detalladas, lo que dificulta su procesamiento en su totalidad debido a su complejidad y tamaño. La división en parches facilita la manipulación de la información y, al mismo tiempo, prepara los fragmentos para ser utilizados como entradas directas a las CNN, que siguen a continuación en el modelo. Además, el procesamiento de parches mejora la eficiencia general del proceso, ya que permite al modelo analizar y clasificar estos fragmentos de manera más rápida y eficiente en comparación con el procesamiento de imágenes completas.

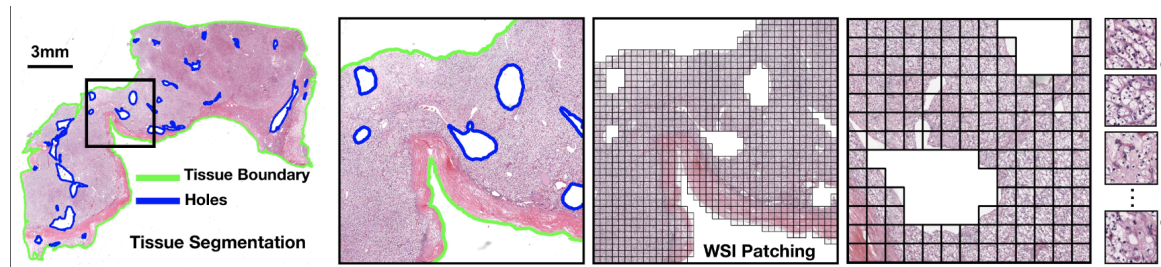


Figura 29: Representación esquemática del proceso de creación de parches.

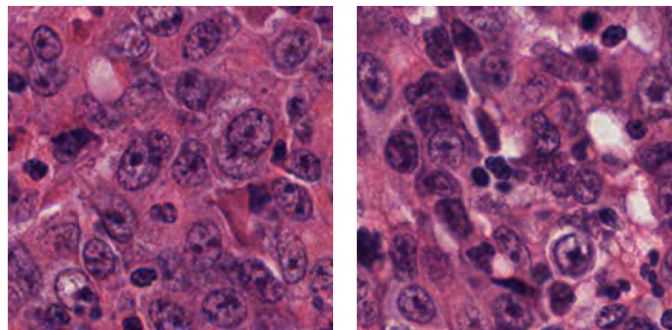


Figura 30: Ejemplos de parches (tiles) generados a partir de una WSI.

3.1.3. Extracción de Características

Después de la creación de parches en cada diapositiva, se utiliza una red neuronal convolucional profunda, conocida como ResNet50, para calcular una representación de características de baja dimensión para cada parche de imagen. ResNet50 es un modelo de red neuronal previamente entrenado en el conjunto de datos ImageNet, que es una amplia colección de imágenes utilizada para entrenar modelos de visión por computadora en tareas de clasificación de objetos.

Este modelo ResNet50 se utiliza para procesar los parches de 256×256 píxeles, y después del tercer bloque residual de la red, se aplica un muestreo adaptativo medio-espacial. Esto significa que, en lugar de realizar un muestreo regular o uniforme, el proceso se adapta al contenido de la imagen y considera la información espacial (ubicación) de los píxeles en el parche. Esto permite transformar cada parche de 256×256 píxeles en un vector de características de 1024 dimensiones. Esta operación reduce la dimensionalidad y extrae características relevantes de manera eficaz para su posterior procesamiento en el modelo.

El uso de características extraídas como entradas al modelo reduce drásticamente el tiempo de entrenamiento y los recursos computacionales necesarios. Además, el uso de características de baja dimensión facilita la gestión y el procesamiento de los datos, ya que permite ajustar todos los parches en una diapositiva (hasta 150,000 o más) en la memoria de la GPU al mismo tiempo. Así, es posible entrenar el modelo con miles de WSI en cuestión de horas una vez que se han extraído las características.

Los pasos mencionados se aplicaron tanto a las imágenes utilizadas en el proceso de entrenamiento del modelo como a las imágenes utilizadas para evaluar su rendimiento. Gracias a la extracción de características, tanto el entrenamiento como la evaluación se realizan en un espacio de características de baja dimensionalidad en lugar del espacio de píxeles de alta dimensionalidad. Esto implica una reducción significativa en el volumen de datos, que se reduce en aproximadamente 200 veces. Como resultado, se logra una drástica reducción en los cálculos necesarios para entrenar el modelo, lo que contribuye a una mayor eficiencia y velocidad en dicho proceso.

3.2. Estructura del Modelo para la Clasificación de MSI-H/MSS

3.2.1. Compresión de Representaciones de Nivel de Parche

En primer lugar, el modelo cuenta con una capa completamente conectada, denominada W1, que toma las representaciones de parches de 1024 dimensiones y las comprime a vectores de 512 dimensiones. Esto se logra mediante una transformación lineal de los datos.

3.2.2. Función de Agregación Basada en Atención

En la siguiente etapa del modelo, se emplea una función de agregación para combinar información de diversas fuentes o elementos en una representación única. Lo que lo hace especialmente poderoso es que en lugar de depender de funciones de agregación estáticas, utiliza una función de agregación basada en la atención, la cual es entrenable y altamente adaptable. En este contexto, la “atención” se refiere a la capacidad de la red para destacar características importantes y relegar las menos relevantes para la tarea de clasificación de IMS que está llevando a cabo.

La red de atención cuenta con dos capas iniciales de atención que funcionan como una “columna vertebral de atención” compartida por las dos clases. Posteriormente, la red de atención se divide en dos ramas de atención paralelas, ya que se están considerando dos clases en el problema: MSI-H y MSS. Cada rama está diseñada para una clase específica, lo que permite a la red calcular puntajes de atención particulares para cada clase. Este enfoque es esencial para que la red pueda evaluar qué características se consideran evidencia positiva o negativa para cada clase en particular.

A continuación, en cada rama siguen dos clasificadores independientes y paralelos. Cada clasificador asigna un puntaje de nivel de WSI sin normalizar para su clase correspondiente. Estos puntajes son fundamentales para la inferencia y la predicción final de la clase a la que pertenece la imagen histológica en su totalidad.

3.2.3. Instance-Level Clustering

Para fomentar aún más el aprendizaje de características específicas de cada clase, se introduce un objetivo adicional de agrupamiento binario a nivel de instancia durante el entrenamiento, también conocido como “*Instance-Level Clustering*”. Para cada una de las clases, se coloca una capa de agrupamiento con 512 unidades ocultas después de la primera capa completamente conectada. Dado que no hay etiquetas disponibles para cada parche, se utilizan las salidas de la red de atención para crear etiquetas simuladas durante el entrenamiento y supervisar el agrupamiento. En lugar de agrupar todos los parches en una diapositiva, el objetivo se optimiza solo para el subconjunto de parches en los que el modelo presta mucha o poca atención.

Para evitar confusiones, cuando se tiene una diapositiva con una etiqueta de clase verdadera, se considera la rama de atención correspondiente a esta clase como “en la clase”, y las otras ramas como “fuera de la clase”. Las puntuaciones de atención en la clase se ordenan, y se asignan etiquetas positivas a los parches con las puntuaciones más altas y etiquetas negativas a los parches con puntuaciones más bajas.

Durante el entrenamiento, se espera que los parches con puntuaciones altas sean evidencia positiva para la clase, y los parches con puntuaciones bajas sean evidencia negativa. La tarea de agrupamiento se interpreta como restringir el espacio de características a nivel de parche de modo que la evidencia positiva sea separable de la evidencia negativa de cada clase. Para la IMS, el modelo asume que los parches representativos de las diferentes clases son mutuamente excluyentes.

Por lo tanto, se impone supervisión adicional para garantizar que los parches fuera de la clase no sean evidencia positiva.

3.2.4. Inferencia y Predicción

Durante la etapa de inferencia, se calcula la distribución de probabilidad predicha para cada clase aplicando una función *softmax* a los puntajes de predicción de nivel de WSI. La función *softmax* es una operación matemática que toma un conjunto de puntajes (en este caso, los puntajes de predicción de nivel de diapositiva) y los convierte en probabilidades. Esta conversión se realiza de tal manera que los puntajes más altos se traducen en probabilidades más altas, lo que significa que la clase con el puntaje más alto tendrá la probabilidad más alta de ser la clase predicha.

Por imagen, en el proceso de clasificación, el algoritmo genera probabilidades de pertenencia para ambas clases: MSS y MSI-H. Se determina a qué clase pertenece la imagen en función de cuál de estas probabilidades es más alta. Este criterio se traduce naturalmente en un punto de corte del 50 %. Este umbral se ha seleccionado para garantizar imparcialidad y objetividad en el proceso de predicción, evitando el sesgo hacia una clase específica. Esta decisión estratégica busca mantener la integridad y el equilibrio del modelo, asegurando una evaluación neutral, libre de influencias externas. La adopción de este umbral estándar no solo preserva la capacidad de generalización del modelo en diferentes contextos clínicos, sino que también garantiza su aplicabilidad en una variedad de situaciones y conjuntos de datos, manteniendo la solidez y adaptabilidad del algoritmo.

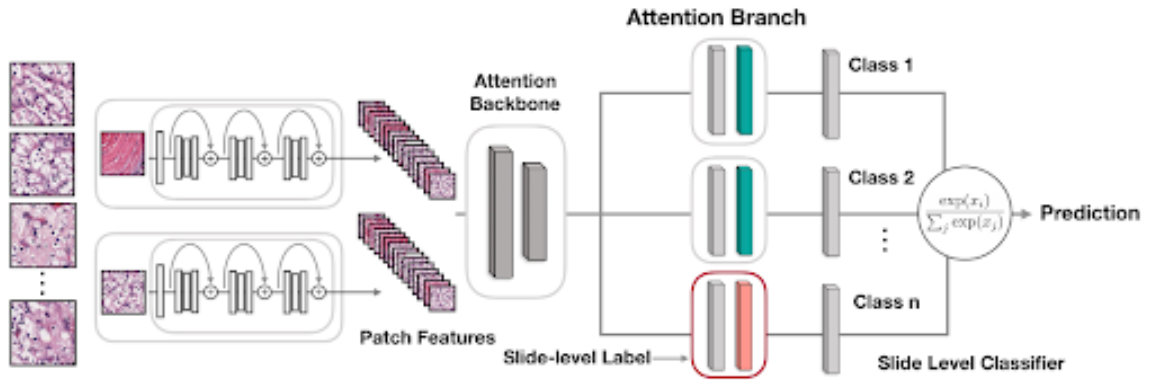


Figura 31: Diagrama de la secuencia de pasos del modelo.

3.3. Entrenamiento y Evaluación del Modelo

El entrenamiento y la evaluación del modelo se realizaron mediante un enfoque de validación cruzada de Monte Carlo con 10 pliegues. Esto quiere decir que se llevaron a cabo 10 iteraciones del proceso de entrenamiento y evaluación, lo que implica que se realizaron 10 modelos diferentes.

Utilizando el mecanismo de validación cruzada de 10-pliegues, se logra reducir la subjetividad y el sesgo del trabajo, dado que los pliegues son generados de manera aleatoria.

Para cada modelo, el conjunto de WSI descargadas del TCGA de Colon, Recto y Estómago se dividió en tres conjuntos: uno de entrenamiento (que representaba el 80 % de los casos), uno de validación (con el 10 % de los casos) y uno de evaluación (también con el 10 % de los casos). Cada partición se creó de manera independiente en cada iteración, lo que garantizó que los datos fueran distintos para cada pliegue. Es fundamental destacar que se mantuvo constante la proporción de casos en cada categoría en todos los conjuntos, garantizando una representación equitativa de las clases en cada etapa.

Durante el entrenamiento en cada uno de los 10 pliegues, los modelos se entrenaron durante al menos 50 épocas. Además, se implementó un mecanismo de parada temprana que se activó si la pérdida en el conjunto de validación no disminuyó durante más de 20 épocas consecutivas. Este enfoque aseguró un entrenamiento efectivo y evitó el sobreajuste del modelo.

En el transcurso del entrenamiento en cada pliegue, el rendimiento del modelo se evaluó en el conjunto de validación. Esto permitió la selección del mejor modelo en cada pliegue. El modelo que presentó la pérdida de validación más baja se guardó en cada pliegue. Finalmente, una vez completado el entrenamiento en los 10 pliegues, se evaluó el modelo seleccionado en el conjunto de pruebas. Este proceso de selección y prueba del modelo garantiza el uso del modelo óptimo y evalúa su rendimiento final en datos que no se utilizaron previamente.

En el proceso de entrenamiento de nuestro modelo, se utilizaron 119 WSI del conjunto de datos del TCGA de Colon y Recto, compuesto por 51 muestras MSI-H y 68 MSS. Además, se proporcionó al modelo un archivo .csv con los nombres de cada imagen y sus etiquetas correspondientes: "MSS" o "MSI-H". Con esta información, el algoritmo dividió el conjunto de muestras en tres partes, siguiendo una proporción de 80/10/10. En consecuencia, de las imágenes del TCGA de Colon y Recto, se asignaron 96 para el entrenamiento, 12 para la validación y 11 para la evaluación.

Con las 11 muestras asignadas al conjunto de evaluación, se probaron los modelos generados en cada uno de los 10 pliegues. La elección del modelo final se fundamentó en las métricas resultantes de estas evaluaciones, en particular de la exactitud (*accuracy*) y el AUC-ROC. En la Figura 32 a continuación se presentan las métricas resultantes para cada modelo, siendo el modelo del pliegue 2 el seleccionado debido a sus mejores métricas.

Pliegue	AUC	Exáctitud
0	0.60	0.58
1	0.77	0.83
2	0.86	0.92
3	0.60	0.67
4	0.49	0.50
5	0.77	0.58
6	0.77	0.67
7	0.94	0.58
8	0.63	0.42
9	0.63	0.67

Figura 32: Métricas para la selección del modelo final.

La disponibilidad limitada de WSI de colon y recto en el conjunto de datos del TCGA, catalogados como MSS o MSI-H en investigaciones externas, ha planteado un desafío en la evaluación de nuestro modelo. La mayoría de estas imágenes se utilizaron para entrenar el modelo, lo que dejó un número reducido para su evaluación, lo cual resulta en métricas poco representativas.

Para abordar la limitación en la disponibilidad de WSI de colon y recto del TCGA, se optó por combinar este conjunto de datos con muestras adicionales del TCGA de estómago. Esta estrategia tiene como objetivo aumentar la representatividad y diversidad del conjunto de evaluación, ofreciendo una visión más robusta del rendimiento del modelo. Este enfoque, al contar con más datos disponibles, minimiza el riesgo de sesgos o variaciones que podrían surgir con un conjunto de evaluación pequeño y proporciona métricas más sólidas y confiables.

La inclusión de muestras de estómago en el conjunto de datos de evaluación no solo amplía la cantidad de datos disponibles, sino que también pone a prueba la capacidad del modelo para desempeñarse de manera efectiva en diversas áreas del TGI. El objetivo principal es evaluar si las métricas de rendimiento del modelo, basadas en estos conjuntos de datos ampliados, son lo suficientemente sólidas como para respaldar la extrapolación de la eficacia del modelo a otros segmentos del tracto. De esta forma, se podría lograr con el modelo una aplicación más amplia y versátil, lo que nos permitiría abordar una variedad de desafíos en diferentes áreas del sistema digestivo.

Dicho todo esto, se realizó una evaluación interna utilizando un conjunto de imágenes de colon, recto y estómago descargas del TCGA. El conjunto resultante constó de un total de 58 WSI, siendo 32 MSS y 26 MSI-H. Esta evaluación interna permitió obtener métricas de rendimiento inicial del modelo en un entorno controlado y con datos específicos de interés.

Por otro lado, también se llevó a cabo una evaluación externa del modelo utilizando imágenes proporcionadas por un laboratorio de anatomía patológica de la Argentina. Este paso de evaluación externa es fundamental ya que esta evaluación con datos del mundo real es esencial para garantizar que el modelo funcione adecuadamente en el entorno en el que se utilizará de manera efectiva. Aquí

también se utilizaron 58 imágenes en total, de las cuales 19 eran MSI-H, y 39 MSS.

3.3.1. Parámetros del Entrenamiento del Modelo

Se especificaron los parámetros de entrenamiento del modelo como se puede visualizar en el Cuadro 2 a continuación:

lr	k	label_frac	inst_loss	model_type
2e-4	10	1	SVM	clam_sb

Cuadro 2: Parámetros de entrenamiento del modelo.

Los parámetros del modelo se actualizan mediante el optimizador Adam con una decaimiento de peso L2 de $1e-5$. En el Cuadro presentado, “k” representa el número de pliegues, “lr” denota la tasa de aprendizaje, “label_frac” corresponde a la fracción de las muestras empleadas en el proceso de entrenamiento del modelo (en este contexto, el 100 % de las muestras), “inst_loss” hace referencia a la función de costo, que, como explicamos previamente, ayuda a maximizar el margen entre las clases y penaliza las predicciones incorrectas que están demasiado cerca del límite de decisión. Por último “model_type” es un parámetro que le permite al usuario elegir entre “clam_sb”, que utiliza un modelo de una única rama, o “clam_mb”, que utiliza un modelo de múltiples ramas.

3.4. Interpretación de la Predicción del Modelo a través del Mapa de Atención

Para interpretar la relevancia de las diferentes regiones en una diapositiva para la predicción final del modelo, se calcularon y guardaron las puntuaciones de atención sin normalizar. Estas puntuaciones, antes de aplicar la función *softmax*, se obtuvieron para todos los parches extraídos, utilizando la rama de atención asociada a la clase predicha. Luego, estas puntuaciones se transformaron en percentiles, escalándolas de 0 a 1.0, donde 1.0 representa la mayor atención y 0.0 la menor. Posteriormente, se convirtieron en colores RGB mediante un mapa de colores divergentes y se visualizaron en la WSI, identificando visualmente regiones de alta atención en rojo (evidencia positiva, alta contribución al modelo) y regiones de baja atención en azul (evidencia negativa, baja contribución al modelo). Un ejemplo de un mapa de atención generado se puede visualizar en la Figura 33.

Para generar mapas de calor más detallados, se dividieron las diapositivas en parches de 256×256 sin superposición. Se calculó la puntuación de atención en bruto para cada parche y se aplicó un procedimiento similar al mencionado anteriormente para convertir la puntuación en bruto de cada parche en colores RGB. Los mapas de calor de las WSI se superpusieron en la diapositiva original con una transparencia del 0.5, permitiendo visualizar simultáneamente las estructuras morfológicas subyacentes.

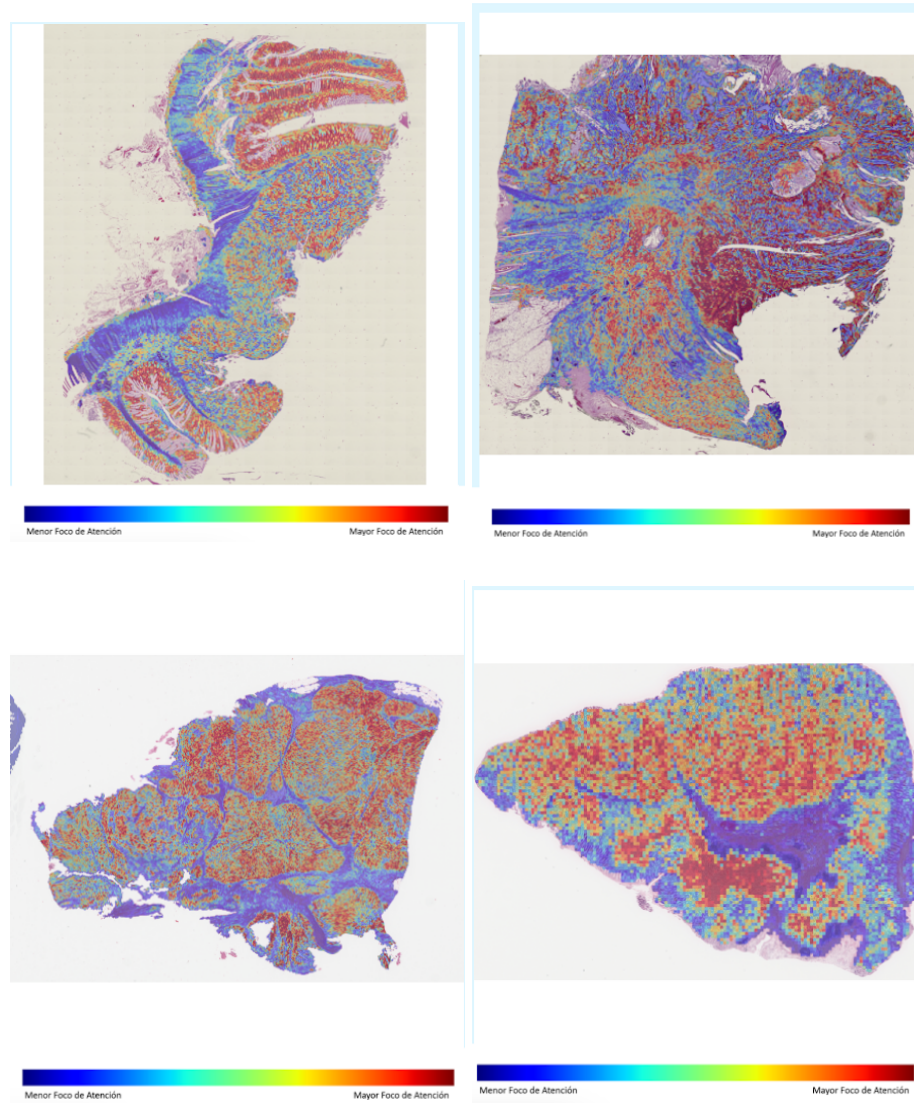


Figura 33: Mapas de atención generados a partir de imágenes histológicas de cohorte interno y externo.

4. Interfaz

Con el objetivo de proporcionar una herramienta práctica y accesible, se desarrolló una interfaz gráfica que incorpora el modelo desarrollado en este proyecto. Esta interfaz tiene como propósito principal acercar el modelo a los patólogos, convirtiéndolo en una herramienta tangible para su uso diario. Diseñada con la premisa de ser amigable para el usuario, intuitiva y eficiente, la interfaz busca simplificar el proceso de aplicación del modelo, ofreciendo una experiencia que transforme el modelo en una herramienta práctica y valiosa para su uso en un laboratorio de anatomía patológica.

En cuanto a sus funciones, la interfaz proporciona la visualización de tres imágenes clave. En primer lugar, permite observar la imagen histológica digitalizada en sí. Además, posibilita la visualización de la segmentación de tejido realizada por el modelo en una segunda imagen, donde se señala el tejido con una línea verde y se identifican agujeros y cavidades con una línea azul. Por último, se brinda la opción de visualizar el mapa de calor (*heatmap*) de la imagen histológica junto con su respectiva escala. Como fue mencionado anteriormente, en azul quedan resaltadas las áreas al que el modelo le otorga menos importancia y en rojo aquellas en las que se enfoca más intensamente.

Adicionalmente, la interfaz facilita la ejecución de la predicción principal del modelo al clasificar la imagen histológica seleccionada como MSI-H o MSS. Asimismo, permite cargar más de una imagen, guardando las imágenes cargadas en una lista de muestras previas para su fácil visualización posterior. Esta funcionalidad no solo amplía la versatilidad de la herramienta, sino que también ofrece la capacidad de comparar y analizar varias imágenes de manera eficiente. A continuación, en la Figura 34, se puede observar una captura de la interfaz en uso. Además, se puede visualizar la misma siendo utilizada en este video.

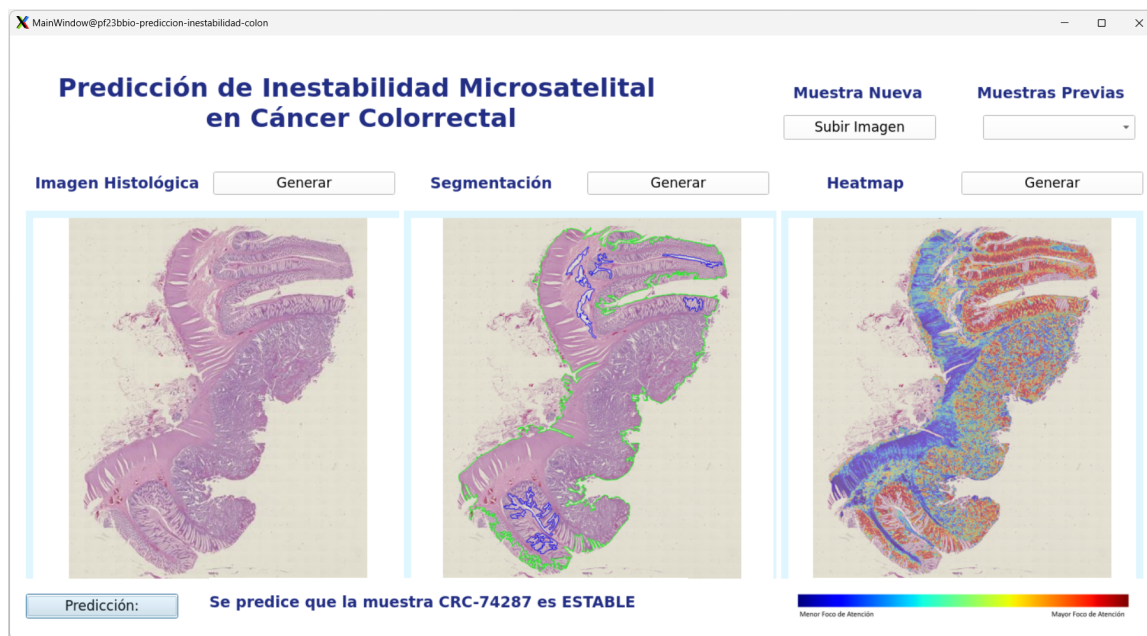


Figura 34: Interfaz gráfica que permite generar la imagen histológica, imagen segmentada, mapa de atención, y la predicción de MSI-H/MSS, como también subir nuevas imágenes.

5. Hardware y Software

5.1. Servidor

Para el desarrollo de este proyecto fue necesario utilizar un servidor de alto rendimiento alojado en el Instituto Tecnológico de Buenos Aires. La necesidad de utilizar este servidor se debió a la alta exigencia de recursos computacionales que requerían las tareas del proyecto. Al trabajar con imágenes en formato SVS, fue fundamental contar con un entorno de cómputo que pudiera manejar eficientemente la manipulación y el análisis de estos archivos de gran tamaño. El servidor seleccionado no sólo proporcionó recursos de procesamiento, sino también una capacidad de almacenamiento adecuada para la gestión de dichas imágenes. El servidor cuenta con una unidad central de procesamiento (CPU, del inglés *Computer Processing Unit*) AMD Ryzen 7 5800X con 8 núcleos, con una memoria de acceso aleatorio virtual (VRAM, del inglés *Virtual Random Access Memory*) de 32 GB de y una unidad de procesamiento gráfico (GPU, del inglés *Graphics Processing Units*) NVidia Quadro RTX 8000 con 48 GB de VRAM.

5.2. Entorno de Desarrollo

El servidor en cuestión está basado en el sistema operativo Linux, específicamente Ubuntu 22.04.2. Este entorno de desarrollo es justamente el requerido para la ejecución de la herramienta CLAM.

Todo el desarrollo y la implementación de los modelos fue llevada a cabo en Python 3.10. Se utilizaron diversas bibliotecas y paquetes, incluyendo: h5py, matplotlib, numpy, opencv, openslide, pandas, pillow, PyTorch, scikit-learn, scipy, tensorflow, tensorboardx, torchvision y smooth-topk.

La creación de la interfaz gráfica se realizó utilizando la aplicación QT Designer y la biblioteca PyQt. Para la ejecución de la misma se utilizaron las aplicaciones XQuartz y MobaXTerm.

Resultados

Se evaluó el rendimiento del modelo mediante el análisis de diversas métricas, tales como exactitud, precisión, sensibilidad, puntuación F1, especificidad y AUC-ROC. Estas métricas fueron calculadas luego de la clasificación de las imágenes en cada clase (MSS, MSI-H), utilizando el umbral de corte del 50%, como se detalla en la sección 3.2.4. *Inferencia y Predicción*.

1. Evaluación Interna con WSI de Colon, Recto y Estómago del TCGA

En primer lugar, se pueden observar los resultados correspondientes a la evaluación realizada de manera interna con las muestras obtenidas de TCGA de colon, recto y estómago. La cantidad de muestras predichas correctamente e incorrectamente por el modelo se presenta a través de una matriz de confusión, la cual se encuentra en la Figura 35. Es importante destacar que la clase definida como positiva es MSI-H, mientras que MSS es definida como negativa. A su vez, en la Figura 36 se puede observar una curva ROC de la evaluación interna para poder plasmar el AUC del mismo. Por último, en el Cuadro 3 se resumen las métricas obtenidas por el modelo.

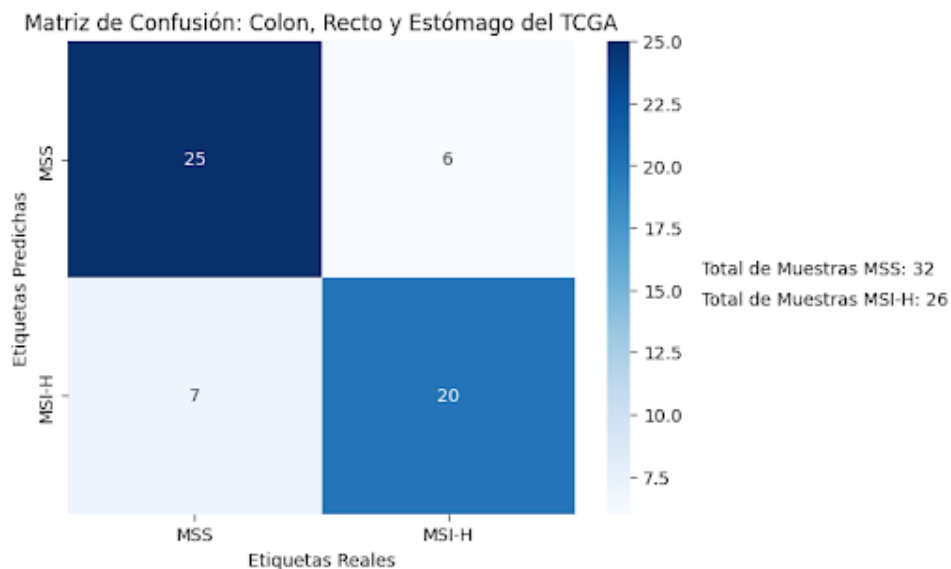


Figura 35: Matriz de confusión de la evaluación interna.

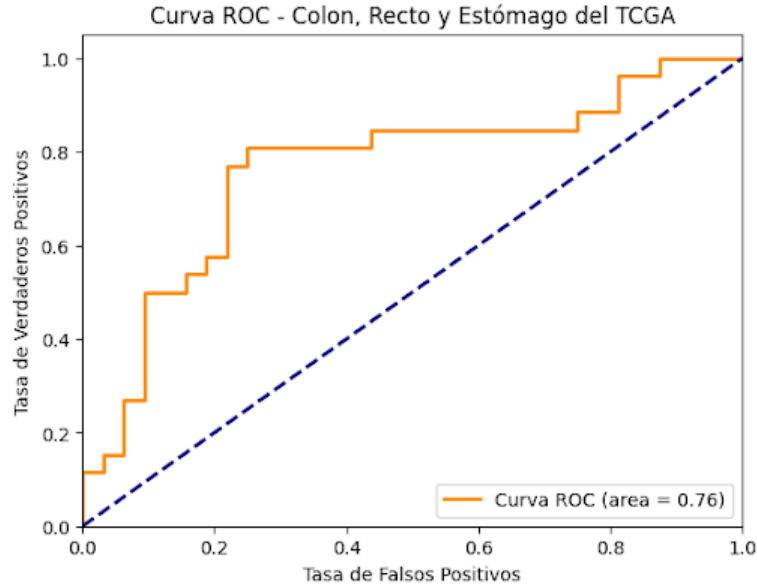


Figura 36: Curva ROC de la evaluación interna.

Métrica	Valor
Exactitud (<i>Accuracy</i>)	77,6 %
Precisión (<i>Precision</i>)	74,1 %
Sensibilidad (<i>Recall</i>)	76,9 %
Puntuación F1 (<i>F1 Score</i>)	75,5 %
Especificidad (<i>Specificity</i>)	78,1 %
AUC-ROC	76,3 %

Cuadro 3: Resultados de métricas de la evaluación interna del modelo.

2. Evaluación Externa con Colon y Recto

En segundo lugar, se presentan los resultados de la evaluación con las muestras obtenidas del cohorte externo de colon y recto, con el objetivo de mostrar cómo el modelo se desempeña en el entorno en el que se busca aplicar. En las Figuras 37 y 38, se puede observar la matriz de confusión y la curva ROC de dicha evaluación. Además, en el Cuadro 4, se presentan sus métricas.

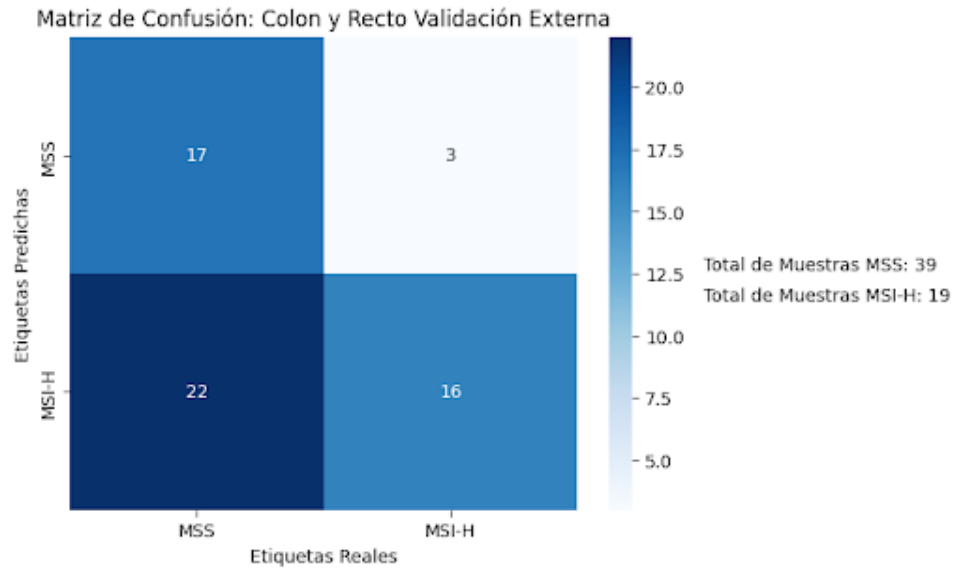


Figura 37: Matriz de confusión de la evaluación externa.

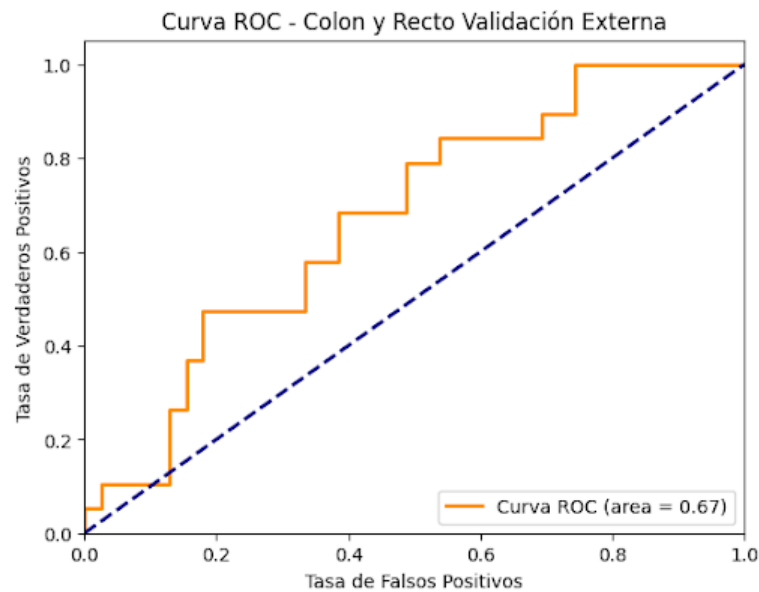


Figura 38: Curva ROC de la evaluación externa.

Métrica	Valor
Exactitud (<i>Accuracy</i>)	56,9 %
Precisión (<i>Precision</i>)	42,1 %
Sensibilidad (<i>Recall</i>)	84,2 %
Puntuación F1 (<i>F1 Score</i>)	56,1 %
Especificidad (<i>Specificity</i>)	43,6 %
AUC-ROC	67,3 %

Cuadro 4: Resultados de métricas de la evaluación externa del modelo.

3. Ejemplos de Resultados Discordantes

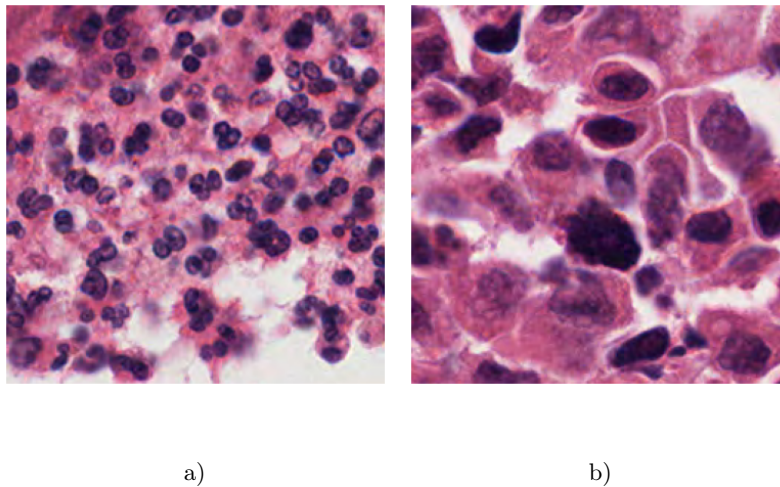


Figura 39: Parches provenientes de dos imágenes diferentes MSI-H con infiltración linfocitaria alta. El modelo predice que son MSS con una probabilidad de a) 59,7 %, y b) 82,5 %.

Discusión

Antes de analizar en detalle los resultados obtenidos en ambas evaluaciones, es crucial destacar que el enfoque principal estará puesto en el análisis de tres métricas claves: sensibilidad, especificidad y AUC-ROC.

Aquellos casos predichos como MSI-H, que son tomados como positivos por el modelo, requerirán confirmación adicional mediante técnicas adicionales de biología molecular, como PCR o NGS. Estos se diferenciarán de aquellos predichos como MSS, los cuales carecerán de esta validación adicional. Es por este motivo que el propósito central de esta investigación reside en la precisa identificación de los casos que genuinamente manifiestan MSI-H, dado que los casos predichos como MSS no serán verificados posteriormente. Si bien la sensibilidad se ha considerado como la métrica más relevante, reflejando la habilidad del clasificador para discernir la presencia de IMS, también es crucial considerar la especificidad. Esta última métrica, con un umbral mínimo aceptable del 65 %, busca equilibrar el desempeño del modelo. No solo se trata de reducir falsos negativos, sino también de evitar una sobreestimación de los casos MSI-H.

El objetivo de este estudio es que el modelo no prediga excesivamente casos MSI-H, manteniendo un equilibrio que garantice la eficiencia y la economía de recursos en el proceso de validación posterior. Es importante resaltar que el mínimo aceptable de sensibilidad se establece en un 70 %, lo que asegura una detección adecuada de los casos positivos, complementando así la necesaria especificidad mínima para una implementación efectiva del modelo en entornos clínicos.

Finalmente, es pertinente destacar la relevancia de la métrica AUC-ROC, ya que no solo describe de manera integral el desempeño del modelo, sino que también ha sido una métrica recurrente y de referencia en investigaciones previas.

1. Evaluación Interna

Los resultados de la evaluación interna confirman la solidez del rendimiento del modelo desarrollado al ser probado con las WSI del TCGA. Se obtuvo una especificidad del 78,1 % y una sensibilidad del 76,9 %, junto con un AUC-ROC del 76,3 %. Estos resultados no alcanzan los objetivos preestablecidos del 85 % de sensibilidad para ser una herramienta plenamente implementable en los laboratorios de anatomía patológica. Sin embargo, respaldan un desempeño satisfactorio del modelo, ya que la sensibilidad supera el umbral mínimo mencionado anteriormente del 70 %.

Por otro lado, en las imágenes que el modelo sí logró clasificar correctamente, en el tejido abundan las células tumorales, siendo las mismas de fácil identificación. Esta limitación se traduce en la necesidad de incluir más muestras MSI-H con gran infiltración linfocitaria en la fase de entrenamiento del modelo, para mejorar la capacidad en la detección de estos casos.

2. Evaluación Externa

Los resultados de la validación externa son esenciales para fortalecer la validez en un modelo de IA, garantizando su aplicabilidad en situaciones del mundo real y su capacidad para tomar decisiones precisas y fiables fuera del entorno de entrenamiento. El uso exclusivo de datos del TCGA, sin la inclusión de otros grupos de datos para aumentar la diversidad, se identifica como un riesgo significativo de sesgo en estudios anteriores [13].

Los resultados arrojaron una sensibilidad elevada del 84,2 %, lo que resalta la habilidad del modelo para reconocer adecuadamente las muestras positivas. Sin embargo, una especificidad del 43,6 % señala las dificultades del modelo para identificar con precisión las muestras negativas, es decir, predecir a las muestras MSS como tales.

Además, se obtuvo un AUC-ROC del 67,3 %, que respalda lo mencionado anteriormente. Este valor indica que el modelo tiene una capacidad moderada para distinguir entre las clases de manera efectiva: muestra una buena identificación de las MSI-H pero una identificación menos precisa de las MSS.

A pesar de que la sensibilidad estuvo muy cercana al 85 %, un valor que justificaría potencialmente la utilización del algoritmo en un contexto real dentro de los laboratorios de anatomía patológica, este resultado no se condice con los valores de especificidad, los cuales no alcanzaron el mínimo del 65 %. Si consideramos las demás métricas, la precisión está por debajo del 50 %, mientras que la exactitud y la puntuación F1 se encuentra en valores por debajo del 60 %. Estas cifras indican que el modelo actual no puede ser utilizado todavía en aplicaciones clínicas, ya que no cumple con los estándares necesarios de fiabilidad requeridos para su aplicación efectiva.

Es crucial destacar que los resultados obtenidos, especialmente la baja especificidad, pueden atribuirse en parte a las diferencias en las técnicas preanalíticas y analíticas entre el conjunto de datos externo y el TCGA. Esta disparidad subraya la falta de concientización sobre la importancia de la medicina de precisión en los laboratorios, para la detección adecuada y confiable de marcadores biológicos. Resulta fundamental estandarizar los procesos preanalíticos y analíticos para garantizar resultados que le den validez al estudio. Medidas como la homogeneización de la tinción, la mejora de la calidad de la fijación y la digitalización inmediata de las muestras se presentan como pasos necesarios para utilizar plenamente el potencial de esta herramienta en aplicaciones futuras. Como se puede observar en la Figura 40, la intensidad de la tinción de Hematoxilina y Eosina (H&E), puede variar tanto entre diferentes instrumentos y protocolos de tinción, como con el tiempo de fijación en formalina, lo que dificultará la lectura de pequeños cambios celulares, por parte de los digitalizadores de alta resolución [38].

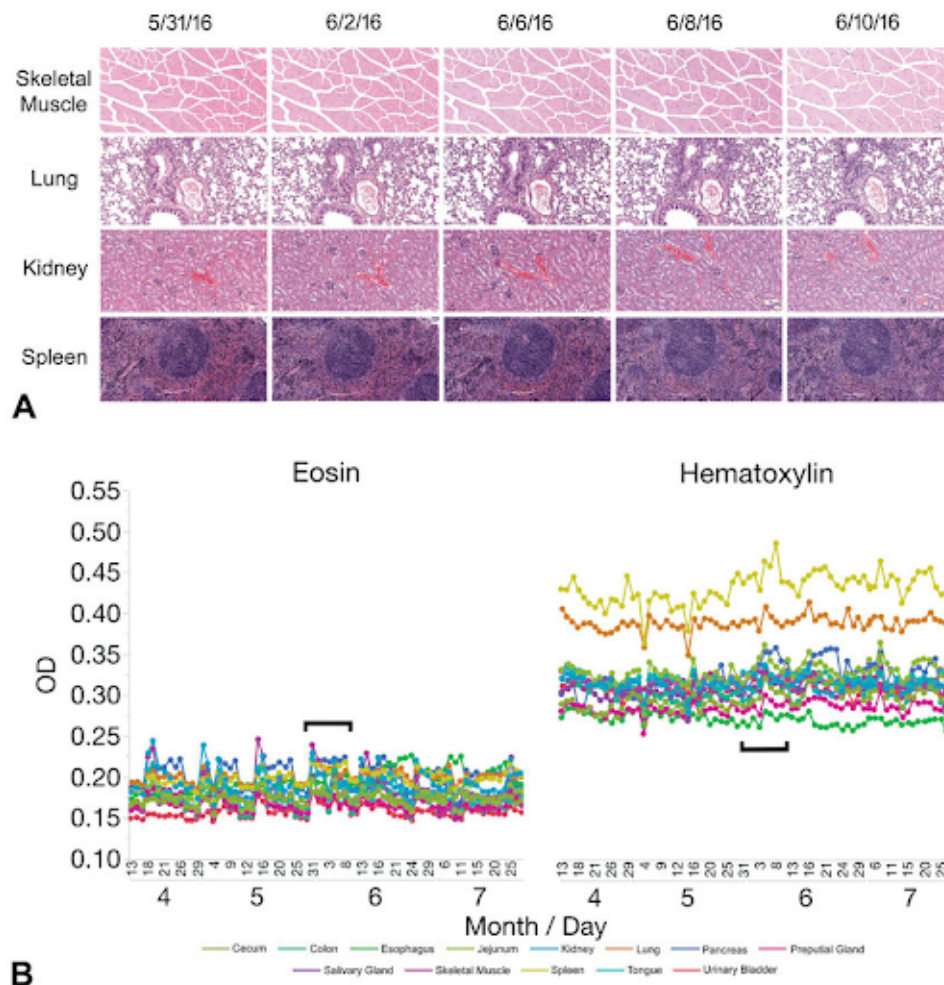


Figura 40: Características de la coloración con H&E según el tiempo de fijación del tejido en formol.

Cabe señalar que aunque se identificaron desafíos, estos resultados se alínean con el principio *"garbage in, garbage out"* que destaca la importancia de mejorar la calidad de los datos de entrada, para obtener resultados clínicamente útiles.

Es relevante destacar que este estudio al igual que la mayoría con datos bibliográficos de IA en salud, es retrospectivo. Resulta imperativo recomendar, donde exista la decisión de la digitalización como herramienta de trabajo, que las muestras sean manejadas adecuadamente, desde que ingresan a un laboratorio que cuente con esta tecnología. Esta práctica no solo optimiza la eficacia de los procesos, sino también, evita correcciones manuales posteriores, simplifica y agiliza la tarea, asegurando la fiabilidad y la validez de los datos obtenidos. Reduce la probabilidad de errores y la necesidad de intervenciones adicionales.

3. Limitaciones Generales

Tras una evaluación exhaustiva de los resultados obtenidos tanto en la evaluación interna como en la externa, y de las diferencias preanalíticas entre ellas, se identifican limitaciones generales que impactan en ambos conjuntos de datos. Es crucial reconocer estas limitaciones, ya que pueden ser la causa de no haber alcanzado, en el caso de la evaluación interna, los objetivos de máxima, y en el caso de la evaluación externa, los objetivos de mínima inicialmente propuestos. En este contexto, se lleva a cabo un análisis detallado de varios factores que pueden afectar los resultados del estudio.

En primer lugar, revisando la bibliografía de Lu et al [36], que desarrolla el algoritmo CLAM, vimos que las tareas de clasificación tumoral que arrojaron valores elevados en las métricas de AUC-ROC y de sensibilidad, probablemente se debieron a que el algoritmo buscaba identificar características visibles para el ojo humano. Esto se diferencia notablemente de esta investigación, ya que la ambición de identificar características a nivel molecular implica la necesidad de una cantidad mayor de muestras, lo que podría haber impactado en la eficacia del modelo. Lamentablemente, esta limitación no pudo ser mitigada, ya que se utilizaron todas las muestras disponibles tanto del repositorio público TCGA como del cohorte externo en esta investigación.

En segundo lugar, al analizar los mapas de atención generados para cada muestra junto con patólogos experimentados, se observan discrepancias entre la localización real del tumor y las áreas de enfoque del algoritmo. Estas diferencias sugieren que el algoritmo podría no estar centrando su atención en las áreas más relevantes para la identificación de IMS. De cara a futuras investigaciones, se proponen dos enfoques. Uno es explorar métodos alternativos, que permitan al algoritmo observar con mayor precisión en las regiones tumorales, con el objetivo de determinar si las características identificadas aportan información más relevante sobre IMS y si ello mejora las métricas de sensibilidad y especificidad. Otra línea de investigación sugerida implica que el algoritmo detalle las características en las que se enfoca actualmente, permitiendo a los patólogos estudiar el patrón identificado y evaluar su relevancia en la detección de IMS.

4. Datos Bibliográficos

El modelo desarrollado en este proyecto se alinea con la tendencia observada en 13 estudios que fueron recopilados hasta septiembre de 2021 en el artículo 'Inteligencia Artificial para Predecir Inestabilidad de Microsatélites Basada en la Histomorfología del Tumor: Una Revisión Sistemática' [13]. En este metanálisis, también se emplea IA para predecir el estado de dMMR-MSI-H a partir de características histomorfológicas. Entre los 12 estudios analizados (uno no reveló el modelo específico), el modelo basado en ResNet fue el más frecuente, lo que le da respaldo a la utilización del mismo en este estudio. Otros modelos también empleados fueron ShuffleNet, Inception-V3, MSInet e InceptionResNetV1.

El algoritmo que obtuvo el mejor desempeño para predecir IMS en el CCR, fue desarrollado por Lee et al [39]. En su estudio, Inception-V3, se entrenó utilizando una validación cruzada de

10-pliegues, con un conjunto de 584 imágenes de WSI digitalizadas del TCGA y 274 del *Hospital Saint Mary (SMH)*. Es esencial resaltar el desequilibrio significativo en las imágenes del TCGA, con 416 asociadas a pacientes MSS y 84 a MSI-H, lo que afecta el rendimiento de los clasificadores de aprendizaje profundo. Esta discrepancia es notoria en los resultados del modelo propuesto por Lee, agregando dudas sobre su fiabilidad.

Cuando el modelo entrenado se probó en un conjunto de validación interna (TCGA), el AUC fue de 89,2 %. Los autores también crearon otro modelo entrenado únicamente con el conjunto de datos de TCGA. Este último modelo puede ser comparado de manera directa con el que fue generado en este proyecto, debido a la naturaleza de los datos y la configuración de entrenamiento de nuestro estudio. Los resultados se pueden ver plasmados en el Cuadro 5 a continuación.

Tipo de Modelo	Evaluación	AUC
Modelo desarrollado por Lee et. al	Cohorte Interno	86,1 %
	Cohorte Externo	78,7 %
Modelo desarrollado en este proyecto	Cohorte Interno	76,3 %
	Cohorte Externo	67,3 %

Cuadro 5: Resumen de los resultados obtenidos por los modelos

Al comparar los resultados del Cuadro 5, se evidencia una disparidad relativamente pequeña en el rendimiento de los modelos citados de la bibliografía y los obtenidos por nuestro grupo de trabajo, considerando la significativa diferencia en el tamaño de sus conjuntos de datos de entrenamiento. Si bien, creemos que deberíamos aumentar el número de muestras, el modelo presentado en esta investigación muestra potencial, especialmente evidenciado por los resultados de la validación interna.

Es fundamental destacar su enorme potencial, sugiriendo que con ajustes y mejoras en los procesos de laboratorio, puede convertirse en una valiosa herramienta costo-efectiva en el ámbito de la medicina de precisión desde hoy hacia el futuro.

Conclusión

La aplicación de IA en la detección de MSI en tumores gastrointestinales representa un avance significativo y sumamente novedoso en el campo del diagnóstico oncológico. Su puesta a punto final, será un gran avance para simplificar y acortar el proceso de diagnóstico y colaborar en un tratamiento a medida, para todos los pacientes con esta patología.

En el transcurso de este proyecto, se ha desarrollado un modelo fundamentado en la estructura CLAM, cuyo propósito radica en la clasificación de WSI en categorías de MSI-H y MSS mediante el empleo de técnicas de aprendizaje profundo. Agregadamente, este esfuerzo investigativo se enriquece con la implementación de una interfaz gráfica, estratégicamente diseñada para optimizar su utilización por parte de los profesionales de la salud en sus tareas diarias, transformando así este avance tecnológico en una herramienta pragmática y de fácil acceso en el ámbito de la anatomía patológica y otros laboratorios en el ámbito hospitalario. La implementación de esta herramienta busca superar las limitaciones de métodos tradicionales, al ofrecer una alternativa accesible y costo-efectiva para la predicción de IMS en TGI.

La inmunoterapia para pacientes con MSI-H, generadores de gran cantidad de antígenos tumorales, ha estado en el centro del paradigma del tratamiento del cáncer en estos últimos años. Se espera que las tecnologías basadas en el aprendizaje profundo capaces de predecir el estado IMS tengan un gran impacto en el flujo de trabajo diagnóstico de la oncología y áreas relacionadas, desempeñando roles importantes como herramienta de diagnóstico. Sin embargo, es esencial destacar que existen limitaciones significativas en la aplicación práctica de estos modelos de diagnóstico en el contexto actual. Esto se debe en parte, a la falta de estandarización en los procesos preanalíticos de las muestras de laboratorios, lo que dificulta en gran medida la calidad del tejido para su digitalización. Por otra parte, el rendimiento del modelo deberá ser refinado mediante la expansión del conjunto de datos de entrenamiento. Además, una tercera limitación se podría adjudicar a la diferencia entre la localización del tumor y las áreas donde el algoritmo focaliza primordialmente su atención.

La presencia de estas limitaciones no despreciables, hace que la herramienta todavía no pueda ser implementada en su totalidad. En caso de ser utilizada, se deberá requerir de otra técnica de confirmación como la biología molecular, en aquellos casos que resulten positivos en la evaluación realizada por el algoritmo.

A pesar de estas limitaciones, los resultados se consideran muy prometedores y representan un sólido punto de partida para futuros avances. Aplicando las mejoras mencionadas previamente, el modelo podrá adaptarse a una gama más amplia de biomarcadores, mejorando su fidelidad y robustez en la detección de la IMS. La tecnología IA es un proceso evolutivo y estos resultados iniciales brindan una base sólida para el desarrollo continuo del modelo.

En resumen, este proyecto se suma al impulso colectivo de avanzar en la aplicación de IA en la predicción de la IMS, reconociendo su potencial en oncología.

Glosario de Contracciones

ADN: Ácido Desoxirribonucleico

ARN: Ácido Ribonucleico

AUC-ROC: Área bajo la Curva ROC (del inglés *Area under the ROC Curve*)

CCR: Cáncer Colorrectal

CIN: Inestabilidad Cromosómica (del inglés *Chromosome Instability*)

CLAM: *Clustering-constrained Attention Multiple Instance Learning*

CNN: Redes Neuronales Convolucionales (del inglés *Convolutional Neural Networks*)

CPU: Unidad Central de Procesamiento (del inglés *Central Processing Unit*)

dMMR: Tumores Deficientes en Reparación de Emparejamiento Incorrecto (del inglés *Deficient Mismatch Repair*)

EBV: Virus de Epstein-Barr (del inglés *Epstein-Barr Virus*)

FDA: Administración de Alimentos y Medicamentos (del inglés *Food and Drug Administration*)

FN: Falsos Negativos

FP: Falsos Positivos

GIST: Tumores Estromales Gastrointestinales (del inglés *Gastrointestinal Stromal Tumor*)

GLOBOCAN: Observatorio Global del Cáncer

GPUs: Unidades de Procesamiento Gráfico (del inglés *Graphics Processing Units*)

GS: Genómicamente estable (del inglés *Genomically Stable*)

H&E: Hematoxilina-Eosina

HSV: Matiz, Saturación, Valor (del inglés *Hue, Saturation, Value*)

IA: Inteligencia Artificial

IARC: Agencia Internacional de Investigación sobre Cáncer (del inglés *International Agency for Research in Cancer*)

ICIs: Inhibidores de Puntos de Control Inmunológico (del inglés *Immune Checkpoint Inhibitors*)

IDL: Inserción-Delección

IHQ: Inmunohistoquímica

IMS: Inestabilidad Microsatelital

lr: Tasa de Aprendizaje (del inglés *Learning Rate*)

MDSCs: Células Supresoras Derivadas de la Médula Ósea (del inglés *Myeloid-Derived Suppressor Cell*)

MHC: Complejo Principal de Histocompatibilidad (del inglés *Major Histocompatibility Complex*)

MMR: Sistema de Reparación del ADN (del inglés *Mismatch Repair*)

MMS: Estabilidad Microsatelital

MSI-H: Inestabilidad Microsatelital Alta (del inglés *Heavy Microsatellite Instability*)

MSI-L: Inestabilidad Microsatelital Baja (del inglés *Low Microsatellite Instability*)

NGS: Secuenciación de Nueva Generación (del inglés *Next Generation Sequencing*)

OMS: Organización Mundial de la Salud

PCR: Reacción en Cadena de la Polimerasa (del inglés *Polymerase Chain Reaction*)

pMMR: Tumores Proficientes en Reparación de Emparejamiento Incorrecto (del inglés *Proficient Mismatch Repair*)

ROC: Característica Operativa del Receptor (del inglés *Receiver Operating Characteristic*)

RMSprop: Propagación de la Raíz Media Cuadrática (del inglés *Root Mean Square Propagation*)

ROI: Regiones de Interés

SGD: Descenso de Gradiente Estocástico (del inglés *Stochastic Gradient Descent*)

SMH: *Hospital Saint Mary*

STR: Tándem Cortas Repetidas (del inglés *Short Tandem Repeats*)

TCGA: The Cancer Genome Atlas

TCR: Receptores de Células T (del inglés *T-cell Receptors*)

TGI: Tracto Gastrointestinal

TH1: Células T Cooperadoras Tipo 1 (del inglés *Type 1 Helper T cells*)

Tregs: Células T Reguladoras (del inglés *Regulatory T cells*)

ViT: *Vision Transformers*

VN: Verdaderos Negativos

VP: Verdaderos Positivos

VRAM: Memoria de Acceso Aleatorio Virtual (del inglés *Virtual Random Access Memory*)

WSI: Imágenes de Láminas Histológicas Completas Digitalizadas (del inglés *Whole Slide Images*)

Referencias

- [1] Instituto Nacional del Cáncer. “Diccionario del Cáncer: Colon”. (s.f.), dirección: <https://www.cancer.gov/espanol/publicaciones/diccionarios/diccionario-cancer/def/colon>.
- [2] American Cancer Society. “¿Qué es el cáncer colorrectal?” (s.f.), dirección: <https://www.cancer.org/es/cancer/tipos/cancer-de-colon-o-recto/acerca/que-es-cancer-de-colon-o-recto.html>.
- [3] Organización Mundial de la Salud. “Cáncer colorrectal”. (2023, 11 de julio), dirección: <https://www.who.int/es/news-room/fact-sheets/detail/colorectal-cancer#:~:text=Se%20estima%20que%2C%20en%202020,enfermedad%20en%20todo%20el%20mundo>.
- [4] U. Gualdrini, L. E. Iummato y M. L. Bidart, *Guía de procedimientos para la consejería de evaluación de antecedentes y riesgo*, 1a ed. Instituto Nacional del Cáncer, 2015. dirección: https://bancos.salud.gob.ar/sites/default/files/2018-10/0000000900cnt-2016-10-28-guia_ccr_consejeria.pdf.
- [5] American Cancer Society. “Recomendaciones de la Sociedad Americana Contra el Cáncer para la detección, diagnóstico y clasificación por etapas del cáncer de colon o recto”. (s.f.), dirección: <https://www.cancer.org/es/cancer/tipos/cancer-de-colon-o-recto/deteccion-diagnostico-clasificacion-por-etapas/recomendaciones-de-la-sociedad-americana-contra-el-cancer.html>.
- [6] A. C. Society, *Factores de riesgo del cáncer de estómago*, Cancer.Net, diciembre 8 de 2023. dirección: <https://www.cancer.net/es/tipos-de-cancer/cancer-de-estomago/factores-de-riesgo>.
- [7] Agencia Internacional de Investigación sobre el Cáncer. “Análisis en línea del cáncer gástrico - GLOBOCAN 2020”. (2020), dirección: https://gco.iarc.fr/today/online-analysis-table?v=2020&mode=cancer&mode_population=continents&population=900&populations=900&key=asr&sex=0&cancer=7&type=0&statistic=5&prevalence=0&population_group=0&ages_group%5B%5D=0&ages_group%5B%5D=17&group_cancer=1&include_nmsc=0&include_nmsc_other=1.
- [8] Instituto Nacional del Cáncer, Ministerio de Salud de Argentina. “Estadísticas de incidencia”. (s.f.), dirección: <https://www.argentina.gob.ar/salud/instituto-nacional-del-cancer/estadisticas/incidencia>.
- [9] Memorial Sloan Kettering Cancer Center. “Tipos de cáncer de colon”. (s.f.), dirección: <https://www.mskcc.org/es/cancer-care/types/colon/types>.
- [10] P. Lauren, “The two histological main types of gastric carcinoma: Diffuse and so-called intestinal-type carcinoma. An attempt at a histo-clinical classification”, *Acta pathologica et microbiologica Scandinavica*, vol. 64, págs. 31-49, 1965. DOI: 10.1111/apm.1965.64.1.31.
- [11] J. Guinney, R. Dienstmann, X. Wang et al., “The consensus molecular subtypes of colorectal cancer”, *Nature medicine*, vol. 21, n.º 11, págs. 1350-1356, 2015. DOI: 10.1038/nm.3967.

- [12] J. N. Nojadeh, S. Behrouz Sharif y E. Sakhinia, “Microsatellite instability in colorectal cancer”, *EXCLI journal*, vol. 17, págs. 159-168, 2018. DOI: 10.17179/excli2017-948. dirección: <https://doi.org/10.17179/excli2017-948>.
- [13] J. H. Park, E. Y. Kim, C. Luchini et al., “Artificial Intelligence for Predicting Microsatellite Instability Based on Tumor Histomorphology: A Systematic Review”, *International journal of molecular sciences*, vol. 23, n.º 5, pág. 2462, 2022. DOI: 10.3390/ijms23052462. dirección: <https://doi.org/10.3390/ijms23052462>.
- [14] R. Motta, S. Cabezas-Camarero, C. Torres-Mattos et al., “Immunotherapy in microsatellite instability metastatic colorectal cancer: Current status and future perspectives”, *Journal of clinical and translational research*, vol. 7, n.º 4, págs. 511-522, 2021.
- [15] S. Magaki, S. A. Hojat, B. Wei, A. So y W. H. Yong, “An Introduction to the Performance of Immunohistochemistry”, *Methods in molecular biology (Clifton, N.J.)*, vol. 1897, págs. 289-298, 2019. DOI: 10.1007/978-1-4939-8935-5_25. dirección: https://doi.org/10.1007/978-1-4939-8935-5_25.
- [16] A. L. Cisyk, Z. Nugent, R. H. Wightman, H. Singh y K. J. McManus, “Characterizing Microsatellite Instability and Chromosome Instability in Interval Colorectal Cancers”, *Neoplasia*, vol. 20, n.º 9, págs. 943-950, 2018. DOI: 10.1016/j.neo.2018.07.007. dirección: <https://doi.org/10.1016/j.neo.2018.07.007>.
- [17] National Cancer Institute. “Tratamiento del cáncer colorrectal (PDQ®)–Versión para profesionales de salud”. (s.f.), dirección: <https://www.cancer.gov/espanol/tipos/colorrectal/paciente/tratamiento-colorrectal-pdq>.
- [18] K. Ganesh, Z. K. Stadler, A. Cercek et al., “Immunotherapy in colorectal cancer: rationale, challenges and potential”, *Nature reviews. Gastroenterology hepatology*, vol. 16, n.º 6, págs. 361-375, 2019. DOI: 10.1038/s41575-019-0126-x. dirección: <https://doi.org/10.1038/s41575-019-0126-x>.
- [19] S. Kang, Y. Na, S. Y. Joung, S. I. Lee, S. C. Oh y B. W. Min, “The significance of microsatellite instability in colorectal cancer after controlling for clinicopathological factors”, *Medicine (Baltimore)*, vol. 97, n.º 9, e0019, 2018. DOI: 10.1097/MD.00000000000010019. dirección: <https://doi.org/10.1097/MD.00000000000010019>.
- [20] S. Popat, R. Hubner y R. S. Houlston, “Systematic review of microsatellite instability and colorectal cancer prognosis”, *Journal of clinical oncology : official journal of the American Society of Clinical Oncology*, vol. 23, n.º 3, págs. 609-618, 2005. DOI: 10.1200/JCO.2005.01.086. dirección: <https://doi.org/10.1200/JCO.2005.01.086>.
- [21] N. Rusk, “Deep learning”, *Nat Methods*, vol. 13, pág. 35, 2016. DOI: 10.1038/nmeth.3707. dirección: <https://doi.org/10.1038/nmeth.3707>.
- [22] F. Chollet, *Deep learning with Python*. Simon y Schuster, 2021.
- [23] Y. LeCun, Y. Bengio y G. Hinton, “Deep learning”, *Nature*, vol. 521, págs. 436-444, 2015. DOI: 10.1038/nature14539. dirección: <https://doi.org/10.1038/nature14539>.

- [24] D. J. Matich, “Redes Neuronales: Conceptos básicos y aplicaciones”, *Universidad Tecnológica Nacional, México*, vol. 41, págs. 12-16, 2001. dirección: https://www.frro.utn.edu.ar/repositorio/catedras/quimica/5_anio/orientadora1/monograias/matich-redesneuronales.pdf.
- [25] A. Dosovitskiy, L. Beyer, A. Kolesnikov et al., “An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale”, 2021. arXiv: 2010.11929 [cs.CV].
- [26] M. Caron, H. Touvron, I. Misra et al., “Emerging Properties in Self-Supervised Vision Transformers”, 2021. arXiv: 2104.14294 [cs.CV].
- [27] M. W. Berry, A. Mohamed y B. W. Yap, *Supervised and Unsupervised Learning for Data Science*. 2020, cap. Unsupervised and Semi-Supervised Learning. DOI: 10.1007/978-3-030-22475-2.
- [28] M. Kozan. “Supervised and Unsupervised Learning: An Intuitive Approach”. (sep. de 2021), dirección: <https://medium.com/@metehankozan/supervised-and-unsupervised-learning-an-intuitive-approach-cd8f8f64b644>.
- [29] Z. H. Zhou, “A Brief Introduction to Weakly Supervised Learning”, *National Science Review*, vol. 5, 2017. DOI: 10.1093/nsr/nwx106.
- [30] G. S. Handelman, H. K. Kok, R. V. Chandra et al., “Peering Into the Black Box of Artificial Intelligence: Evaluation Metrics of Machine Learning Methods”, *AJR. American journal of roentgenology*, vol. 212, n.º 1, págs. 38-43, 2019. DOI: 10.2214/AJR.18.20224. dirección: <https://doi.org/10.2214/AJR.18.20224>.
- [31] M. Hossin y S. M.N., “A Review on Evaluation Metrics for Data Classification Evaluations”, *International Journal of Data Mining Knowledge Management Process*, vol. 5, págs. 01-11, 2015. DOI: 10.5121/ijdkp.2015.5201.
- [32] L. Torrey y J. Shavlik, “Transfer Learning”, págs. 242-264, 2010. DOI: 10.4018/978-1-60566-766-9.ch011.
- [33] K. Weiss, T. Khoshgoftaar y D. Wang, “A survey of transfer learning”, *J Big Data*, vol. 3, pág. 9, 2016. DOI: 10.1186/s40537-016-0043-6. dirección: <https://doi.org/10.1186/s40537-016-0043-6>.
- [34] Bultin. “A Beginner’s Guide to Transfer Learning in Machine Learning”. (sep. de 2022), dirección: <https://bultin.com/data-science/transfer-learning>.
- [35] J. N. Kather, A. T. Pearson, N. Halama et al., “Deep learning can predict microsatellite instability directly from histology in gastrointestinal cancer”, *Nat Med*, vol. 25, n.º 7, págs. 1054-1056, 2019. DOI: 10.1038/s41591-019-0462-y.
- [36] M. Y. Lu, D. F. K. Williamson, T. Y. Chen et al., “Data-efficient and weakly supervised computational pathology on whole-slide images”, *Nat Biomed Eng*, vol. 5, págs. 555-570, 2021. DOI: 10.1038/s41551-020-00682-w.
- [37] R. Yamashita, J. Long, T. Longacre et al., “Deep learning model for the prediction of microsatellite instability in colorectal cancer: a diagnostic study”, *The Lancet. Oncology*, vol. 22, n.º 1, págs. 132-141, 2021. DOI: 10.1016/S1470-2045(20)30535-0.

- [38] E. Chlipala, M. Butters, M. Brous et al., “Impact of Preanalytical Factors During Histology Processing on Section Suitability for Digital Image Analysis”, *Toxicologic Pathology*, 2020. DOI: 10.1177/0192623320970534.
- [39] S. H. Lee, I. H. Song y H. J. Jang, “Feasibility of deep learning-based fully automated classification of microsatellite instability in tissue slides of colorectal cancer”, *International journal of cancer*, vol. 149, n.º 3, págs. 728-740, 2021. DOI: 10.1002/ijc.33599.