



CLASIFICACIÓN DE DESNUDEZ BASADA EN APRENDIZAJE PROFUNDO

Tesis presentada para obtener el título de Licenciatura en Ingeniería en Ciencias
de la Computación



Presenta:

Brisa Isabel Mendoza Romero

Matrícula no. 201425867

Asesor:

Dr. Luis Enrique Colmenares Guillén

20 DE SEPTIEMBRE DE 2021

BENEMÉRITA UNIVERSIDAD AUTÓNOMA DE PUEBLA
Facultad de Ciencias de la Computación

Dedicatoria

A mí mamá:

“Ser madre soltera es el doble de trabajo, el doble de estrés y el doble de lágrimas, pero también el doble de abrazos, el doble de amor y el doble de orgullo”.

Prólogo

La presente tesis lleva el título de “CLASIFICACIÓN DE DESNUDEZ BASADA EN APRENDIZAJE PROFUNDO”. presentada a continuación forma parte del proyecto de investigación 201903073C inscrito en el *Laboratorio Nacional de Supercómputo del Sureste de México* (LNS), perteneciente al padrón de laboratorios nacionales CONACYT, aceptado en otoño 2019, se llevó a cabo en el “Laboratorio de análisis forense digital para la detección de delitos en internet”, registrado en <http://lms.org.mx/?q=content/proyectos-aceptados-2019>, por lo que se agradecen los materiales y soportes brindados.

Este trabajo ha sido escrito como parte de los requisitos de graduación para el programa de Licenciatura en Ingeniería en Ciencias de la Computación de la Benemérita Universidad Autónoma de Puebla. El periodo de investigación y redacción de este trabajo de fin de grado ha durado desde noviembre de 2019 hasta junio de 2021.

El proyecto se llevó a cabo con el propósito de mejorar el algoritmo para la clasificación de desnudos del *Laboratorio de Análisis Forense Digital para la Detección de Delitos en Internet* a cargo del Dr. Colmenares. La tarea principal del trabajo se basó en lo encontrado en el estado del arte. El proceso de investigación y desarrollo ha sido arduo, pero me ha permitido cumplir con dicho propósito. Afortunadamente, el Dr. Luis Enrique Colmenares Guillen, mi asesor, ha estado disponible y dispuesto a ayudarme con todas mis cuestiones.

Me gustaría, por tanto, dar las gracias a mi asesor por su excelente orientación y soporte durante todo el proceso de realización de mi trabajo. También me gustaría dar las gracias a Gustavo Aponte Pérez, ya que su soporte para la interacción con el LNS fue clave para la realización de este proyecto.

También me ayudó discutir sobre varios asuntos de mi tesis con mi novio, amigos y familia. Si alguna vez perdí el interés, ustedes me ayudaron a permanecer motivada.

Por último, pero no menos importante, quisiera agradecer a la Dra. Hilda Castillo Zacatelco. Por brindarme la oportunidad de trabajar con ella en mi servicio social y prácticas profesionales, así como por sus clases de sistemas operativos y compiladores; al Dr. Iván Olmos Pineda, por introducirme al mundo de la inteligencia artificial y al Dr. Gerardo Martínez Guzmán. Por sus clases de cálculo, probabilidad y estadística. A todos ellos gracias por inspirarme a ir más allá en mi aprendizaje.

Espero que disfruten la lectura.

Brisa Isabel Mendoza Romero

Resumen

Hoy día se consume contenido más rápido que nunca, con **4.57 billones** [1] de personas usando internet un promedio de **6.42**[1] horas a nivel mundial. Este crecimiento trae muchos beneficios y permite a las personas continuar conectadas y organizarse colectivamente en todo el mundo, pero también trae problemas como la propagación de noticias falsas, agresiones en línea como el odio racial y étnico, burlas sobre apariencias físicas y amenazas de violación; contenido explícito como la pornografía infantil, tortura animal, mutilación, asesinatos y suicidio, ya que el escudo del anonimato facilita que las personas se comporten a placer, muchas veces sin tener que afrontar consecuencias. En México se reporta que un **72.9%** de un total de **101.5** millones de personas usan internet y **23.9%**—equivalente a **17.7** millones— de la población de 12 años y más que utilizó internet entre julio de 2018 y agosto de 2019 fue víctima de ciberacoso, de los cuales **8.3%** fueron hombres y **9.4%** mujeres [2]. Además, México ocupa el primer lugar en abuso sexual infantil, con **5.4 millones** de casos por año [33].

Actualmente existen moderadores de contenido humanos, APIs que proveen dicho servicio y empresas que ofrecen uno o ambos servicios, sin embargo, muchas veces esto queda fuera del alcance de pequeñas organizaciones y/o empresas que no pueden costear dichos gastos.

La presente tesis proporciona un modelo clasificador de desnudez de código abierto en con un 95% de exactitud que puede ser implementado en sistemas que coadyuven a la reducción de la violencia digital (abuso infantil, pornografía de venganza y acoso) mediante moderación de contenido.

Índice

Dedicatoria -----	2
Prólogo-----	3
Resumen -----	4
Índice-----	5
Índice de Figuras -----	7
Índice de Tablas -----	8
Índice de abreviaturas -----	9
1 Introducción -----	10
1.1 Estructura del documento -----	11
1.2 Objetivo General-----	12
1.2 Objetivos Específicos -----	12
2 Conceptos fundamentales-----	13
2.1 Aprendizaje Profundo -----	13
2.1.1 Otros conceptos-----	17
2.2 Moderación de Contenido -----	18
3 Estado del Arte -----	20
3.1 Investigación-----	20
3.2 Comercial-----	21
3.2.1 Moderadores Humanos -----	22
3.2.2 APIs -----	23
3.3 Nociones generales -----	25
4 Metodología -----	26
4.1 Hardware -----	26
4.2 Software -----	26
4.3 Flujo de trabajo-----	27
4.3.1 Definición del Problema y Adquisición de conjunto de datos -----	27
4.3.2 Medida de éxito -----	30
4.3.3 Protocolo de evaluación -----	31
4.3.4 Preparación de los datos -----	31
4.3.5 Función de pérdida y activación de última capa-----	31

4.3.6 Modelo sobreajustado -----	33
4.3.7 Regularización del modelo -----	33
5 Resultados -----	35
5.1 Conjunto de datos para el desarrollo de los experimentos -----	35
5.2 Experimentos desarrollados -----	35
5.3 Modelo propuesto -----	38
5.4 Modelo propuesto vs APIs -----	39
5.5 Modelo propuesto vs modelos pre-entrenados -----	39
5.6 Comparativa CPUs vs GPUs -----	40
6 Conclusiones y trabajo a futuro -----	41
6.1 Conclusiones -----	41
6.2 Trabajo a futuro -----	41
7 Referencias -----	43
8 Apéndices -----	48

Índice de Figuras

Figura 1 - IA, AA, RN, AP-----	13
Figura 2 - Programación clásica vs Aprendizaje Automático -----	14
Figura 3 - Red Neuronal-----	15
Figura 4 - Aprendizaje automático-----	16
Figura 5 - Aprendizaje Profundo-----	16
Figura 6 - Fuentes de datos -----	16
Figura 7 - Unidades de Procesamiento -----	17
Figura 8 - Toolboxes-----	17
Figura 9 - Moderación de Contenido -----	18
Figura 10 - Exactitud APIs-----	24
Figura 11 - Flujo de trabajo -----	27
Figura 12 - Problema de clasificación, ¿frío o calor? -----	28
Figura 13 - Ejemplo de problema de regresión, temperatura -----	29
Figura 14 - Proceso de limpieza realizado al conjunto de datos seleccionado -----	35
Figura 15 - Resultados de experimentos, exactitud -----	37
Figura 16 - Resultados de experimentos, pérdida -----	38
Figura 17 - Arquitectura del modelo propuesto -----	39
Figura 18 - Modelo Propuesto vs APIs -----	39
Figura 19 - Modelo propuesto vs modelos pre-entrenados -----	40

Índice de Tablas

Tabla 1 - Abreviaturas -----	9
Tabla 2 - Contenidos que se moderan-----	18
Tabla 3 - Pautas respecto a la desnudez en Facebook, Instagram, Snapchat y Reddit -----	19
Tabla 4 - Costo en dólares de moderación de contenido en imágenes y video -----	22
Tabla 5 - Costos por llamada a APIs -----	24
Tabla 6 - Tipos de problemas -----	28
Tabla 7 - Activación de Última Capa y Función de Pérdida -----	32
Tabla 8 - Arquitecturas experimentadas -----	36
Tabla 9 - Relación épocas, tamaño de la imagen, conjunto de datos y tipo de problema --	36
Tabla 10 - Comparativa CPUs vs GPUs -----	40

Índice de abreviaturas

Tabla 1 - Abreviaturas

Abreviatura	Significado
AA	Aprendizaje Automático.
ACORDE	Reconocimiento De Contenido Para Adultos Con Redes Neuronales profundas, por sus siglas en inglés.
AP	Aprendizaje Profundo.
CNN o ConvNet	Red neuronal convolucional, por sus siglas en inglés.
CPU's	Unidades de procesamiento central, por sus siglas en inglés.
GPUs	Unidades de procesamiento de gráficos, por sus siglas en inglés.
IA	Inteligencia Artificial
MAP	Usuarios Activos al mes, por sus siglas en inglés
MIL	Aprendizaje de múltiple instancia, por sus siglas en inglés.
NSFW	No apto para el trabajo, por sus siglas en inglés
ROI	extracción de regiones de interés, por sus siglas en inglés
SPIEC	Imágenes De Explotación Sexual De Niños, por sus siglas en inglés
SVM	Máquinas de vectores de soporte, por sus siglas en inglés.
TPUs	Unidades de procesamiento tensorial, por sus siglas en inglés.
UMA	Unidad de masa atómica.

1 Introducción

En poco más de una década, los encuentros con información y personas que estuvieron una vez esparcidos por la web, han sido reunidos en gran parte por un pequeño grupo de empresas en unas cuantas redes sociales. Hoy día se habla de plataformas, de hecho, cuando estas plataformas se comparan, lo hacen con sus contrapartes menos reguladas, grandes vs las más pequeñas, convencionales vs marginales, ya no se comparan con la web abierta. El sueño de la web abierta enfatizaba oportunidades nuevas, ampliadas y sin obstáculos para el conocimiento y la socialización. El acceso al público no estaría mediado por los editores y emisoras que desempeñaron papeles de guardianes en el siglo anterior. El poder de hablar se distribuiría más ampliamente, con más oportunidades para responder, deliberar, criticar, burlarse y contribuir. Esta cultura participativa, muchos esperaban, sería más igualitaria, más global, más creativa y más inclusiva. Las comunidades podrían basarse en un interés compartido, y podrían establecer sus propias reglas y prioridades, mediante cualquier forma de consenso democrático. *“La web misma iba a ser la plataforma”*[32].

Así como las redes sociales y otras plataformas crecen en popularidad y accesibilidad, la cantidad de contenido que se sube a ellas también crece. Incluso antes de la pandemia, se consumía contenido más rápido que nunca, con **4.57 billones** [1] de personas usando internet un promedio de **6.42**[1] horas a nivel mundial, de los cuales **3.96**[27] son usuarios de redes sociales, sumando el **51%** de la población mundial [27]. Entre el top 20 de sitios más visitados, 3 son sitios pornográficos y 5 redes sociales [27]. Facebook por su parte tiene alrededor de **2.70** billones de usuarios activos al mes [25], aunque dicha cifra puede relacionarse con el aislamiento por COVID-19, en periodos anteriores, se habían registrado entre **2.10** y **2.40 MAP** [25] y aproximadamente **15 mil** moderadores de contenido [26]. Mientras este crecimiento trae muchos beneficios y permite a las personas continuar conectadas y organizarse colectivamente en todo el mundo, también trae problemas como la propagación de noticias falsas, agresiones en línea como el odio racial y étnico, burlas sobre apariencias físicas y amenazas de violación; contenido explícito como la pornografía infantil, tortura animal, mutilación, asesinatos y suicidio, ya que el escudo del anonimato facilita que las personas se comporten a placer, muchas veces sin tener que afrontar consecuencias.

En México se reporta que un **72.9%** de un total de **101.5** millones de personas usan internet y **23.9%**—equivalente a **17.7** millones— de la población de 12 años y más que utilizó internet entre julio de 2018 y agosto de 2019 fue víctima de ciberacoso, de los cuales **8.3%** fueron hombres y **9.4%** mujeres [2]. Dentro de estas situaciones se ha reportado que **32.8%** y **19.4%**—por mujeres y hombres respectivamente— se tratan de la recepción de contenido de carácter sexual, dañando no solo la salud física y mental de las víctimas sino también de los menores que se encuentran expuestos a dicho contenido. Además, México ocupa el primer lugar en abuso sexual infantil, con **5.4 millones** de casos por año [33].

Durante 2014, fue publicada la Ley General de los Derechos de Niñas, Niños y Adolescentes [3], donde se reconoce a niñas, niños y adolescentes como titulares de derechos humanos y se establece la creación del Sistema Nacional de Protección Integral de los Derechos de Niñas, Niños y Adolescentes y se respalda la generación de medidas contra el abuso/violencia sexual en contra de los menores, así como su protección en los artículos 47 (incisos I, III), 57 (XI), 103 (VII), 105 (III) y 148 (II). Además, a la fecha la denominada Ley Olimpia, que tipifica los actos sexistas y la difusión del discurso de odio contra las mujeres en medios de comunicación, ha sido aprobada por 25 estados de la república, con una pena que va de **3 a 8 años de cárcel** y hasta **2mil UMAs** a quien difunda contenido sexual sin autorización [4]. Sin embargo, aún se requiere que se comiencen a generar herramientas que se adecuen al contexto político-social que existe en México para la prevención y mitigación de este tipo de delitos.

Casi todas las plataformas de redes sociales son empresas comerciales y deben encontrar una manera de obtener ganancias, tranquilizar a los anunciantes y respetar un espectro internacional de leyes. Debido a ello, lo que termina siendo sus valores no es un núcleo moral que existe debajo de la presión pública, sino cualquier solución que pueda resolver sus conflictos actuales teniendo en cuenta las demandas económicas e institucionales competentes, presentada en un lenguaje de lo correcto [32].

Actualmente, grandes empresas ofrecen servicios para la clasificación de contenido, otras desarrollan sus propias herramientas, sin embargo, no hay versiones de código abierto disponibles. Por lo que es necesario desarrollar una metodología que contribuya a la moderación de contenido, en este caso específico, para la detección de desnudez.

Como antecedente se realizó una tesis previa, denominada *Sistema para detectar desnudos en Imágenes Digitales* (Tello Flores, 2011). A partir de este trabajo se propone la presente tesis con nombre **CLASIFICACIÓN DE DESNUDEZ BASADA EN APRENDIZAJE PROFUNDO**.

1.1 Estructura del documento

Para lograr los objetivos de este trabajo de tesis se seguirá un proceso de prueba y análisis que será la base de los capítulos que conforman esta tesis.

En el capítulo 1 se describen los objetivos del trabajo de tesis y la estructura del mismo.

En el capítulo 2 se describen conceptos relevantes para el aprendizaje profundo y la importancia de la moderación de contenido.

En el capítulo 3 se discutirá el estado del arte, donde se estudian los trabajos previos en el campo de la investigación y comercial para la clasificación y moderación de desnudos.

En el capítulo 4 se presentan los recursos de hardware y software utilizados, así como la metodología con la que se desarrolló el presente trabajo.

En el capítulo 5 se presentan los resultados obtenidos del presente trabajo de tesis.

Finalmente, en el capítulo 6 se exponen las conclusiones y trabajos futuros que surgen de esta tesis.

1.2 Objetivo General

Este trabajo tiene como objetivo emplear métodos del aprendizaje profundo para proporcionar un modelo clasificador de desnudez de código abierto en imágenes digitales que pueda ser implementado en sistemas que coadyuven a la reducción de la violencia digital (abuso infantil, pornografía de venganza y acoso) mediante moderación de contenido.

1.2 Objetivos Específicos

1. Obtener un conjunto de datos para el desarrollo de los experimentos.
2. Diseñar y ejecutar una serie de experimentos para proponer un modelo clasificador.
3. Propuesta de un modelo clasificador con una exactitud mayor al 80% [29].
4. Comparativa del modelo propuesto y APIs.
5. Comparativa del modelo propuesto y de modelos pre-entrenados.
6. Comparativa del mejor modelo ejecutado en CPUs y GPUs.

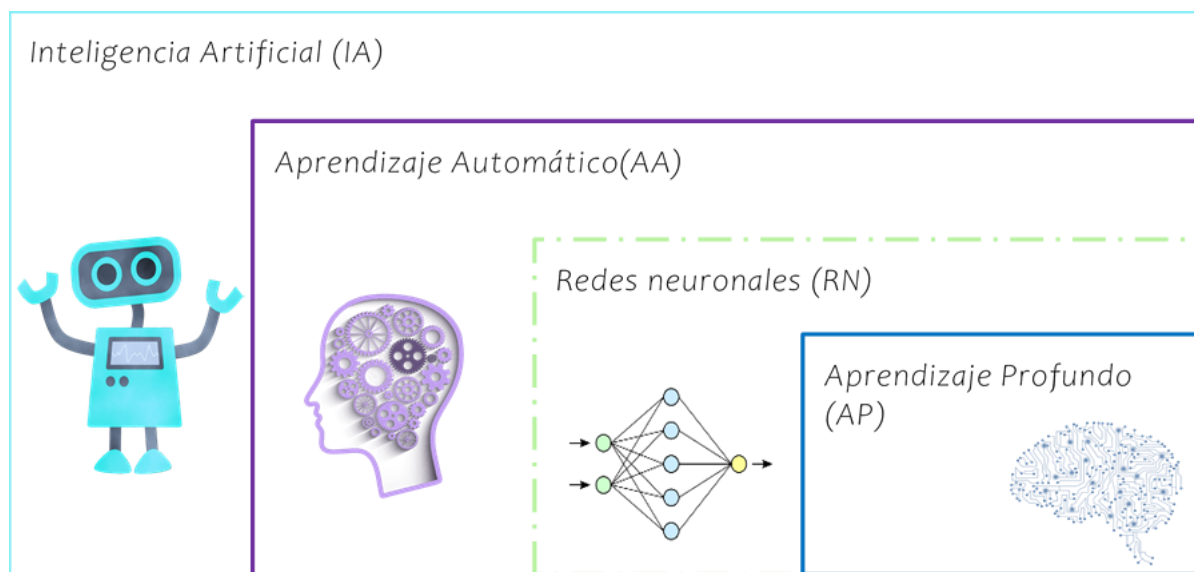
2 Conceptos fundamentales

Para comenzar, es necesario establecer los conceptos principales que se utilizan a lo largo de la presente tesis, dividiéndolos en dos categorías: aprendizaje profundo y moderación de contenido.

2.1 Aprendizaje Profundo

Antes de definir al Aprendizaje Profundo, es necesario conocer de dónde viene y para ello es necesario definir los términos de Inteligencia Artificial (en adelante IA) y Aprendizaje Automático (AA en adelante).

Figura 1 - IA, AA, RN, AP

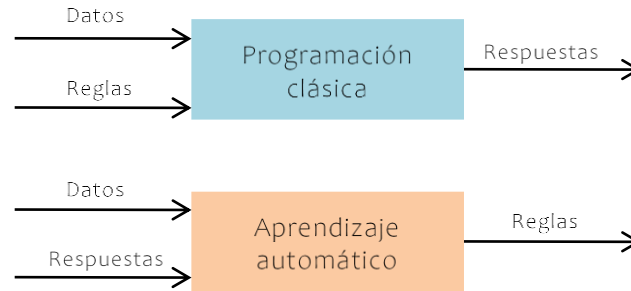


Por Inteligencia Artificial, se entiende cualquier técnica de programación que permita asemejar la inteligencia humana, por ejemplo, Watson, Siri, Cortana, Alexa entre otras aplicaciones que reconocen el lenguaje natural [34].

Como subconjunto de la IA, el Aprendizaje Automático comprende algoritmos capaces de completar una tarea específica, como la regresión lineal, árboles de decisión, SVM, redes bayesianas, bosques aleatorios algunos ejemplos de aplicación pueden ser las recomendaciones sobre un tema en específico o producto, diagnósticos médicos, detección de spam.

Una de las características más importantes de los algoritmos de AA consiste en el cambio de paradigma de la programación clásica, en la que el programador introducía un conjunto de reglas que producirían determinada respuesta dado un conjunto de datos.

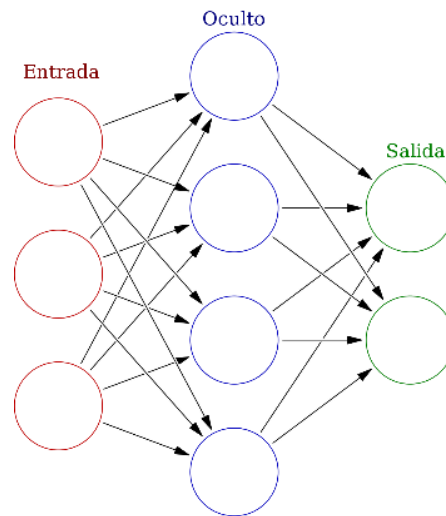
Figura 2 - Programación clásica vs Aprendizaje Automático



El AA rompe el paradigma de la programación clásica, dando como entrada a los algoritmos las respuestas esperadas de un determinado conjunto de datos y obteniendo como resultado las reglas para identificar estas respuestas.

Las **Redes Neuronales**, son modelos matemáticos que asemejan el funcionamiento de una neurona *real* en el cerebro. Consisten en un conjunto de unidades de procesamiento, llamadas neuronas artificiales y se organizan en capas. Hay tres partes normalmente en una red neuronal: una capa de entrada, con unidades que representan los campos de entrada; una o varias capas ocultas; y una capa de salida, con una unidad o unidades que representa el campo o los campos de destino. La red aprende examinando los registros individuales, generando una predicción para cada registro y realizando ajustes a las ponderaciones cuando realiza una predicción incorrecta. Este proceso se repite muchas veces y la red sigue mejorando sus predicciones hasta haber alcanzado uno o varios criterios de parada. Este proceso es lo que se conoce como entrenamiento y a medida que progresa, la red se va haciendo cada vez más precisa en la replicación de resultados conocidos. Una vez entrenada, la red se puede aplicar a casos futuros en los que se desconoce el resultado, lo que se conoce como clasificación [22].

Figura 3 - Red Neuronal



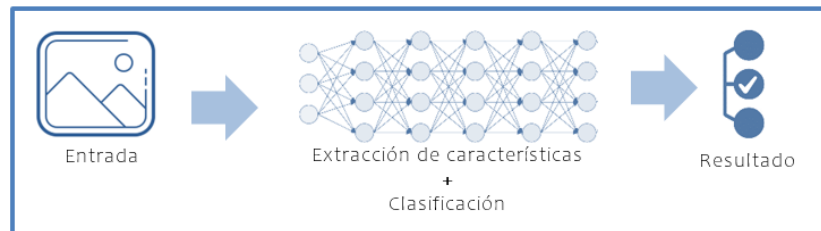
El **Aprendizaje Profundo** (en adelante AP), es un subconjunto del AA que utiliza RN y lo profundo de estos algoritmos viene del hecho de que entrenan en NN profundas, es decir, redes que tienen varias capas ocultas. El tipo más común de estas redes son denominadas convolucionales (CNN o ConvNet), además, al utilizar filtros bidimensionales, hace que esta arquitectura sea muy adecuada para procesar datos 2D, como imágenes. Algunas de las aplicaciones del AP son: colorear fotografías, traducciones automáticas, clasificación de objetos y/o personas en una imagen, generación de etiquetas que describen una imagen, etc.

La diferencia más significativa entre el AA y AP radica en que, el AA requiere de intervención humana para realizar *ingeniería de características* (simplificar un problema al expresarlo de una manera distinta, usualmente requiere conocimiento a profundidad del dominio), mientras que el AP elimina mayormente la necesidad de realizar dicho proceso debido a que las CNNs son capaces de extraer automáticamente características relevantes de datos crudos, convirtiéndolo en una mejor alternativa para un sistema de producción, ya que se puede proporcionar la entrada directamente al clasificador [5].

Figura 4 - Aprendizaje automático



Figura 5 - Aprendizaje Profundo



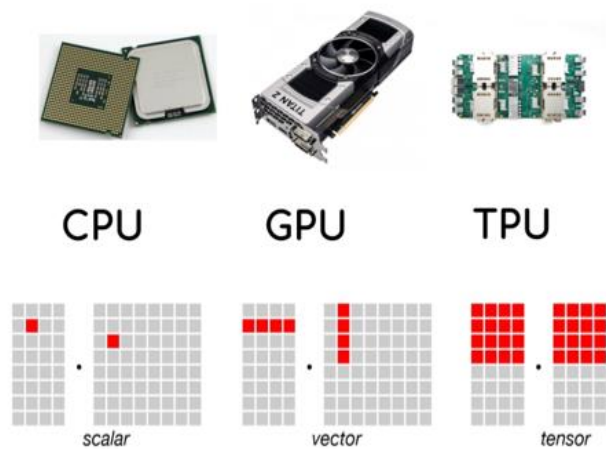
Además, existen 3 puntos técnicos [23] que son clave para el uso de AP en problemas de clasificación. El primero de ellos son los *macrodatos*, también llamados datos masivos, inteligencia de datos, datos a gran escala o big data (terminología en idioma inglés utilizada comúnmente) es un término que hace referencia a conjuntos de datos tan grandes y complejos que precisan de aplicaciones informáticas no tradicionales de procesamiento de datos para tratarlos adecuadamente [36], tales como el AA.

Figura 6 - Fuentes de datos



El segundo de ellos es el hardware puesto que ahora además de poder utilizar múltiples CPU's o GPUs, es posible utilizar TPUs, circuitos integrados de aplicación específica y acelerador de IA, desarrollado por Google para el AP y más específicamente optimizado para usar TensorFlow [37].

Figura 7 - Unidades de Procesamiento



La tercera razón por la que actualmente el AP ha experimentado un auge es el software, ya que otro beneficio que han traído las mejoras en el hardware ha sido la optimización de algoritmos. Además, se han creado toolboxes (cajas de herramientas por su traducción del inglés), que permiten la manipulación del hardware, los datos e incluso implementaciones en entornos de producción en un solo software.

Figura 8 - Toolboxes



La presente tesis, se basa en el AP para proponer un modelo de clasificación, es necesario definir algunos términos que se utilizan durante el desarrollo de este trabajo.

2.1.1 Otros conceptos

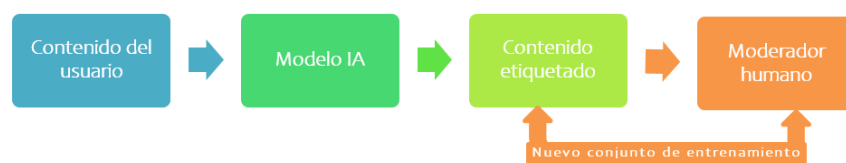
- *Época* se refiere a una iteración en la que el modelo procesa todo el conjunto de datos.
- *Tolerancia* es la cantidad de épocas que se permiten transcurrir sin un decremento en el valor de pérdida para el conjunto de validación. Es una función de TensorFlow.

- Se entenderá por truncación del conjunto de datos el evento en el que el conjunto de datos agota la memoria disponible y comienza a leer una y otra vez el conjunto de datos.
- Transfer learning, significa aprendizaje por transferencia por su traducción del inglés y es una técnica de AA en la que un modelo entrenado en una tarea (conjunto de datos) se reutiliza en una segunda tarea relacionada, por ejemplo, un conjunto que fue entrenado con fotos de gatos y perros reales que después se utiliza para identificar las mismas clases, pero ahora en caricatura [31].

2.2 Moderación de Contenido

A medida que maduran los entornos legales y sociales que rodean estos espacios de interacción con contenido generado por sus usuarios (UGC por sus siglas en inglés), existe una presión cada vez mayor para estas plataformas de vigilar, mantener y dar calidad a los contenidos que hospedan. Las acciones que se realizan en una plataforma para gestionar su UGC es lo que se conoce como Moderación de Contenido. Sin embargo, con el flujo de usuarios y contenido publicado en estas plataformas, es imposible para estas compañías inspeccionar cada una manualmente. En su lugar suelen aprovechar sistemas de AA para analizar automáticamente el contenido subido a sus sitios web. El contenido es reportado como que incumple una regla propia de la plataforma, por ejemplo, una imagen con contenido sexual explícito y moderadores humanos proveen el juicio final sobre si ese contenido debiera o no ser permitido en el sitio (Ver Figura 9) [30].

Figura 9 - Moderación de Contenido



A continuación, se describen las cuatro categorías del contenido que es moderado según el siguiente trabajo [28].

Tabla 2 - Contenidos que se moderan

Contenido abusivo	Agresión cibernética, discursos de odio y publicación de información privada o identificante sobre un individuo u organización, con el propósito de intimidar, humillar o amenazar.
Contenido falso o engañoso	Se trata de noticias o artículos que tienen como objetivo utilizar una red social para acelerar la propagación de información falsa, más comúnmente, lo conocemos como <i>fake news</i> .

Desnudez/contenido explícito	Desnudos y contenido sexualmente explícito.
Estafas/piratería	Este tipo de contenido intenta que los usuarios abandonen un sitio web y utilicen uno externo. En este sitio web, el usuario normalmente será engañado para que envíe información personal o dinero a una parte no deseada. Esto a menudo se logra imitando las URL del sitio web o prometiendo al usuario una mejor oferta en la compra mediante el uso del sitio web externo.

El presente trabajo de tesis se encuentra dentro de la Tabla 2, la categoría *desnudez/contenido explícito*. Antes de continuar es necesario establecer, ¿Qué constituye desnudez/pornografía? Incluso la definición básica es diferente para distintas plataformas. A continuación, se muestran las pautas de contenido para algunas redes sociales según [28]:

Tabla 3 - Pautas respecto a la desnudez en Facebook, Instagram, Snapchat y Reddit

	Facebook	Instagram	Snapchat	Reddit
Actividad sexual	La mayor parte del contenido es removido salvo que sea contenido digital con fines educativos, humorísticos o satíricos.	Prohibida en cualquier forma.	Prohibida en cualquier forma.	Cualquier forma de pornografía y desnudez a menos que sea ilegal.
Desnudez	Sin desnudez explícita, salvo la que representa actos de protesta, mujeres amamantando, cicatrices post mastectomía, esculturas y otras representaciones de arte.	Sin desnudez explícita, salvo las mujeres amamantando, cicatrices post mastectomía, esculturas y otras representaciones de arte	Solo se permite si no se representa en contextos sexuales.	

Para los fines del presente trabajo se considera desnudez como la exposición de partes íntimas femeninas, masculinas o alguna forma de acto sexual.

3 Estado del Arte

El presente estudio del estado del arte se divide en dos partes principales, *Investigación*, en donde se listan modelos o metodologías propuestas por estudiantes e investigadores, mientras que en la sección *Comercial* se ahonda en las opciones empresariales existentes en el mercado.

3.1 Investigación

En [9] se observa que los primeros trabajos se enfocaron en encontrar personas desnudas en imágenes basadas en un modelo de estructura humana y que la mayoría de los algoritmos tradicionales se descomponen en dos frases: extracción de regiones de interés (ROI) de fondos complejos y cálculo de características híbridas de ROI extraídos para clasificación. Sin embargo, debido a la complejidad de la detección de la piel, problema usualmente atacado con técnicas relacionadas al color, [9] y [7] concuerdan en que los enfoques basados en el ROI generalmente tienen una capacidad de generalización muy limitada debido a la alta tasa de falsos positivos, especialmente en el contexto de playas o la práctica de deportes acuáticos. Además [9], implementa ACORDE (Reconocimiento De Contenido Para Adultos Con Redes Neuronales profundas por sus siglas en inglés), la arquitectura de ACORDE comprende una ConvNet [15] para la extracción de características y una red de memoria a corto plazo (LSTM) para el aprendizaje de secuencias [16].

En 2014 Santos y col [6], proponen la extracción de características usando información de color y textura, así como selección de características, enfocándose en señalar las más relevantes, usando SVM para la clasificación de imágenes. Además, remarcan Kalva y col. [11] la introducción de la idea de que, al dividir una imagen en zonas independientes, uno podría extraer características locales y descartar información irrelevante, técnica que utilizan en su trabajo.

Durante 2017, [7] remarca la existencia de herramientas forenses y comerciales para el reconocimiento de imágenes de carácter sexual y reconoce la previa aproximación a la problemática utilizando Redes Neuronales Convolucionales, sin embargo, ofrece como distintivo, la capacidad de no solo ejecutar sobre discos o equipos de cómputo como las herramientas previamente mencionadas, sino que también es pionera en detectar contenido SPIEC (Imágenes De Explotación Sexual De Niños, por sus siglas en inglés) [12] con la clasificación arriba de L8 según [13][14].

En 2018, [8], señala que:

- Srisaan [17] estableció el modelo de piel usando los parámetros en el espacio YCbCr, como la proporción del área de la tez, luego detectó imágenes pornográficas usando un **árbol de decisión**.

- Li Daxiang [18] adoptó el método **MIL** (Aprendizaje de Instancia Múltiple, por sus siglas inglés), que describe características visuales de bajo nivel como su bolsa de palabras visuales que entrena a un clasificador.
- Deselaers [19] presentaron un método para aprender un vocabulario visual específico de la tarea con el empleo de un modelo **log-linear** para discriminar imágenes pornográficas.
- Moustafa [20] entren las redes **AlexNet** y **GoogleNet** por separado, entregando como resultado el promedio ponderado de los resultados de las dos redes como resultado de la clasificación de imágenes.

Para posteriormente proponer un módulo de detección de órganos sensibles (pecho femenino expuesto y los genitales masculinos y femeninos), que está diseñado para el problema de imágenes pornográficas no detectadas sin una gran área de color de piel desnuda, presentando algunos problemas con la detección de desnudos en imágenes animadas (anime).

En 2019, [9] utiliza la metodología MIL, argumentando la innovación del uso de la misma en el área de reconocimiento de imágenes con contenido sexual, sin embargo, como se observa en [18], ya existía evidencia de su uso. El trabajo se basa en las anotaciones de los contenidos pornográficos clave en las imágenes durante el entrenamiento. En base a estas anotaciones, presenta una medida cuantitativa del grado de pornografía de las regiones en bolsas positivas. Con la medida de representación y pornografía basada en la región, formula la tarea de reconocimiento como un problema ponderado de MIL bajo el marco de una red neuronal convolucional.

En el mismo año, [10] utiliza el filtro liberado por Yahoo en 2016, ya que no propone una metodología o clasificador, sino que su objetivo es demostrar que se puede construir un sistema de filtrado de pornografía de venganza práctico y automatizado utilizando las técnicas forenses existentes. El prototipo tiene dos usos principales:

- 1) Encontrar imágenes porno de venganza de cuentas de redes sociales conocidas y
- 2) Encontrar nuevas cuentas de redes sociales que albergan imágenes pornográficas de venganza potenciales.

3.2 Comercial

Un ser humano puede decidir instintivamente si lo que está viendo es inapropiado o no. Sin embargo, este problema está lejos de resolverse cuando se trata de tener una IA que pueda decidir si una imagen o video es inapropiado. Muchas empresas están compitiendo por ser las mejores en lo que respecta a la aplicación de técnicas automatizadas, utilizando moderadores

humanos o un híbrido (de moderadores humanos e IA) para identificar si el contenido de una plataforma es seguro para difundir o debe eliminarse.

3.2.1 Moderadores Humanos

Nanonets [30] realizó un estudio con 7 agencias de subcontratación de moderadores. En los resultados del estudio se puede observar que no hay un estándar de costo con respecto a la moderación de contenido humana y que dicha variación coloca a esta, como una alternativa poco viable para organismos con bajo presupuesto.

Tabla 4 - Costo en dólares de moderación de contenido en imágenes y video

Proveedor	Costo en dólares	Tiempo de respuesta por imagen (tiempo para procesar 1 imagen)	Tiempo de respuesta por video (tiempo para procesar 1 min de video)	Tiempo de entrenamiento para moderadores y como manejan la subjetividad/ casos extremos.
Taskus	\$0.03/imagen	NA	NA	NA
Scale.ai	Express: \$0.20/imagen Standar:\$0.08/imagen	"Nuestro tiempo de respuesta depende del volumen de los datos enviados. Para nuestro servicio bajo demanda puede encontrar las estimaciones de tiempo de respuesta del mejor esfuerzo en nuestra página de precios. En general, son unos días hábiles. También podemos ofrecer acuerdos de nivel de servicio que aseguran la respuesta para mayores volúmenes de datos bajo nuestro servicio empresarial "		"Formamos a nuestros trabajadores en base a las instrucciones que definirías en la API. Puede echar un vistazo a nuestra documentación para ver cómo funciona docs.scale.ai. Todas las reglas y casos extremos se definen también en la API"
Webpurify	Image: 0.02; Video: 15¢/minuto	<2 mins	<1 hora	Capacidad para capacitar a los moderadores sobre criterios únicos del cliente; o criterios llave en mano disponibles para la moderación de contenido estándar
Foiwe	NA	NA	NA	NA
Olapic	NA	NA	NA	NA
Assivo	Proyectos de precio fijo; según las necesidades del cliente	NA	NA	Sin criterios iniciales; el cliente define las necesidades primero y luego se le da un piloto; Si el cliente está satisfecho, el equipo de moderación de contenido se activa.

UGC Moderators	0.01\$/imagen o 2.00\$/hora por moderador	sin tiempo específico	sin tiempo específico	Capacidad para capacitar a los moderadores sobre criterios únicos del cliente; o criterios clave disponibles para la moderación de contenido estándar
----------------	---	-----------------------	-----------------------	---

3.2.2 APIs

En 2018, Aditya Ananthram [29], realizó una comparación para encontrar la mejor API en el mercado para clasificar contenido NSFW (No apto para el trabajo, por sus siglas en inglés), las APIs evaluadas fueron: *Algorithmia*, *Clarifai*, *DeepAI*, *Google*, *Imagga*, *Microsoft*, *Nanonets*, *Amazon Rekognition*, *Sightengine* y *X-Moderator*, basándose su rendimiento en las siguientes categorías:

- Desnudez explícita
- Desnudez sugerente
- Pornografía / acto sexual
- Pornografía animada / simulación
- Sangriento / Violencia

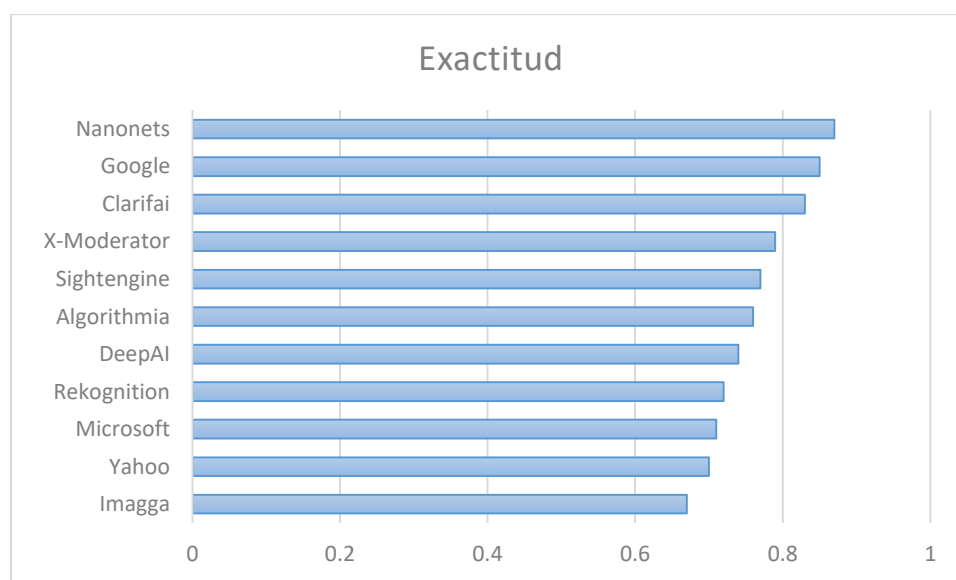
Para la evaluación, creó un conjunto de datos personalizado con el mismo peso (cantidad), para cada subcategoría. El conjunto de datos constó de 120 imágenes con 20 imágenes positivas para cada una de las cinco categorías descritas anteriormente y 20 imágenes negativas.

En general, el autor destaca lo siguiente:

- **Nanonets** no tiene el mejor desempeño en general en ninguna categoría. Sin embargo, es el más equilibrado en todas las categorías. Podría mejorar en identificar más imágenes como SFW (Aptas para el trabajo, por sus siglas en inglés). Es sensible a cualquier tono piel.
- **Google** tiene el mejor desempeño en la mayoría de las categorías de NSFW, pero el peor en la detección de SFW. Un punto a tener en cuenta es que las imágenes usadas para la prueba pertenecen a Google, lo que significa que las debería reconocerlas y clasificarlas correctamente. Usar imágenes conocidas para Google podría ser la razón del buen rendimiento de la API en este experimento.
- **Clarifai** destaca en la identificación de Gore (contenido violento) y lo hace mejor que la mayoría de las otras API, está bien equilibrado y lo hace bien en la mayoría de las categorías. Carece de exactitud en la identificación de desnudez sugerente y pornografía.

- **X-Moderator** es otra API bien equilibrada. Además de identificar a Gore, identifica bien la mayoría de los otros tipos de contenido NSFW. Obtuvo una exactitud del 100% en SFW, lo que lo distingue de sus competidores.
- **Sightengine** tiene una puntuación casi perfecta al tratar de identificar contenido NSFW. Sin embargo, no identificó correctamente ni una sola imagen de Gore.

Figura 10 - Exactitud APIs



De manera comercial, se considera el precio para decidir qué API utilizar. A continuación, se muestra una comparación de los precios que tiene cada uno de los proveedores. La mayoría de las API tienen una prueba gratuita con un uso limitado. Yahoo es el único que es de uso completamente abierto, pero al momento de redactar el presente trabajo los creadores han discontinuado su uso.

Tabla 5 - Costos por llamada a APIs

Proveedor	Llamadas gratis al API	Costo por 10k llamadas al API/mes (dólares)	Costo por 1M llamadas al API/mes (dólares)
Algorithmia	184	\$27	\$2,771
Calrifai	5000	\$12	\$1200
DeepAI	1000	\$10	\$1000

Google	1000	\$15	\$1500
Imagga	2000	\$14	\$1163
Microsoft	5000	\$10	\$1000
Nanonets	1000	\$9	\$999
Amazon Rekognition	1000	\$10	\$1000
Sightengine	2000	\$29	\$1599
X-Moderator	4629	\$18	\$1800

3.3 Nociones generales

La naturaleza subjetiva del contenido sexual hace que sea difícil declarar una API como la API de referencia para la moderación del contenido. Además, las APIs aún no están bien preparadas para manejar la subjetividad de sesgos raciales y contextos sociales, proporcionan etiquetas predeterminadas que pueden no ajustarse a lo que se requiere, sin embargo, una aplicación general, orientada a la distribución de contenido y desee un clasificador equilibrado preferiría utilizar la API de Nanonets, por otro lado, una aplicación dirigida a niños definitivamente sería lo más cautelosa posible y preferiría ocultar incluso el contenido mínimamente inapropiado, por lo que lo óptimo sería usar la API de Google.

Por su parte los algoritmos descritos en el Capítulo 3.1, son específicos para los conjuntos de datos con los que fueron entrenados y es poco probable que tengan buenos resultados al exponerse a otros conjuntos de datos puesto que no son modelos desarrollados para implementarse en un entorno de alta demanda.

El presente trabajo de tesis pretende proporcionar un modelo de código abierto que pueda ser implementado en sistemas que coadyuven a la reducción de la violencia digital (ciberacoso, pornografía de venganza, abuso infantil y/o pornografía infantil.) mediante moderación de contenido.

4 Metodología

4.1 Hardware

Para la implementación del presente trabajo de tesis, se contó con algunos recursos pertenecientes al Laboratorio Nacional de Supercómputo del Sureste de México, los cuáles se listan a continuación:

- 8 GPU (Graphics Processing Unit)'s GeForce GTX 1080 Ti con las siguientes características:
 - computeCapability: 6.1
 - coreClock: 1.6325GHz
 - coreCount: 28
 - deviceMemorySize: 10.92GiB
 - deviceMemoryBandwidth: 451.17GiB/s
- 16 CPU's Intel(R) Xeon(R) CPU (Central Processing Unit) E5-2630 v3 @ 2.40GHz

Para la implementación y visualización de los datos se utilizó un equipo propio, una Laptop Lenovo IdeaPad L340, con las siguientes características:

- Intel(R) Core(TM) i5-9300H CPU, 2.4GHz, 4 núcleo físicos, 8 lógicos
- 8 gb DDR3 de memoria RAM.
- 219 gb de disco duro de estado sólido.
- Intel HD Graphics.

4.2 Software

- **TensorFlow**¹. Es el framework con la comunidad y documentación más grande hoy día. Cuenta con versiones especializadas para trabajar con múltiples CPUs y/o GPUs, utilizando por defecto los recursos disponibles al momento de la ejecución. Además, es posible optimizar el tiempo de entrenamiento de los modelos convirtiendo los datos a un conjunto nativo de TF y cargándolo en memoria, es decir, realizando la lectura una sola vez. La presente tesis utiliza la versión **2.4** de este framework junto con la versión 11.0 de la librería de Cuda. Las librerías utilizadas de TensorFlow más importantes utilizadas para el desarrollo de modelos son:
 - **Keras**², una biblioteca de Redes Neuronales de Código Abierto escrita en Python. Está especialmente diseñada para posibilitar la experimentación en más o menos poco tiempo con redes convolucionales.

¹ <https://www.tensorflow.org/>

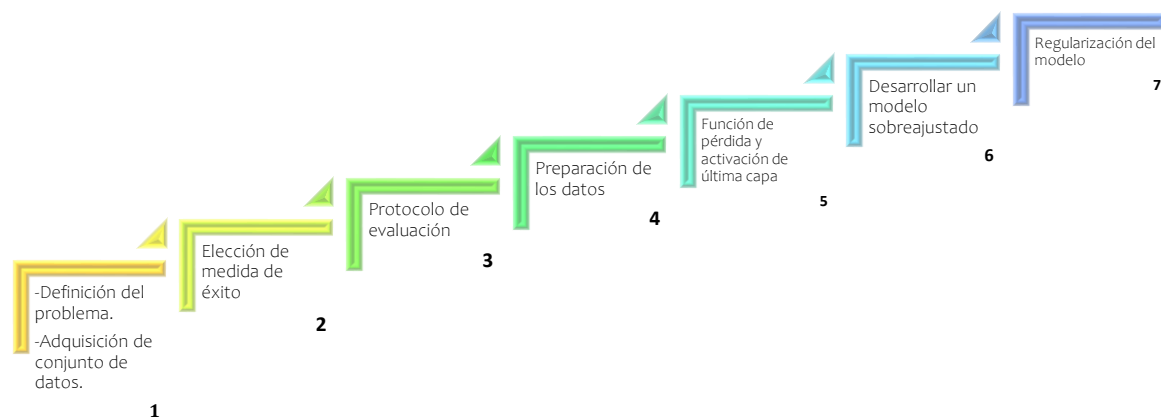
² <https://keras.io/>

- *Tensorboard*³, proporciona la visualización y las herramientas necesarias para experimentar con el aprendizaje automático como: seguir y visualizar métricas tales como la pérdida y la exactitud del modelo entre otras.
- *Python*⁴. Todo el código implementado en el presente trabajo de tesis se realizó con este lenguaje de programación en su versión 3.8 debido a la simpleza de su sintaxis y su fácil incorporación con el API de TensorFlow [21].

4.3 Flujo de trabajo

François Chollet, creador de la librería Keras, define en [23] el flujo de un proyecto de AA, y puede apreciarse en la siguiente figura:

Figura 11 - Flujo de trabajo



4.3.1 Definición del Problema y Adquisición de conjunto de datos

En esta etapa se decide de qué clase de problema se trata, tomando en cuenta el conjunto de datos con el que se cuenta, ya que los pasos posteriores se basan en la hipótesis de que las clasificaciones pueden predecirse a partir de las entradas dadas y que el conjunto de datos contiene suficiente información para que la red aprenda la relación entre ambas.

Existen cuatro tipos de problemas en el AA:

³ <https://www.tensorflow.org/tensorboard?hl=es-419>

⁴ <https://www.python.org/>

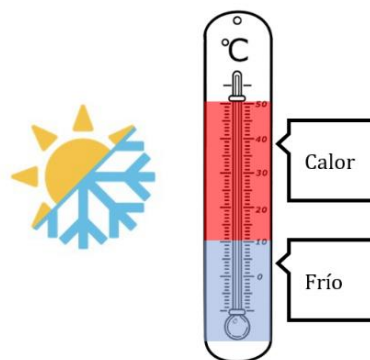
1. Un problema **binario**, es aquel en el que solo hay dos posibles salidas, tal es el caso de la clasificación de correos electrónicos como *spam* o *no spam*, ver columna 1 de tabla 6.
2. Cuando se tienen más de dos posibles salidas, pero solo se quiere identificar una en cada entrada, estamos hablando de un problema **multiclase, etiqueta única** (ver columna 2 de tabla 6).
3. Por otro lado, si en una sola entrada se quiere identificar más de una clase, estamos hablando de un problema **multiclase, multi-etiqueta** (ver columna 3 de tabla 6).

Tabla 6 - Tipos de problemas

Clasificación binaria	Multiclase, etiqueta única	Multiclase, multi-etiqueta
		
(0) Spam (1) No Spam	(1) Perro (2) Gato (3) Pez (4) Caballo (5) Ave (6) ...	1. Perro 2. Gato 3. Pez 4. Caballo 5. Ave 6. ...

Un problema de clasificación tiene etiquetas definidas, por ejemplo, si se habla de clima, hace calor o frío (problema binario), ver figura 12.

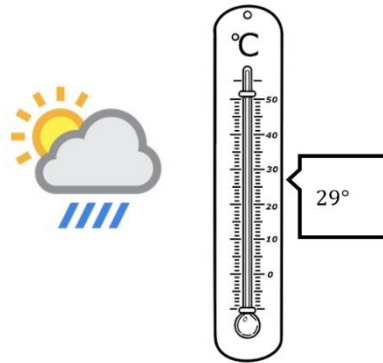
Figura 12 - Problema de clasificación, ¿frío o calor?



4. Por otro lado, si se habla de predecir la temperatura en grados, se trata de un problema de regresión, la cual es el cuarto tipo de problema, que puede ser *binaria*, cuando se

busca predecir un valor entre 0 o 1, y arbitraria cuando el rango a predecir es más grande a este último.

Figura 13 - Ejemplo de problema de regresión, temperatura



En cuanto al conjunto de datos, se realizó una investigación para encontrar uno con el mayor tamaño y calidad posibles, sin embargo, la disponibilidad de los recursos de desnudez como conjunto de datos para la investigación es limitada, por la sensibilidad del tema.

En enero de 2019 el científico de datos de Synced, Alexander Kim, introdujo una manera de eliminar la necesidad humana de la recolección de imágenes que incluyen o sugieren desnudos, a través de su proyecto *NSFW Datascraper* [24], que consiste en un conjunto de scripts que permiten la recopilación de decenas de miles de imágenes que pueden usarse para entrenar clasificadores de una variedad de contenido sensible. Como resultado de ejecutar los scripts, se obtiene un conjunto de datos con más de 220,000 imágenes en cinco categorías "vagamente definidas" según el propio Alexander, las cuales son:

- Pornografía
- Manga y otros estilos de dibujos pornográficos
- Sexy: imágenes sexualmente explícitas, pero no pornografía (playboy, bikini, voleibol de playa, etc.).
- Neutral: imágenes neutrales de personas cotidianas.
- Dibujos: dibujos neutrales (incluido el anime).

Un mes después, Evgeny Bazarov desarrolló otro proyecto inspirado en *NSFW Datascraper*, *NSFW data source URLs* [35], con este último es posible obtener hasta 1,300,000 imágenes de 159 categorías de pornografía, sin embargo, éste último conjunto no contiene muestras de imágenes neutrales, es decir, no es posible identificar un no desnudo, por lo que se decidió utilizar el conjunto de datos que proporciona [24].

En el Capítulo 4.3.1 de este capítulo se plantean dos problemas, binario y multietiqueta, con este último, ocurrieron truncamientos del conjunto de datos y como consecuencia la exactitud

que lograban los modelos que se habían implementado hasta el momento no superaron el 80%, por lo que, con motivo de obtener una exactitud más alta y sin truncamientos del conjunto de datos, se decidió particionar el conjunto de datos y priorizar la clasificación binaria, dado que el objetivo de la presente tesis es decidir si una imagen contiene desnudez o no (sí o no), para esta clasificación se utilizaron las clases etiquetadas como pornográfico y neutral, del conjunto de datos base, entrenando solo las arquitecturas que mejor exactitud habían tenido anteriormente.

4.3.2 Medida de éxito

Para determinar una métrica, primero es necesaria la observación y el análisis, por lo que para declarar como exitoso al presente proyecto de tesis, es necesario definir qué significa éxito para los términos de este. El objetivo específico número 2 nace a partir de esta decisión, que a su vez da paso a la elección de la función de pérdida seleccionada.

Existen numerosas métricas para evaluar modelos de clasificación, a continuación, se mencionan las más generales. La *exactitud*, es la fracción de predicciones que el modelo realizó correctamente, la *precisión* mide qué proporción de identificaciones positivas fue correcta, el *recuerdo*, qué proporción de positivos reales se identificó correctamente y el *f1-score* es la media armónica de precisión y recuerdo. Formalmente, sus definiciones son las siguientes según [38]:

- $exactitud = \frac{VP+VN}{VP+VN+FP+FN}$
- $precisión = \frac{VP}{TP+FP}$
- $recuerdo = \frac{VP}{VP+FN}$
- $F1 - score = \frac{2*precisión*recuerdo}{precisión+recuerdo}$

Donde:

- $VP = Verdaderos Positivos$
- $VN = Verdaderos Negativos$
- $TP = Total de Predicciones$
- $FN = Falsos Negativos$
- $FP = Falsos Positivos$

En el presente trabajo de tesis se utiliza la exactitud, ya que es la métrica proporcionada por el creador del conjunto de datos seleccionado.

4.3.3 Protocolo de evaluación

El protocolo de evaluación es una forma de medir el progreso que se ha logrado y puede usarse una de las siguientes tres opciones según [23]:

- a) *Mantener o retener un conjunto de validación*: el camino a seguir cuando tiene suficientes datos.
- b) *Validación cruzada de K iteraciones*: la elección correcta cuando se tienen muy pocas muestras para que la validación de retención sea confiable.
- c) *Validación cruzada repetida de K iteraciones*: para realizar una evaluación de modelo de alta precisión cuando hay pocos datos disponibles.

Dado el tamaño del conjunto de datos que se tiene, se decidió usar el protocolo a), puesto que se considera hay suficientes datos.

4.3.4 Preparación de los datos

En esta parte del proceso se realiza la extracción de características, en casos en los que se tiene un conjunto de datos muy pequeño, de lo contrario simplemente se da paso a que las entradas puedan ser entendidas por la red neuronal, es decir, los tensores (objetos matemáticos que almacenan valores numéricos y que pueden tener distintas dimensiones).

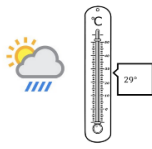
Dado que se utilizan redes convolucionales para el desarrollo del modelo clasificador, no es necesario realizar ingeniería de características como se menciona en el *Capítulo 2.1*, sin embargo, si es necesario limpiar un poco el conjunto de datos para no truncar el entrenamiento de los modelos.

4.3.5 Función de pérdida y activación de última capa

Los valores por elegir para la activación de la última capa y función de pérdida se relacionan con el tipo de problema. De acuerdo con [23] y retomando los ejemplos del *Capítulo 2.1*, en un problema de clasificación binaria como la de correos electrónicos en spam o no spam, se deberá utilizar la función *sigmoide* para la activación de la última capa y *entropía cruzada binaria* como función de pérdida, ver fila 2 de tabla 7. De la misma forma, en una clasificación de fotografías de animales, con una sola etiqueta, se deberá utilizar la función *sigmoide* para la activación de la última capa y *entropía cruzada categorizada* como función de pérdida, ver fila 3 de tabla 7, sin embargo, para una clasificación multietiqueta, se utilizaría la misma configuración que si se tratara de un problema binario, ver fila 4 de tabla 7. Para un problema de regresión binario, es posible utilizar la función *sigmoide* para la activación de la última capa y el ECM o la *entropía cruzada binaria* como función de pérdida, ver fila 6 de tabla 7. En cambio, para un

intervalo de valores arbitrarios no es necesario utilizar una función de activación y la función de pérdida se calcula con el ECM, ver fila 4 de tabla 7.

Tabla 7 - Activación de Última Capa y Función de Pérdida

Tipo de problema	Activación de última capa	Función de pérdida	Ejemplo
Clasificación binaria	sigmoide	entropía cruzada binaria	 Spam No spam
Multiclase, etiqueta única	softmax	entropía cruzada categorizada	 (1) Perro (2) Gato (3) Ave (4) ...
Multiclase, multietiqueta	sigmoide	entropía cruzada binaria	 1. Perro 2. Gato 3. Ave 4. ...
Regresión, valores arbitrarios	ninguna	ECM	 0° - 100°
Regresión, valores entre 0 y 1	sigmoide	ECM, entropía cruzada binaria	 Aprobado (1) Reprobado (0)

En el presente trabajo de tesis, plantearon dos problemas en el Capítulo 4.3.1, siendo el principal uno binario, por lo que se usó la función *sigmoide* para la activación de la última capa y la *entropía cruzada binaria* como función de pérdida.

4.3.6 Modelo sobreajustado

Para averiguar si un modelo es suficientemente grande, es decir, ¿tiene suficientes capas y parámetros para realizar la tarea de clasificación que se está trabajando? Por ejemplo, una red de una sola capa oculta es suficiente en un conjunto de datos pequeño, pero es poco probable que resuelva el problema en uno más grande; es necesario recordar que la tensión universal en el AA (y por consecuencia el AP) es entre optimización y generalización; el modelo ideal es aquel que se encuentra justo en el centro del desajuste y el sobreajuste. Para encontrar este balance, primero hay que cruzarlo, es decir, crear un modelo con sobreajuste. Esto se puede lograr experimentando con las siguientes técnicas:

- a) agregar más capas,
- b) hacer las capas más grandes,
- c) entrenar por más tiempo (aumentar las épocas)

Se debe monitorear constantemente la métrica previamente elegida, cuando esta comienza a bajar en el conjunto de validación, se ha alcanzado el sobreajuste y puede comenzar la regularización del modelo.

4.3.7 Regularización del modelo

Este es el paso más extenso de todo el proceso, puesto que hay que modificar, entrenar y evaluar el modelo repetidamente hasta que se haya logrado el valor deseado en la métrica elegida. Algunas cosas que se pueden hacer en esta etapa son:

- a) Agrandar el conjunto de datos (en caso de que se tengan más datos a disposición).
- b) Para los casos en los que no es posible realizar lo descrito en a), es posible simular un conjunto de datos más grande mediante la técnica denominada *data augmentation*⁵ que consiste en aplicar una serie de modificaciones a las imágenes en el tiempo de entrenamiento (sin necesidad de almacenarlas permanentemente) para brindarle entradas más diversas al modelo clasificador, ejemplos de estas modificaciones son: rotaciones, acercamientos, recortes, cambios en espacios de color y/o deformaciones.
- c) Añadir *dropout*⁶, que consiste en definir un porcentaje de valores que se abandonan u omiten durante el entrenamiento de manera aleatoria.
- d) Agregar o quitar capas.

Para el presente trabajo, se utilizaron todas las técnicas de la siguiente forma:

⁵ Aumento de datos, por su traducción del inglés.

⁶ Abandono por su traducción del inglés.

- a) Se utilizaron distintos tamaños del conjunto de datos, de manera que los modelos que obtuvieron mejor exactitud con conjuntos de datos más pequeños se entrenaron en el conjunto más grande para observar si el modelo obtenía una mejor exactitud.
- b) Los modelos que se entrenaron usando conjuntos de datos menores al más grande disponible fueron aumentados mediante rotaciones y acercamientos.
- c) Los mejores modelos fueron entrenados con y sin dropout de 0.5 previo a la capa de clasificación.
- d) Conforme se iba refinando el modelo se fue observando el desempeño de las distintas arquitecturas propuestas con más o menos capas.

5 Resultados

En este capítulo se detallan los resultados de la presente tesis con relación a los objetivos definidos en el Capítulo 1.

5.1 Conjunto de datos para el desarrollo de los experimentos

Cómo se mencionó previamente, al tratarse de un proyecto basado en AA, no fue necesario utilizar extracción de características, sin embargo, si se realizaron labores de limpieza en el conjunto obtenidos como producto de ejecutar los scripts de [24]. En la figura 14 se ilustra el proceso que se siguió para poder utilizar los datos, eliminando los archivos corruptos y los archivos con extensión .gif para posteriormente redimensionarlos y separarlos en el set que se usó para el entrenamiento y el que se usó para la validación.

Figura 14 - Proceso de limpieza realizado al conjunto de datos seleccionado



Al final del proceso, se obtuvo un conjunto de entrenamiento de 107,206 imágenes, 80% para el entrenamiento y 20% para la validación.

Posteriormente las imágenes se alimentaron a los modelos por medio de lotes, que se van optimizando a medida que se cargan en memoria y que no es necesario volver a cargar durante el entrenamiento, finalizando en la creación de un conjunto de datos nativo de Tensorflow que se compone de pares de imagen-etiqueta.

5.2 Experimentos desarrollados

Como resultado del uso de las técnicas descritas en el Capítulo 4.3.6, se obtuvieron un total de 13 arquitecturas, las cuales se pueden apreciar en la tabla 8. Es importante remarcar que para las arquitecturas b-i, se trata de redes en las que cada capa tiene el mismo número de filtros, mientras que para las arquitecturas j-m el orden el número de filtros por capa corresponde a estas de manera ascendente, ejemplo: para j, la capa 1 tiene 16 filtros, la 2da 32 y la 3ra cuenta con 64.

Tabla 8 - Arquitecturas experimentadas

ID	No. Capas	No. Filtros x capa
a	1	16
b	2	16
c	3	16
d	4	16
e	3	32
f	4	32
g	3	64
h	4	64
i	3	128
j	3	16_32_64
k	4	16_32_64_128
l	4	32_64_128_128
m	5	32_64_128_128_256

El no. de épocas se fue incrementando conforme se alcanzaba el límite previamente establecido; se empezó en 500 épocas y se fue doblando este límite, (ver columna 2 de la tabla 9), sin embargo, ningún experimento sobrepasó las 2mil épocas. De la misma forma incrementó la tolerancia, sin embargo, se limitó el valor a 250 debido para evitar largos periodos de entrenamiento sin mejoras en los resultados, ver columna 3 de la tabla 9.

Tabla 9 - Relación épocas, tamaño de la imagen, conjunto de datos y tipo de problema

Ronda	Épocas	Tolerancia	Tamaño img	Tamaño conjunto de datos	Tipo de problema
1	500	100	450	107,206	Multiclase, etiqueta única
2	1000	200	300		
3					
4	2000	250	150	11,582	
5				96,478	
6				56,975	Binario
8					
9	15	0	150	11,582	Multiclase, etiqueta única
10					
7					

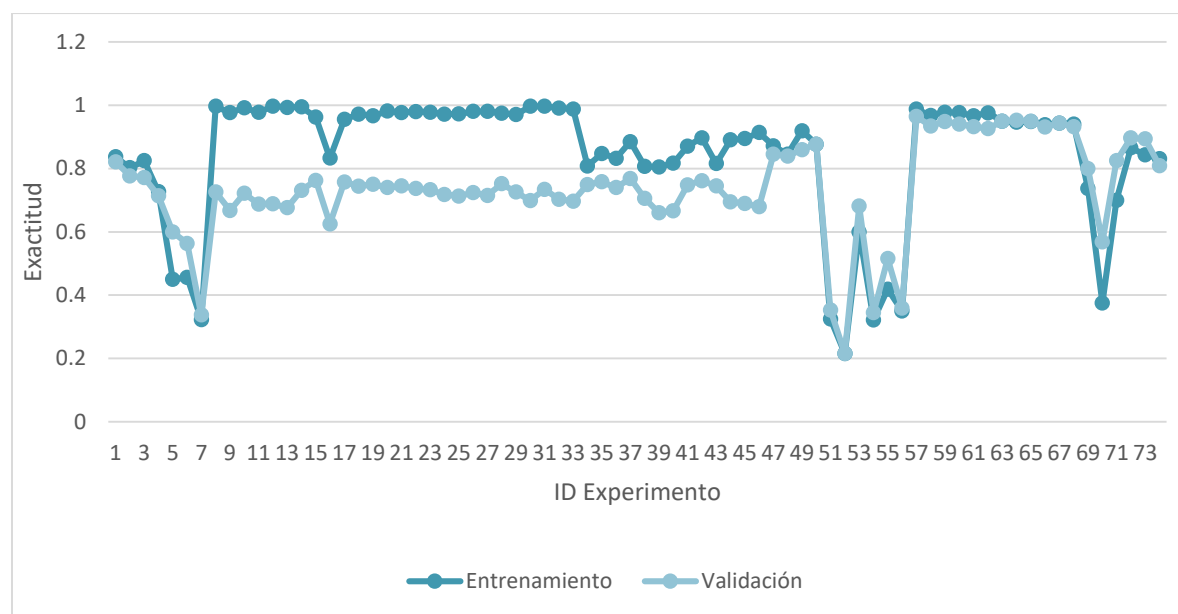
Las modificaciones relacionadas al tamaño de la imagen y del conjunto de datos tienen que ver con la capacidad del hardware para tenerlo en memoria, ya que en las primeras rondas

(conjunto de experimentos con la misma configuración) se tuvieron casos de truncamiento del conjunto de datos, lo que dio paso a experimentar con distintos tamaños del mismo.

Como resultado de combinaciones entre las arquitecturas descritas en la tabla 8 con las distintas configuraciones descritas en la tabla 9 se obtuvieron **74 experimentos**.

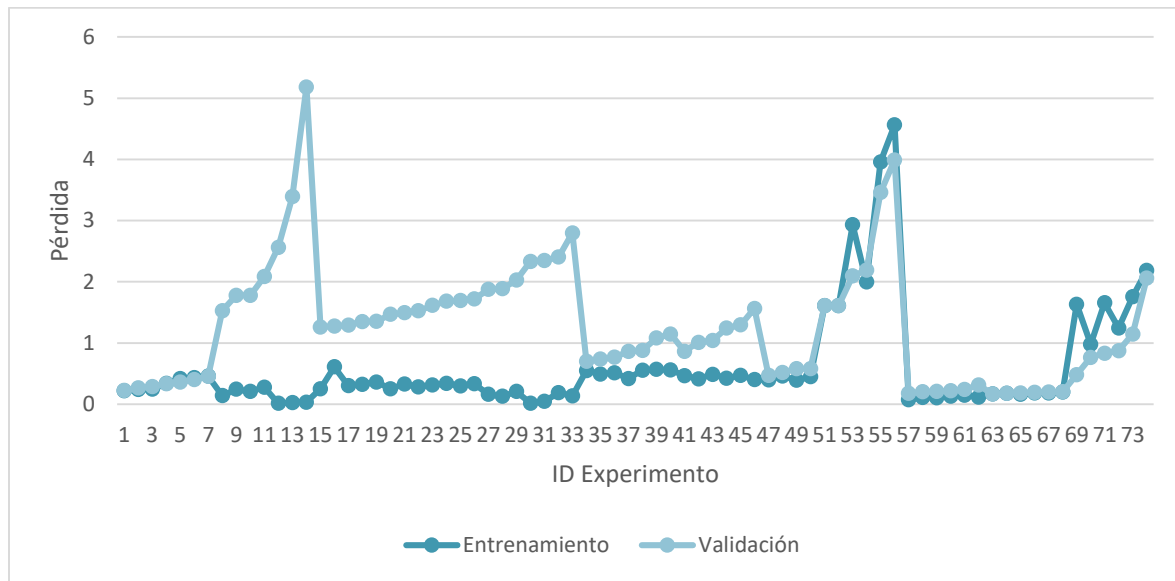
En la figura 15 se observan los resultados obtenidos en cada experimento con respecto a la exactitud, recordar que los experimentos en los que el valor obtenido en el conjunto de entrenamiento no es mayor al obtenido con el conjunto de validación, son aquellos que no presentan sobreajuste.

Figura 15 - Resultados de experimentos, exactitud



En la Figura 16 se observan los valores de pérdida para cada modelo, recordar que entre más pequeños sean los valores de la pérdida, más confiables son los resultados de exactitud.

Figura 16 - Resultados de experimentos, pérdida

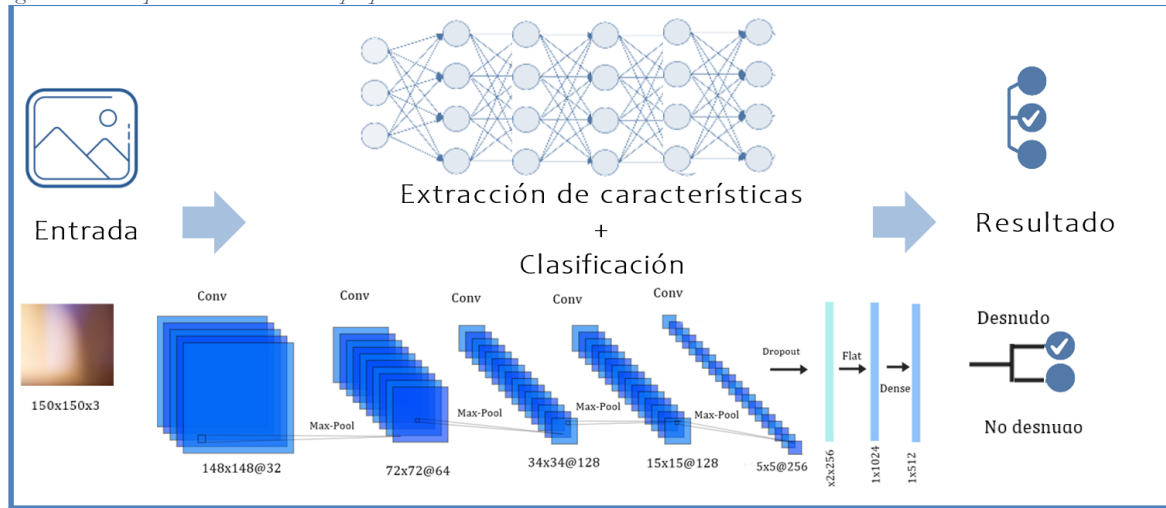


El desglose completo de los experimentos realizados se puede observar en el **Anexo 1**.

5.3 Modelo propuesto

El mejor modelo obtenido es resultado del experimento con ID 63, producto de la combinación de la arquitectura m, (descrita en la tabla 8), la configuración 9 (descrita en la tabla 9), y un entrenamiento de 1443 épocas, en la figura 17, se muestra su arquitectura, puesto en línea con la metodología del AA, tal como se describió en el Capítulo 2.1. Este modelo clasificador tiene una exactitud de 95%, (con 1.676% de pérdida), considerando como referencia el 91% reportado originalmente por el creador del script para generar el conjunto de datos base [24], el resultado obtenido es más confiable [40] debido a la limpieza y transformaciones (data augmentation) que se realizaron en el presente trabajo de tesis.

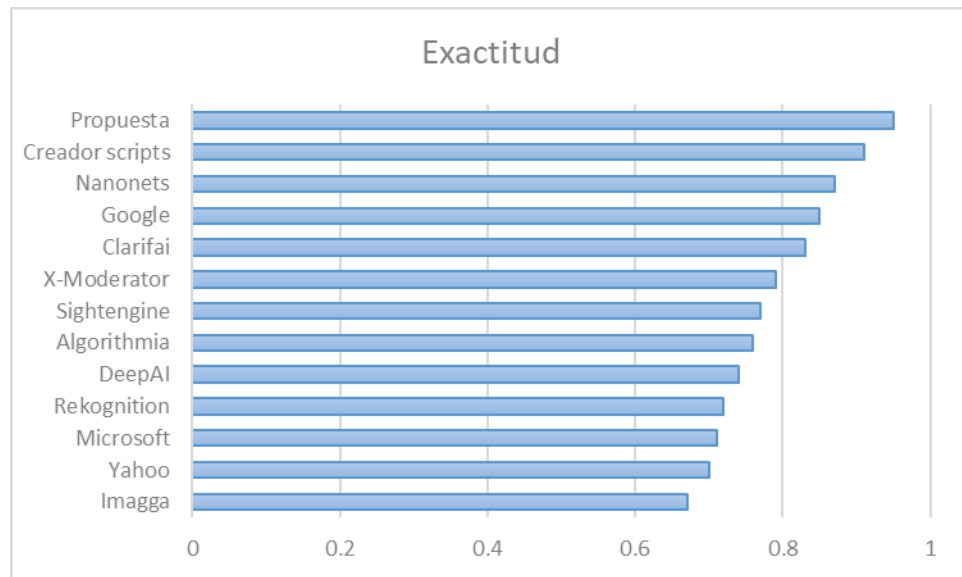
Figura 17 - Arquitectura del modelo propuesto



5.4 Modelo propuesto vs APIs

En el Capítulo 3.2.2 se presentaron los resultados de exactitud que [29] reporta para algunas APIs, mismos que se retoman en la figura 18, añadiendo el resultado reportado por [24] y el del modelo propuesto en la presente tesis, siendo este último el que reporta mayor exactitud.

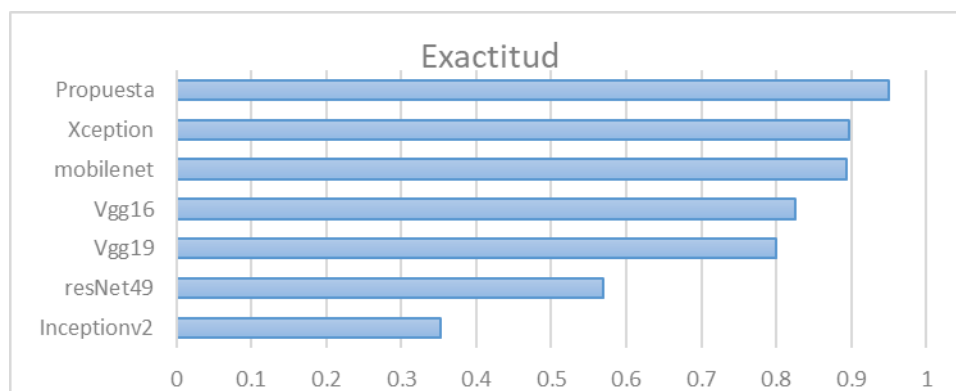
Figura 18 - Modelo Propuesto vs APIs



5.5 Modelo propuesto vs modelos pre-entrenados

En la figura 19 se muestra la comparativa de 69 de la exactitud de los experimentos con IDs, 70,71, 72, 73 y 51 (ver apéndice 1), con el modelo propuesto, sin embargo, la pérdida de los mismos es bastante alta.

Figura 19 - Modelo propuesto vs modelos pre-entrenados



5.6 Comparativa CPUs vs GPUs

En la *tabla 10*, se puede observar la **comparativa** de la exactitud y tiempo que tomó el mejor modelo encontrado en su ejecución en CPUs y GPUs, que a pesar de ser el mismo modelo se obtuvieron resultados con un margen de 2% en exactitud y 1h 24s en tiempo, siendo más confiable el resultado para los GPU's, ya que es menor la pérdida en el conjunto de validación y la exactitud del conjunto de entrenamiento no excede la del de validación por lo que no hay sobreajuste. Por último, se observa que con las GPU's se puede recorrer un mismo conjunto de datos poco más de 4 veces que con CPUs en un tiempo similar.

Tabla 10 - Comparativa CPUs vs GPUs

	CPU's	GPU's
Exactitud Entrenamiento	97%	95%
Exactitud Validación	94%	95%
Pérdida Entrenamiento	1.397%	1.691%
Pérdida Validación	2.355%	1.676%
Épocas	381	1443
Tiempo	4h 49m 36s	3h 49m 12s

6 Conclusiones y trabajo a futuro

6.1 Conclusiones

En primer lugar, el **conjunto de datos depurado** en la presente tesis con **107,206** imágenes etiquetadas como *pornografía*, *manga*, *sexy*, *neutral* y *dibujos*, así como los scripts para particionarlo pueden utilizarse para trabajos que sigan en la misma línea de investigación o similar.

En segundo lugar, se considera que el diseño y ejecución de los experimentos fue exitoso ya que, a pesar de los truncamientos del conjunto de datos, se logró obtener un modelo que cumple con el objetivo 3.

En tercer lugar, se concluye de manera exitosa la generación de un modelo clasificador con una exactitud por encima del 80%, con respecto a [29] tal como se muestra en el Capítulo 5.4, lo que demuestra que el modelo propuesto puede ser usado en un entorno de alta demanda como las APIs con las que se comparó. Además, al ser un modelo de código abierto puede ser implementado en sistemas que coadyuven a la reducción de la violencia digital (acoso, pornografía de venganza, abuso infantil, etc.) mediante moderación de contenido.

En cuarto lugar, se experimentó con la transferencia de conocimiento para realizar una comparación con las redes previamente entrenadas que provee el API de TensorFlow, sin embargo, se requiere refinar dichos experimentos para obtener resultados más confiables, es decir, con menor pérdida.

Finalmente, se obtuvo la comparación del mejor modelo ejecutado en CPUs y GPUs, siendo este último el de mejor desempeño y demostrando que al utilizar GPUs es posible recorrer un conjunto de datos más veces en un menor tiempo.

6.2 Trabajo a futuro

El trabajo a futuro consiste en mejorar los resultados del modelo para detectar desnudos con una pérdida menor al 1.676% reportado en el presente trabajo e integrarlo en el *Laboratorio de Análisis Forense Digital* y en ensambles de clasificadores de sistemas de moderación de contenido en México.

Para alcanzar esto se propone como primera instancia crear un conjunto de etiquetas más amplio y conciso, así como la recolección de muestras de dichas etiquetas ya que se considera esto puede hacer al clasificador más flexible a la hora de incrustarse en aplicaciones con distintos propósitos.

La *segunda instancia* es *agregar muestras del dominio en México* al conjunto de datos para disminuir la parcialidad racial en el desempeño general del modelo.

Finalmente, habría que hacer pruebas exhaustivas del desempeño del modelo en el sistema en que se integre y un análisis comparativo actualizado con relación desempeño de otras alternativas (APIs).

7 Referencias

- 1) Simon Kemp, S.K. (2020, 21 Julio). DIGITAL 2020: JULY GLOBAL STATSHOT. [Weblog]. Recuperado el 10 September 2020, de <https://datareportal.com/reports/digital-2020-july-global-statshot>
- 2) INEGI (2020). MÓDULO SOBRE CIBERACOSO 2019. Comunicado de prensa num. 163/20. Recuperado el 28 de junio desde: <https://www.inegi.org.mx/contenidos/saladeprensa/boletines/2020/EstSociodemo/MOCIB-A-2019.pdf>
- 3) Gobierno federal. (diciembre 2014). Ley General de los Derechos de Niñas, Niños y Adolescentes. Recuperado el 12 de November, 2019, de <https://www.unicef.org/mexico/informes/ley-general-de-los-derechos-de-ni%C3%B1as-ni%C3%B1os-y-adolescentes>
- 4) Wikipedia. (2021, 28 feb). Ley Olimpia (México). Recuperado el 3 de marzo 2021, de [https://es.wikipedia.org/wiki/Ley_Olimpia_\(M%C3%A9xico\)](https://es.wikipedia.org/wiki/Ley_Olimpia_(M%C3%A9xico))
- 5) Rafael zambrano. (2019, 13 octubre). Diferencias entre Machine Learning y Deep Learning. [Weblog]. Recuperado el 1º de junio 2021, de <https://openwebinars.net/blog/diferencias-entre-machine-learning-y-deep-learning/>
- 6) Clayton santos, Eulanda M. dos Santos, Eduardo Souto. (2014). Nudity detection based on image zoning. Recuperado el noviembre, 2019, de <https://ieeexplore.ieee.org/abstract/document/6310454>
- 7) Paulo Vitorino, sandra avila, mauricio perez, anderson rocha. (2017). Leveraging deep neural networks to fight child pornography in the age of social media. Recuperado el 9 Noviembre, 2019, de <https://www.sciencedirect.com/science/article/pii/S1047320317302377>
- 8) Ruolin zhu, xiaoyu wu, beibei zhu, liuyihan song. (2018). Application of Pornographic Images Recognition Based on Depth Learning. Recuperado el 9 de noviembre, 2019, de <https://dl.acm.org/citation.cfm?id=3209946>
- 9) Xin jin, yuhui wang, xiaoyang tan. (2019). Pornographic Image Recognition via Weighted Multiple Instance Learning. Recuperado el 8 noviembre, 2019, de <https://ieeexplore.ieee.org/abstract/document/8464261>

- 10) Manoranjan mohanty, ming zhang, giovanni russello . (2019). A Photo Forensics-Based Prototype to Combat Revenge Porn. Recuperado el 8 noviembre, 2019, de <https://ieeexplore.ieee.org/abstract/document/8695395>
- 11) P. R. Kalva, F. Enembreck, and A. L. Koerich, "Web Image Classification using Combination of Classifiers." IEEE Latin America Transactions, pp. 661-671, 2008. Consultado el 8 noviembre, 2019.
- 12) N. Al Mutawa, J. Bryce, V. N. Franqueira, A. Marrington, Behavioural evidence analysis applied to digital forensics: An empirical analysis of child pornography cases using p2p networks, in: IEEE Intl. Conference on Availability, Reliability and Security (ARES), 2015, pp. 293-302. Consultado el 8 noviembre, 2019.
- 13) M. Taylor, E. Quayle, Child pornography: An internet crime, Brunner-Routledge, 2003. Consultado el 9 noviembre, 2019.
- 14) M. Taylor, E. Quayle, G. Holland, Child pornography, the internet and offending, The Canadian Journal of Policy Research 2 (2) (2001) 94-100. Consultado el 9 noviembre, 2019.
- 15) Y. Le Cun, B. Boser, J. S. Denker, R. E. Howard, W. Hubbard, L. D. Jackel, D. Henderson, Handwritten digit recognition with a back-propagation network, Advances in neural information processing systems 2. Consultado el 9 noviembre, 2019.
- 16) S. Hochreiter, J. Schmidhuber, Long short-term memory, Neural computation 9 (8) (1997) 1735-1780. Consultado el 9 noviembre, 2019.
- 17) Srisaan C. A classification of internet pornographic images[J]. International Journal of Electronic Commerce Studies, 2016, 7(1):95-104. Consultado el 9 noviembre, 2019.
- 18) Daxiang Li, Na Li, Jing Wang, Tingge Zhu. "Pornographic images recognition based on spatial pyramid partition and multi-instance ensemble learning". Knowledge-Based Systems, 2015, 84: 214-223. Consultado el 9 noviembre, 2019.
- 19) T. Deselaers, L. Pimenidis, H. Ney. Bag-of-visual-words models for adult image classification and filtering. Proceedings of the International Conference on Pattern Recognition (ICPR), 2008.1-4. Consultado el 9 noviembre, 2019.

- 20) Mohamed N. Moustafa. Applying deep learning to classify pornographic images and videos[C]. Proceedings of the 7th Pacific- Rim Symposium on Image and Video Technology, 2015. Consultado el 9 noviembre, 2019.
- 21) Jakub protasiewicz. (Aug 31, 2018). Why Is Python So Good for AI, Machine Learning and Deep Learning. Recuperado el 7 noviembre, 2019, de <https://www.netguru.com/blog/why-is-python-so-good-for-ai-machine-learning-and-deep-learning>
- 22) Ibm corporation. (2021). El modelo de redes neuronales. Recuperado el 11 julio, 2021, de <https://www.ibm.com/docs/es/spss-modeler/SaaS?topic=networks-neural-model>
- 23) Chollet, Francois. 2017. Deep Learning with Python. New York, NY: Manning Publications. Opendgenus foundation.
- 24) Alexander kim. (2019). NSFW Data Scraper. Recuperado el 26 julio, 2020, de https://github.com/alex00okim/nsfw_data_scraper.
- 25) Dan noyes, D.M. (2020, August). The Top 20 Valuable Facebook Statistics – Updated Agosto 2020. [Weblog]. Recuperado el 3 septiembre 2020, de <https://zephoria.com/top-15-valuable-facebook-statistics/>
- 26) Charlotte jee archive page, C.J. (2020, 8 Junio). Facebook needs 30,000 of its own content moderators, says a new report. [Weblog]. Recuperado el 3 septiembre 2020, de <https://www.technologyreview.com/2020/06/08/1002894/facebook-needs-30000-of-its-own-content-moderators-says-a-new-report/>
- 27) Simon kemp, S.K. (2020, 21 July). MORE THAN HALF OF THE PEOPLE ON EARTH NOW USE SOCIAL MEDIA. [Weblog]. Recuperado el 10 septiembre 2020, de <https://datareportal.com/reports/more-than-half-the-world-now-uses-social-media>
- 28) Devin soni, D.S. (2019, July 22). Machine Learning for Content Moderation — Introduction an overview of machine learning systems for online content moderation. [Weblog]. Recuperado el 17 diciembre 2020, de <https://towardsdatascience.com/machine-learning-for-content-moderation-introduction-4e9353c47ae5>
- 29) Aditya Ananthram, A.A. (2018, 22 November). Comparison of the best NSFW Image Moderation APIs 2018. [Weblog]. Recuperado el 30 diciembre 2020, de

<https://towardsdatascience.com/comparison-of-the-best-nsfw-image-moderation-apis-2018-84be8da65303>

- 30) Arun Gandhi. (2019, 03 Enero). Content Moderation in 2020: Human vs AI. [Weblog]. Recuperado el 2 enero 2021, de <https://nanonets.com/blog/nsfw-content-moderation-in-2020-humans-vs-ai/>.
- 31) Jason brownlee. (2017, December 20,). A Gentle Introduction to Transfer Learning for Deep Learning. [Weblog]. Recuperado el 5 marzo 2021, de <https://machinelearningmastery.com/transfer-learning-for-deep-learning/>
- 32) Tarleton Gillespie (2018). CUSTODIANS OF THE INTERNET platforms, content moderation, and the hidden decisions that shape social media. London: Yale University Press.
- 33) Daniela vega. (2020,). Pandemia de COVID-19 agudiza abuso sexual infantil en México. [Weblog]. Recuperado el 29 noviembre 2020, de <https://www.unotv.com/nacional/pandemia-de-covid-19-agudiza-abuso-sexual-infantil-en-mexico/>
- 34) Alexander amini, ava soleimany. (2021). MIT 6S191 Introduction to Deep Learning. Recuperado el 22 abril, 2021, de <http://introtodeeplearning.com/>
- 35) Evgeny bazarov. (2019). NSFW data source URLs. Recuperado el 4 mayo, 2021, de https://github.com/EBazarov/nsfw_data_source_urls
- 36) Wikipedia. (2021). Macrodatos. Recuperado el 25 julio, 2021, de <https://es.wikipedia.org/wiki/Macrodatos>
- 37) Wikipedia. (2021). Unidad de procesamiento tensorial. Recuperado el 25 julio, 2021, de https://es.wikipedia.org/wiki/Unidad_de_procesamiento_tensorial
- 38) Shervin minaei. (2019, Oct 28). 20 Popular Machine Learning Metrics Part 1: Classification & Regression Evaluation Metrics. [Weblog]. Recuperado el 18 septiembre 2020, de <https://towardsdatascience.com/20-popular-machine-learning-metrics-part-1-classification-regression-evaluation-metrics-1ca3e282a2ce>
- 39) Marta carrasquilla. (2017, 4 enero). Cómo hacer un TFG. [Weblog]. Recuperado el 18 Marzo 2021, de <https://www.scribbr.es/category/estructura/>

40) Andrew ng. (2021, 24 Marzo). MLOps: From Model-centric to Data-centric AI. [Weblog]. Recuperado el 18 Julio 2021, de <https://www.deeplearning.ai/wp-content/uploads/2021/06/MLOps-From-Model-centric-to-Data-centric-AI.pdf>

8 Apéndices

Los 1s en las columnas *Lr*, *Dropout* y *Aumento* significan que se usó la característica en el experimento y los 0s indican lo contrario.

E.E. – Exactitud del conjunto de Entrenamiento

P.E. – Pérdida del conjunto de Entrenamiento

E.V. – Exactitud del conjunto de Validación

P.V. – Pérdida del conjunto de Validación

ID	Config	Arq	Clases	Lr	Dropout	Aumento	Épocas	E.E	P.E	E.V.	P.V.	Tiempo
1	1	d	5	1	0	0	361	0.8372	0.2179	0.8216	0.2247	22h 17m 41s
2	1	c	5	1	0	0	221	0.8028	0.2405	0.7768	0.2657	14h 4m 7s
3	1	b	5	1	1	0	321	0.8248	0.2453	0.7712	0.2881	18h 54m 0s
4	1	a	5	1	0	0	363	0.7269	0.3388	0.7144	0.3339	22h 46m 3s
5	1	d	5	1	1	0	324	0.4496	0.4163	0.5996	0.3617	18h 58m 25s
6	1	b	5	0	1	0	205	0.4564	0.4335	0.5636	0.4037	11h 58m 35s
7	1	c	5	1	1	0	200	0.3228	0.4597	0.3376	0.4562	13h 1m 78s
8	2	g	5	1	0	0	625	0.9966	0.1409	0.7256	1.528	3h 58m 50s
9	2	c	5	1	0	0	455	0.9766	0.2458	0.6675	1.779	1h 1m 56s
10	2	e	5	1	0	0	438	0.9922	0.2084	0.7219	1.779	1h 43m 9s
11	2	i	5	1	0	0	576	0.9775	0.2793	0.6881	2.084	7h 19m 22s
12	2	j	5	1	0	0	414	0.9972	0.0142	0.6888	2.562	1h 12m 23s
13	2	j	5	0	0	0	226	0.9934	0.027	0.6762	3.391	39m 50s
14	2	k	5	0	1	0	506	0.995	0.033	0.7312	5.183	1h 32m 14s
15	3	d	5	1	1	0	999	0.9625	0.2524	0.7625	1.261	2h 17m 33s
16	3	c	5	1	1	0	350	0.8328	0.6109	0.625	1.275	49m 29s
17	3	d	5	1	1	0	268	0.9556	0.3028	0.7575	1.294	39m 17s
18	3	f	5	1	1	0	260	0.9719	0.3269	0.7444	1.35	56m 21s
19	3	h	5	1	1	0	265	0.9666	0.362	0.7506	1.354	1h 42m 37s
20	3	c	5	1	1	0	873	0.9819	0.2517	0.74	1.468	2h 9m 51s
21	3	g	5	1	1	0	416	0.9769	0.3283	0.745	1.493	2h 40m 34s
22	3	e	5	1	1	0	537	0.9803	0.2828	0.7369	1.525	1h 51m 59s
23	3	e	5	1	1	0	258	0.9781	0.3129	0.7331	1.614	54m 32s
24	3	g	5	1	1	0	260	0.9722	0.3414	0.7181	1.683	1h 41m 45s
25	3	c	5	1	1	0	499	0.9728	0.2997	0.7131	1.694	1h 7m 58s
26	3	i	5	1	1	0	388	0.9806	0.3327	0.7237	1.721	4h 57m 11s
27	3	f	5	1	1	0	261	0.9806	0.1608	0.715	1.876	55m 31s
28	3	d	5	1	1	0	274	0.9753	0.1307	0.7519	1.889	37m 59s
29	3	h	5	1	1	0	258	0.9706	0.2069	0.7256	2.027	1h 38m 16s

30	3	j	5	1	1	0	496	0.9966	0.0183	0.6988	2.334	1h 34m 51s
31	3	k	5	1	1	0	329	0.9969	0.0448	0.7344	2.349	1h 0m 19s
32	3	g	5	1	1	0	257	0.9912	0.1905	0.7031	2.404	1h 38m 2s
33	3	c	5	1	1	0	258	0.9884	0.1341	0.6969	2.799	36m
34	4	d	5	1	1	1	819	0.8084	0.5498	0.7494	0.6989	42m 13s
35	4	f	5	1	0	1	516	0.8478	0.492	0.7581	0.7357	41m 30s
36	4	f	5	1	0	1	534	0.8325	0.5136	0.74	0.7678	49m 7s
37	4	h	5	1	0	1	366	0.8853	0.419	0.7688	0.8641	37m 27s
38	4	d	5	1	1	1	1012	0.8066	0.557	0.7056	0.8805	52m 18s
39	4	c	5	1	1	1	981	0.8053	0.5693	0.6606	1.082	50m 6s
40	4	c	5	1	1	1	708	0.8169	0.5604	0.6669	1.147	36m 30s
41	5	m	5	1	1	1	360	0.8709	0.4679	0.7487	0.8641	33m 53s
42	5	m	5	1	0	1	318	0.8972	0.4138	0.7619	1.011	30m 5s
43	5	m	5	1	1	1	394	0.8156	0.4877	0.745	1.041	36m 36s
44	5	m	5	1	0	1	346	0.8909	0.4245	0.695	1.242	32m 40s
45	5	l	5	1	0	1	276	0.8947	0.4706	0.69	1.296	25m 17s
46	5	l	5	1	0	1	289	0.9137	0.4048	0.6794	1.564	26m 52s
47	6	c	5	1	1	0	751	0.8719	0.3961	0.845	0.4732	1h 56m 45s
48	6	d	5	1	1	0	831	0.845	0.4617	0.8394	0.5181	2h 9m 7s
49	6	l	5	1	0	0	406	0.9194	0.3905	0.86	0.5799	1h 15m 21s
50	6	m	5	1	1	0	431	0.8769	0.4486	0.8769	0.5819	1h 19m 45s
51	7	n	5	0	1	1	14	0.325	1.609	0.3531	1.609	1m 4s
52	7	s	5	0	1	1	14	0.2156	1.609	0.2156	1.609	1m 2s
53	7	o	5	0	1	1	14	0.6	2.931	0.6812	2.099	55s
54	7	q	5	0	1	1	14	0.3219	1.998	0.3444	2.19	1m 2s
55	7	r	5	0	1	1	14	0.4187	3.96	0.5156	3.459	1m 8s
56	7	p	5	0	1	1	14	0.35	4.562	0.3594	3.99	1m 2s
57	8	c	2	1	1	1	395	0.9875	0.074	0.9644	0.1773	1h 19m 12s
58	8	d	2	1	1	1	782	0.9681	0.1117	0.9344	0.2058	2h 36m 36s
59	8	f	2	1	1	1	495	0.9781	0.1074	0.9488	0.2081	1h 50m 11s
60	8	h	2	1	1	1	364	0.9766	0.1304	0.9406	0.2213	1h 29m 41s
61	8	m	2	1	1	1	384	0.9663	0.147	0.9325	0.2416	1h 30m 23s
62	8	m	2	1	0	1	324	0.9759	0.1145	0.9262	0.3135	1h 18m 5s
63	9	m	2	1	1	1	1443	0.95	0.1691	0.95	0.1676	3h 49m 12s
64	9	h	2	1	1	1	795	0.9469	0.1756	0.9525	0.1756	2h 10m 43s
65	9	f	2	1	1	1	946	0.9481	0.1613	0.95	0.1835	2h 25m 7s
66	9	d	2	1	1	1	733	0.9381	0.1825	0.9312	0.1914	1h 48m 21s
67	9	m	2	1	0	1	707	0.9438	0.1853	0.9438	0.1972	1h 56m 4s
68	9	c	2	1	1	1	891	0.9406	0.1975	0.9319	0.2026	2h 11m

69	10	r	2	0	1	1	14	0.7375	1.633	0.8	0.4832	1m 26s
70	10	p	2	0	1	1	14	0.375	0.9782	0.5688	0.766	1m 6s
71	10	q	2	0	1	1	14	0.7	1.657	0.825	0.8316	1m 7s
72	10	s	2	0	1	1	14	0.8656	1.246	0.8969	0.8738	1m
73	10	o	2	0	1	1	14	0.8438	1.758	0.8938	1.146	55s
74	10	n	2	0	1	1	14	0.8313	2.186	0.8094	2.059	1m 25s