



UNIVERSIDAD AUTÓNOMA CHAPINGO

POSGRADO EN INGENIERÍA AGRÍCOLA Y USO INTEGRAL DEL AGUA



PREDICCIÓN SIMULTÁNEA DE CALIDAD INTERNA EN MÚLTIPLES FRUTAS CON ESPECTROSCOPIA VIS-NIR Y APRENDIZAJE PROFUNDO

TESIS PROFESIONAL

**Que como requisito parcial
para obtener el grado de:**

MAESTRO EN INGENIERÍA AGRÍCOLA Y USO INTEGRAL DEL AGUA

PRESENTA:

CARLOS JUÁREZ GONZÁLEZ

Bajo la supervisión de: CARLOS ALBERTO VILLASEÑOR PEREA, DOCTOR



Chapingo, Estado de México, noviembre de 2022

PREDICCIÓN SIMULTÁNEA DE CALIDAD INTERNA EN MÚLTIPLES
FRUTAS CON ESPECTROSCOPÍA VIS-NIR Y APRENDIZAJE PROFUNDO

Tesis realizada por **Carlos Juárez González** bajo la supervisión del Comité Asesor indicado, aprobada por el mismo y aceptada como requisito parcial para obtener el grado de:

MAESTRO EN INGENIERÍA AGRÍCOLA Y USO INTEGRAL DEL AGUA



DIRECTOR: _____
DR. CARLOS ALBERTO VILLASEÑOR PEREA



ASESOR: _____
DR. GILBERTO DE JESÚS LÓPEZ CANTEÑS



ASESOR: _____
DR. ARTEMIO PÉREZ LÓPEZ

CONTENIDO

ÍNDICE DE CUADROS	V
ÍNDICE DE FIGURAS	VI
DEDICATORIA	VIII
AGRADECIMIENTOS	IX
DATOS BIOGRÁFICOS	X
RESUMEN GENERAL	1
GENERAL ABSTRACT	2
CAPÍTULO 1. INTRODUCCIÓN GENERAL	3
1.1 Objetivo general	5
1.2 Objetivos específicos	5
1.3 Estructura de la tesis	5
CAPÍTULO 2. REVISIÓN DE LITERATURA	6
2.1 Fundamentos de la espectroscopía Vis-NIR	6
2.1.1 El espectro electromagnético	6
2.1.2 Espectroscopía UV-Vis	7
2.1.3 Espectroscopía NIR	7
2.2 Aplicaciones de la espectroscopía en la medición de calidad en frutas	9
2.2.1 Instrumentación	9
2.2.2 Evaluación de calidad en frutas	12
2.3 Modelación de datos espectrales	14
2.3.1 Preprocesamiento espectral	14
2.3.2 Métodos de calibración	15
2.4 Redes neuronales artificiales para análisis espectral	18
2.4.1 Redes neuronales artificiales	18
2.4.2 Aprendizaje profundo para el análisis espectral	20
2.5 Estado del problema	24
2.6 Literatura citada	28

CAPÍTULO 3. ARTÍCULO CIENTÍFICO.....	34
3.1 Resumen	34
3.2 Introducción	36
3.3 Materiales y métodos.....	38
3.3.1 Descripción de los datos	38
3.3.2 Aumento y separación de datos	39
3.3.3 Modelos 1D-CNN multisalida	40
3.3.4 Optimización de hiperparámetros.....	43
3.3.5 Evaluación de resultados	46
3.3.6 Implementación	46
3.4 Resultados y discusión	46
3.4.1 Análisis de los datos.....	46
3.4.2 Regularización de modelos	48
3.4.3 Optimización de los modelos.....	50
3.5 Conclusiones	58
3.6 Literatura citada.....	59

ÍNDICE DE CUADROS

CAPÍTULO 2

Cuadro 1. Bandas de absorción típicas a grupos químicos relacionados a atributos de calidad interna.....	13
Cuadro 2. Principales métodos de preprocesamiento espectral. Recuperado de Sirisomboon (2018).....	15
Cuadro 3. Efecto de algunos hiperparámetros en la capacidad del modelo. Recuperado de GoodFellow et al. (2016).	22

CAPÍTULO 2

Cuadro 1. Intervalo de valores usados para la optimización de hiperparámetros y de arquitectura 1D-CNN.	45
Cuadro 2. Referencias de calidad interna en seis frutas de la familia Cucurbitaceae.....	47
Cuadro 3. Hiperparámetros optimizados para tres modelos 1D-CNN	51
Cuadro 4. Errores de predicción de tres modelos 1D-CNN antes y después de la optimización Bayesiana. Los resultados se presentan como la media y desviación estándar de los errores combinados para las dos predicciones, en unidades normalizadas.....	52
Cuadro 5. Rendimiento de predicción de tres modelos 1D-CNN en comparación con PLSR y PCR. Los resultados se expresan como el promedio y la desviación estándar de 50 repeticiones del entrenamiento.	55
Cuadro 6. Hiperparámetros óptimos para la red de una capa convolucional de salida única para predicción de CA y CSS.	57

ÍNDICE DE FIGURAS

CAPÍTULO 2

Figura 1. El espectro electromagnético. Recuperado de Murray (2004).....	6
Figura 2. Bandas de absorción en el infrarrojo cercano. Recuperado de Murray (2004).	9
Figura 3. Modos de detección espectral: a) reflectancia, b) transmitancia y c) transreflectancia. Azul: detector, negro: base, gris: reflector. Modificado de Xie, Wang, Xu, Fu, & Ying (2016)	10
Figura 4. Arquitectura de una 1D-CNN para datos espectrales. Una neurona en la capa Conv2 cubre un campo receptivo de seis variables en el espectro original. Una vez que se activa la neurona el patrón del campo receptivo asociado se aprende como característica. Modificado de X. Zhang et al. (2019, 2021).	21
Figura 5. Comparación de tres métodos de ajuste automático de hiperparámetros. Las regiones coloreadas representan la pérdida. Modificado de Yu & Zhu (2020).....	24

CAPÍTULO 2

Figura 1. Separación y aumento de datos	40
Figura 2. Estructuras 1D-CNN de salida múltiple: (a) red 1; (b) red 3; (c) red 2. Los cuadros pequeños dentro de los rectángulos representan tensores, los módulos cian representan convolución general, los verdes convolución 1x1 y el gris agrupación máxima. Los colores diferente en los filtros representa su tamaño y las flechas en múltiples direcciones representan conexión entre capas.	42
Figura 3. Espectros promedio de seis productos de la familia <i>Cucurbitaceae</i> ..	48
Figura 4. Muestra de 50 espectros aumentados y asignados al conjunto de validación, antes (A) y después (B) de la estandarización.....	48
Figura 5. Criterios de selección de mínima fuerza de regularización de pesos para la optimización con datos RS (a, b, c) y K-S (d, e, f), de las redes 1 (a, d), 2 (b, e) y 3 (c, f). En rojo, el RMSE de calibración, azul RMSE de ajuste y negro RMSE	

de prueba. La línea discontinua verde representa el punto donde las redes comienzan a generalizar.	50
Figura 6. Evolución de la pérdida de validación combinada para CA y CSS durante un entrenamiento de modelos 1D-CNN antes (línea gruesa) y después (línea fina) de la optimización.	53
Figura 7. Mejores rendimientos de modelos de DL para predecir simultáneamente el contenido de agua (azul) y contenido de sólidos solubles (rojo) en seis productos de la familia Cucurbitaceae con espectros Vis-NIR. Resultados con modelo 1D-CNN de una capa convolucional (A, B), de 3 capas convolucionales (C, D) y de dos ramas (E, F).	56

DEDEDICATORIA

A mi mejor amiga, Reyna Citlali. Porque la amistad es la verdadera fortuna que un hombre puede tener.

AGRADECIMIENTOS

A la Universidad Autónoma Chapingo por el espacio brindado en sus instalaciones y personal docente que guio la formación de mi conocimiento.

Al Consejo Nacional de Ciencia y Tecnología (CONACyT) por el eficiente uso de los recursos del pueblo de México para apoyarme con la beca que me permitió realizar mis estudios de maestría.

A todos los miembros del comité asesor por la disposición de su tiempo, la confianza y paciencia en la supervisión de este trabajo.

A toda la comunidad científica por sus aportes al conocimiento humano que retomé para realizar esta modesta contribución.

Es especial a todos los miembros de mi apreciable familia, Padre, Madre y Hermano, por la incansable motivación e indudable confianza que tienen en mí.

DATOS BIOGRÁFICOS



Datos personales

Nombre	Carlos Juárez González
Fecha de nacimiento	10 de junio de 1996
Lugar de nacimiento	Huehuetla, Hidalgo
CURP	JUGC960610HHGRNR08
Profesión	Ingeniero Mecánico Agrícola
Cédula profesional	12317160

Desarrollo académico

Bachillerato	Centro de Estudios de Bachillerato CEB 6/7 "Gabino Barreda". Huehuetla, Hidalgo
Licenciatura	Universidad Autónoma Chapingo. Texcoco, México

RESUMEN GENERAL

PREDICCIÓN SIMULTÁNEA DE CALIDAD INTERNA EN MÚLTIPLES FRUTAS CON ESPECTROSCOPIA VIS-NIR Y APRENDIZAJE PROFUNDO

Acceder a la calidad interna de la fruta fresca de manera rápida y no destructiva es de gran interés en toda la cadena de producción, desde el momento de cosecha hasta el destino del producto. La espectroscopía de punto Vis-NIR permite relacionar atributos de calidad con absorciones vibracionales, pero se necesitan calibrar modelos específicos para frutas o atributos de interés. El aprendizaje profundo (DL) tiene el potencial para el modelado global, por ello, el objetivo de este estudio fue investigar el potencial del DL en la calibración de modelos multiproducto-multisalida. Se usó una base de datos de acceso libre con 300 espectros de absorbancia Vis-NIR de seis frutas de la familia Cucurbitaceae con mediciones de referencia de contenido de agua (CA) y contenido de sólidos solubles (CSS). Se realizó un aumento de datos y con ellos se optimizaron automáticamente tres redes neuronales convolucionales unidimensionales (1D-CNN) de diferente profundidad y de salida múltiple para comparar su rendimiento con modelos de salida única y regresión por mínimos cuadrados parciales (PLSR). Al comparar los modelos base y óptimos, la raíz del error cuadrático medio combinado (RMSE) mejoró hasta 13.5%. El modelo de una capa convolucional obtuvo el mejor rendimiento en comparación con los modelos de dos y tres capas convolucionales, logró un RMSE ~ 3% menor que un estudio previo con PLSR de salida única y espectros preprocesados y hasta un 38% menor que con espectros en bruto. La optimización de modelos de salida única solo mejoró menos del 1% el RMSE del mejor modelo de salida múltiple y demoró el doble de tiempo. Aunque los tiempos de optimización superan por mucho la simplicidad de PLSR, el DL permite desarrollar modelos globales de manera más fácil y precisa.

Palabras clave: calibración, CNN unidimensional, modelo global, multiproducto-multisalida.

GENERAL ABSTRACT

SIMULTANEOUS PREDICTION OF INTERNAL QUALITY IN MULTIPLE FRUITS WITH VIS-NIR SPECTROSCOPY AND DEEP LEARNING

Accessing the internal quality of fresh fruit quickly and non-destructively is of great interest throughout the production chain, from the moment of harvest to the destination of the product. Vis-NIR point spectroscopy allows quality attributes to be related to vibrational absorptions, but specific models need to be calibrated for fruits or attributes of interest. Deep learning (DL) has the potential for global modeling, therefore, the objective of this study was to investigate the potential of DL in the calibration of multi-product-multi-output models. A free access database with 300 Vis-NIR absorbance spectra of six fruits of the *Cucurbitaceae* family with reference measurements of water content (WC) and soluble solids content (SSC) was used. Data augmentation was performed and three one-dimensional convolutional neural networks (1D-CNN) of different depth and multiple output were automatically optimized to compare their performance with single output models and partial least squares regression (PLSR). When comparing the base and optimal models, the combined root mean square error (RMSE) improved until 13.5%. The one-convolutional layer model got the best performance compared to the two- and three-convolutional layer models, achieved ~3% lower RMSE than a previous study with single-output PLSR and preprocessed spectra and up to 38% lower than with raw spectra. Optimizing the single output models only improved the RMSE of the best multiple output model by less than 1% and took twice as long. Although optimization times far outweigh the simplicity of PLSR, DL makes global model development easier and more accurate.

Keywords: calibration, one-dimensional CNN, global model, multi-product-multi-output.

CAPÍTULO 1. INTRODUCCIÓN GENERAL

Con el aumento de la calidad de vida de la población la industria alimentaria enfrenta el reto de examinar las características externas e internas de las frutas y verduras para garantizar su calidad y satisfacer las demandas de los consumidores. El consumidor decide su compra por primera vez en base a la calidad externa (Butz, Hofmann, & Tauscher, 2005) y por segunda vez en base a la calidad interna del producto (Gutiérrez-Rico, 2017). La calidad externa se puede distinguir en productos uniformes y sin defectos aparentes en su superficie, pero la calidad interna depende de propiedades físico-químicas como el contenido de sólidos solubles (CSS), la acidez titulable (AT), la relación CSS/AT y la textura (Arendse, Fawole, Magwaza, & Opara, 2017). El CSS y el contenido de materia seca (CMS) comúnmente definen la fecha de cosecha, lo cual está relacionado con la futura calidad del producto (Mishra, Woltering, Brouwer, & Echtelt, 2021). El primer paso para el control es la medición y las propiedades internas son más difíciles de determinar por su carácter intrínseco en la fruta o verdura (definido solo fruta).

La determinación de parámetros de calidad interna ya se realiza por diversos métodos que requieren de una extracción del interior de la fruta, seguido de análisis de laboratorio, lo cual además de ser destructivo, es un proceso complicado, laborioso y lento. En contraste, se han desarrollado otras técnicas de análisis de calidad interna no destructivas, dentro de las cuales la espectroscopía visible e infrarroja cercana (Vis-NIR) es la más prometedora por sus ventajas de mínima preparación de la muestra, rapidez y amigable con el ambiente (Munawar, 2014), además de permitir la valoración de cada muestra desde el momento de cosecha y todo el proceso postcosecha (Walsh, Mcglone, & Han, 2020).

El desarrollo de la técnica Vis-NIR se basa en la recopilación de espectros de las muestras de las frutas y obtención de las referencias de calidad con los métodos tradicionales, para el posterior desarrollo de modelos de calibración que relacionen ambas mediciones. Este método presenta algunos problemas como

el gran volumen de datos (desde 250 hasta 2000 puntos por espectro), las múltiples relaciones entre ellos (multicolinealidad) y la presencia de ruido en los espectros, principalmente (Magwaza et al., 2012). El preprocesamiento de datos y los métodos quimiométricos permiten extraer la información subyacente para transformarla en los atributos de calidad, principalmente se usa la regresión lineal múltiple (MLR), regresión de mínimos cuadrados parciales (PLSR), regresión por componentes principales (PCR), redes neuronales artificiales (ANN) y métodos basados en kernel. Sin embargo, los métodos quimiométricos tradicionales presentan el inconveniente de requerir conocimiento previo sobre el problema, ante esto, el aprendizaje automático vino a brindar solución, especialmente el aprendizaje profundo (DL) que permite un análisis de extremo a extremo con mínimo conocimiento humano previo al análisis (Malek, Melgani, & Bazi, 2017).

Rápidamente se está comprobando que el DL supera los métodos quimiométricos y de aprendizaje automático tradicionales para el análisis espectral cualitativo y cuantitativo de gran variedad de frutas y otros productos de dominio comestible (X. Zhang, Yang, Lin, & Ying, 2021; Zhou, Zhang, Liu, Qiu, & He, 2019), aunque sigue siendo desafiante la configuración óptima de los hiperparámetros que involucran estos modelos. Los algoritmos de optimización automática pueden funcionar de manera comparable a los expertos humanos, a veces mejor o fallar catastróficamente en otros problemas (GoodFellow, Bengio, & Courville, 2016). Por tanto, se necesita investigar nuevos aportes del DL en el análisis espectral para contribuir en la aceptación generalizada de la espectroscopía Vis-NIR como método no destructivo de evaluación de calidad interna en las frutas.

Las calibraciones locales generan modelos sólidos, pero resultan costosas y tardadas. Por su parte, un modelo global puede contener muestras de diferentes orígenes, temporadas, cultivares o muestras con diversas fuentes de variabilidad biológica, físico-química y ambientales externas, mejorando así la solidez generalizada del modelo. El DL puede modelar simultáneamente dos atributos de calidad interna de múltiples productos de una familia vegetal. Por ello, en este

trabajo se utilizó una base de datos espectrales de la región Vis-NIR de seis productos de la familia *Cucurbitaceae* y dos referencias de calidad, CSS y contenido de agua (CA) (Kusumiyati, Hadiwijaya, Putri, & Munawar, 2021b), para investigar el potencial del DL en la calibración de un modelo global de predicción múltiple para múltiples productos.

1.1 Objetivo general

- Comprobar el rendimiento del aprendizaje profundo en el modelado global multiproducto para la predicción simultánea de contenido de sólidos solubles y contenido de agua en frutas con espectros de la región visible e infrarroja cercana.

1.2 Objetivos específicos

- Comparar el rendimiento de tres modelos de red CNN unidimensional (1D-CNN) de diferente profundidad.
- Comparar el rendimiento del modelado 1D-CNN de salida múltiple y de salida única.
- Comparar el rendimiento del modelado de extremo a extremo con métodos quimiométricos clásicos.
- Validar una metodología de optimización automática de hiperparámetros.

1.3 Estructura de la tesis

El presente documento se divide en tres partes principales. Comenzando con la presente introducción general para el tema de estudio (Capítulo 1), seguido de una breve revisión sobre los conceptos de espectroscopía Vis-NIR, sus aplicaciones en la evaluación no destructiva de calidad interna en frutas y el estado actual del problema (Capítulo 2). Por último, el Capítulo 3 presenta el artículo científico titulado: Predicción simultánea de calidad interna en múltiples frutas con espectroscopía Vis-NIR y aprendizaje profundo.

CAPÍTULO 2. REVISIÓN DE LITERATURA

2.1 Fundamentos de la espectroscopía Vis-NIR

Antes de 1800 solo se conocía la región visible del espectro, pero en esa fecha William Herschel descubrió que la temperatura de la luz aumentaba del violeta al rojo, incluso que era más alta en una región invisible más allá del rojo. Solo un año después Johann Ritter encontró otro componente invisible más allá del violeta. Habían descubierto la radiación infrarroja y ultravioleta, comenzando una nueva era en la comprensión de la luz y pronto se aclaró que esas regiones formaban parte de un espectro electromagnético (EM) (Ozaki & Huck, 2021).

2.1.1 El espectro electromagnético

El EM es una representación continua del flujo de paquetes de energía denominados fotones, que se propagan en forma de onda a la velocidad de la luz, pero con diferentes frecuencias ν proporcional a la energía de los fotones e inversamente proporcional a la longitud de la onda λ (Figura 1). Normalmente se utilizan los nanómetros para referirse a una banda del EM, pero en ocasiones resulta conveniente utilizar el número de onda $\bar{\nu}$ que se puede calcular como $\bar{\nu} = \frac{10,000,000}{\lambda \text{ (nm)}}$, y se expresa en cm^{-1} .

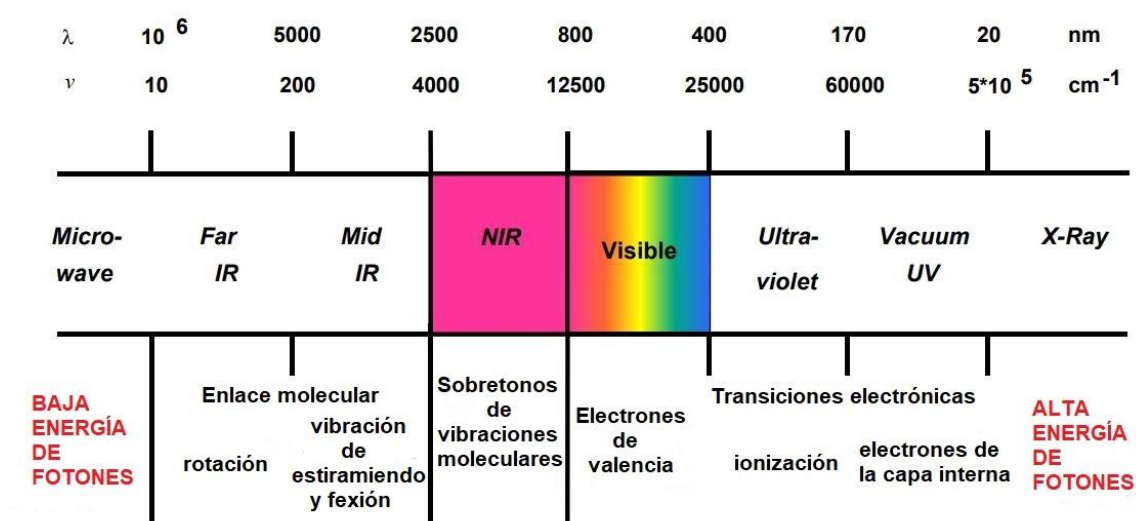


Figura 1. El espectro electromagnético. Recuperado de Murray (2004).

La capacidad de la luz para interactuar con la materia permite conocer la estructura molecular y atómica de los materiales. El estudio de esta interacción se llama espectroscopía y de acuerdo al espacio del EM en el que se trabaje se han desarrollado diferentes tipos de espectroscopía. La ultravioleta-visible (UV-Vis) es espectroscopía electrónica, la infrarroja (IR) espectroscopía vibratoria y la infrarroja cercana (NIR) se considera especial porque se encuentra entre la electrónica y la vibratoria (Ozaki & Huck, 2021)

2.1.2 Espectroscopía UV-Vis

La energía de los fotones en la región visible puede excitar electrones en el átomo y moverlos a un estado de mayor energía si tienen la energía equivalente a la diferencia entre los dos estados electrónicos. Sin embargo, rápidamente volverán al estado inicial y emitirán un fotón de la misma energía. Este movimiento se conoce como transición electrónica y determina la absorción de luz con la energía correspondiente, es decir, la longitud de la onda específica en la región UV-Vis.

Las transiciones electrónicas son únicas para cada elemento, por lo que se pueden usar el espectro de emisión de un material para identificar sus elementos constituyentes. Si embargo, en materiales compuestos la identificación de elementos con transiciones electrónicas es compleja, pero pueden relacionarse con transiciones entre orbitales moleculares y en muchos casos, identificar su naturaleza, especialmente, identificar pigmentos, colorantes, productos de alteración, entre otros compuestos orgánicos (Picollo, Aceto, & Vitorino, 2019).

2.1.3 Espectroscopía NIR

La energía de la región infrarroja no puede excitar electrones, pero puede provocar vibraciones de los enlaces químicos entre moléculas. Para que una molécula absorba radiación, su energía debe coincidir con la energía de vibración de la molécula. Los modos normales vibratorios dentro de una molécula pueden ser de dos tipos, de estiramiento y de flexión con variaciones en su simetría y su plano de movimiento. Todas las moléculas se encuentran en un estado vibratorio fundamental ($V=0$), dependiente de la temperatura ambiente, pero si un fotón con la energía correcta impacta con la molécula, esta cambiará al primer estado

excitado ($V=1$) y ocurrirá una transición fundamental. La energía requerida, y por tanto la frecuencia ν , para una transición fundamental en una molécula diatómica es clásica, se deriva de ley de Hooke como un modelo de bola y resorte que actúan como un armónico simple (Murray, 2004):

$$\nu = \frac{1}{2\pi} \sqrt{\frac{k}{\mu}}$$

Donde $\mu = m_1 m_2 / (m_1 + m_2)$ y k es la constante de fuerza del enlace. Esto indica que la masa y la rigidez del enlace determinan la longitud de onda de absorción.

En realidad, los enlaces no se comportan exactamente como en el caso del resorte, sino que se comportan como un oscilador anarmónico. La mecánica cuántica solo permite niveles discretos de energía vibratoria, y para este caso las transiciones entre los niveles de energía consecutivamente son (Meritxell-Gago, 2011):

$$\nu_{n+1} = \nu_{fundamental} \left[1 - y \left(\nu_n + \frac{1}{2} \right) \right]$$

Donde y es conocida como constante de anarmonicidad y n es el nivel de energía. Aunque es posible cualquier frecuencia vibratoria, solo se pueden observar transiciones $\Delta\nu = \pm 2, \pm 3, \dots$ conocidas como sobretonos. Cada sobretono consecutivo se vuelve más anarmónico, más débil, e involucra saltos de energía más grandes que corresponden a los fotones más energéticos de la región NIR.

El espectro NIR consta de dichos sobretonos y combinaciones de vibraciones fundamentales del infrarrojo medio (MIR), especialmente de los enlaces más anarmónicos como C-H, O-H y N-H, debido a la diferencia de peso atómico entre los átomos. Por ello, esta región es la más adecuada para en análisis de la composición bruta de los alimentos ya que sus componentes como agua, lípidos, carbohidratos y proteínas se atribuyen directamente con los enlaces HOH, CH₂, CHOH y CONH, respectivamente (Murray, 2004). Otros grupos funcionales en un

producto químico puro se han podido correlacionar con bandas específicas en el EM (Figura 2).

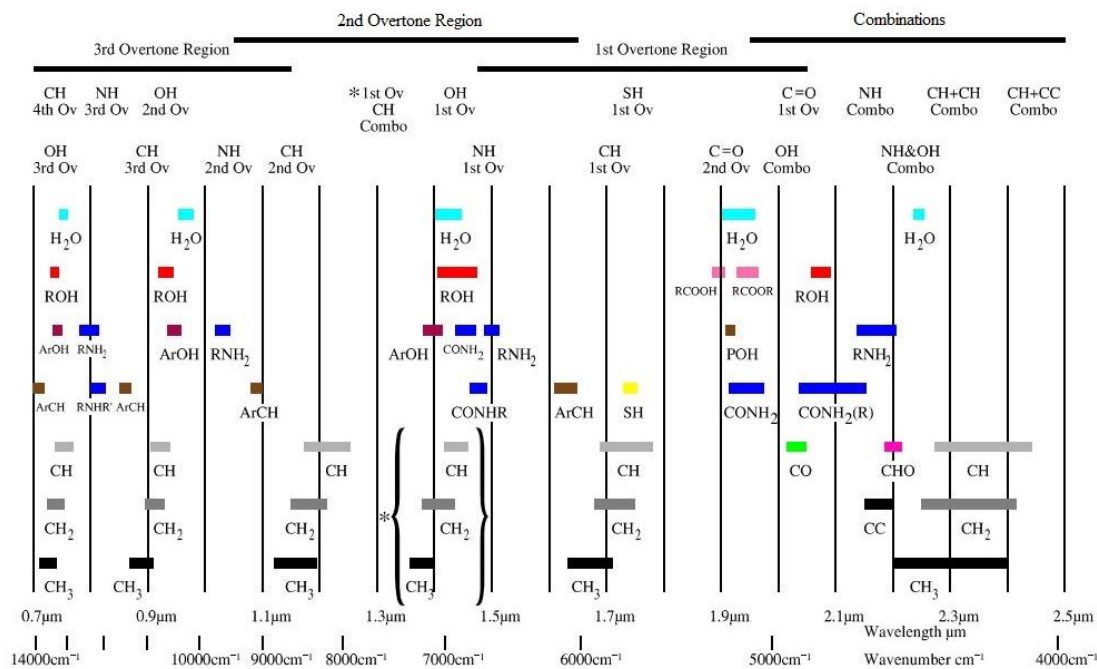


Figura 2. Bandas de absorción en el infrarrojo cercano. Recuperado de Murray (2004).

La absorción en NIR de un compuesto orgánico tiene muchas bandas superpuestas y alta multicolinealidad, por lo que asignarlas a un enlace es difícil. Además, las combinaciones de segundo y tercer orden aparecen en longitudes de onda más cortas, y el tercer sobretono del estiramiento de C-H se puede encontrar hasta la región visible (Ozaki & Morisawa, 2021). Por ello, es conveniente el uso de espectros Vis-NIR en análisis de compuestos orgánicos.

2.2 Aplicaciones de la espectroscopía en la medición de calidad en frutas

2.2.1 Instrumentación

La luz incidente (I_0) sobre la materia puede ser reflejada (I_r), transmitida (I_t) o absorbida (I_a). Existen tres modos principales de detección de cualquiera de estas interacciones: transmitancia, reflectancia y transreflectancia (Figura 3). El modo reflectancia que se debe principalmente a la reflectancia difusa es el que

se ha utilizado en los espectrómetros comerciales, debido a la simplicidad de la posición de la muestra (Vidal, 2019).

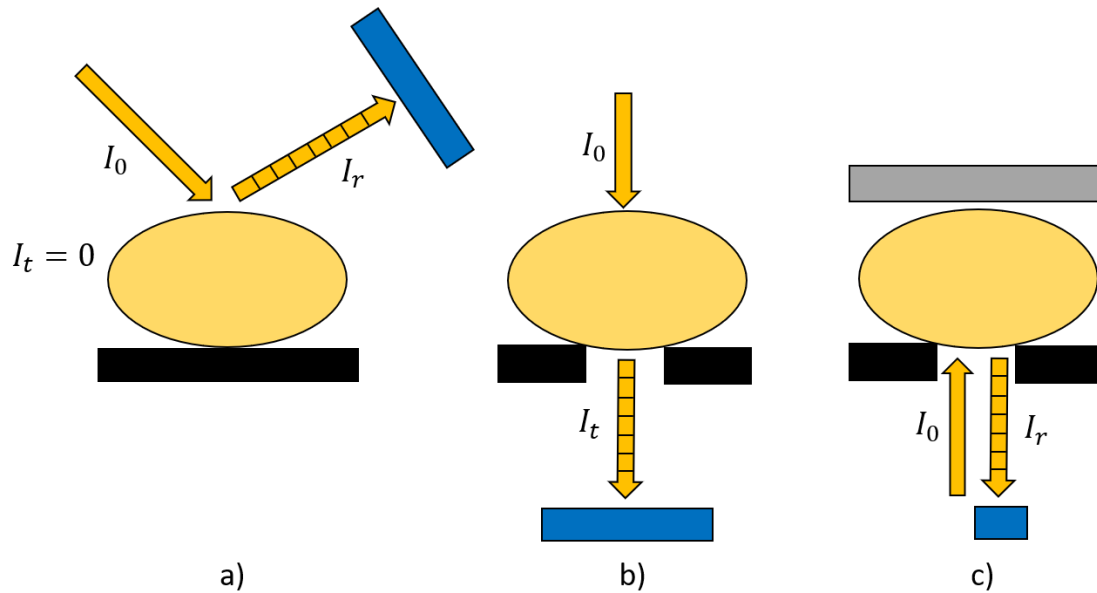


Figura 3. Modos de detección espectral: a) reflectancia, b) transmitancia y c) transfectancia. Azul: detector, negro: base, gris: reflector. Modificado de Xie, Wang, Xu, Fu, & Ying (2016)

En transmisión con muestras transparentes y una trayectoria de la luz monocromática generalmente de 1 cm, la absorbancia se puede obtener por la ley de Beer-Lambert, mientras que en reflexión se puede obtener como (Murray, 2004):

$$\text{Absorbancia } (A) = \text{Log } 1/R = \text{Log } (I_0/I_r)$$

La A es adimensional y tiene una escala logarítmica por lo que si $A=1$ significa que se registra el 10% de I_0 , mientras que una $A=2$, significa la detección del 1%.

El tiempo de obtención de los espectros es poco y no requiere ningún reactivo, los equipos se han simplificado y abaratado (Vidal, 2019), por lo que la espectroscopía NIR y visible tiene una gran versatilidad, pudiendo aplicarse tanto en laboratorios como en campo.

Por lo general un sistema NIR se compone de cinco componentes principales: una fuente de luz, un monocromador, un detector, un compartimento para la muestra y algunos accesorios ópticos, que varían ligeramente dependiendo del

propósito, por ejemplo, para dispositivos portátiles, a menudo se integran los cinco componentes en un solo módulo (Xie et al., 2016).

La fuente de luz provee los fotones para irradiar la muestra. Normalmente, la fuente de radiación usada debe cubrir regiones del espectro Vis y NIR. Las principales son las lámparas halógenas de tungsteno, los diodos de emisión de luz y los diodos láser, de las cuales las lámparas halógenas de tungsteno tienen preferencia por el amplio espectro que cubren, el bajo costo y la fácil regulación en la práctica.

El monocromador es el elemento dispersivo de la luz en una amplia gama de longitudes de onda. De acuerdo a este elemento el espectrómetro puede ser de filtro óptico, de rejilla, de transformada de Fourier o espectrómetro de filtro sintonizable electrónicamente (ETF). En estos últimos, predominan los filtros sintonizables acústico-ópticos (AOTF) y los filtros sintonizables de cristal líquido (LCTF). Los AOTF tienen características ventajosas como la ausencia de partes móviles, rapidez y amplia banda y rango espectral, sin embargo, su precio es mayor que los LCTF por lo que raramente se usan en productos agrícolas.

El detector se encarga de registrar la intensidad de luz por medio de señales eléctricas. Actualmente se utilizan semiconductores hechos a base de silicio (Si), arseniuro de indio y galio (InGaAs) y de sulfuro de plomo (PbS), de los cuales los dos primeros son lo más utilizados. Diferentes detectores tienen diferente sensibilidad, los de Si pueden cubrir de 700-1100 nm, los InGaAs de 800-1700 nm, los de PbS de 1000-2500 nm, mientras que para la región visible se puede utilizar un tubo fotomultiplicador o un detector CMOS.

El compartimento para la muestra se adapta al tipo de muestras, sólidas, líquidas o gaseosas y al modo de detección. La mayoría de dispositivos portátiles también incluyen haces de fibra óptica para conducir la energía eléctrica que genera la radiación utilizada.

2.2.2 Evaluación de calidad en frutas

La calidad de las frutas es un término definido por los humanos que engloba muchas propiedades o características, principalmente, las relacionadas con la apariencia, la textura, el sabor y el valor nutricional. La apariencia se juzga por el tamaño, la forma, brillo o ausencia de defectos, mientras que la textura incluye la firmeza, frescura, jugosidad, harinosidad o dureza, asimismo, el sabor incluye la dulzura, acidez, astringencia o el aroma, y la calidad nutricional involucra el aporte de vitaminas, minerales o fibra dietética a la nutrición humana (Brasil & Siddiqui, 2018).

El consumidor hace una valoración de la calidad con ayuda de sus sentidos, especialmente de la vista, el olfato, el gusto, el tacto e incluso el oído, por lo que el control de la calidad está orientado a la aceptabilidad de los consumidores. Para controlar la calidad se deben medir los atributos relacionados a ella y los instrumentos de medición que fueron basados de estos sentidos, por ejemplo, la textura se detecta midiendo las propiedades ópticas, la consistencia por las propiedades mecánicas y el sabor por las propiedades químicas, proporcionan mediciones más precisas y pueden formalizar un lenguaje común entre los investigadores, la industria y los consumidores (Abbott, 1999). La evaluación de la calidad interna presenta mayores dificultades porque tradicionalmente involucra un muestreo destructivo, sin embargo, se han desarrollado tecnologías no invasivas, dentro de las cuales la espectroscopía Vis-NIR es de las que más adopción ha tenido en la práctica comercial (Walsh, Mcglone, et al., 2020). Se busca optimizar la instrumentación y los métodos para relacionar los atributos de calidad con las mediciones espectrales.

Del 2015 a inicios de 2020 se ha reportado el uso de la espectroscopía Vis-NIR en la medición de atributos de calidad como el contenido de humedad, la pigmentación, niveles de reserva de almacenamiento (azúcar soluble, almidón o aceite), ácidos orgánicos, varios defectos internos, índices de madurez y afirmaciones para la medición de la firmeza (Walsh, Blasco, Zude-Sasse, & Sun, 2020). Los estudios científicos han logrado caracterizar bandas de absorción

típicas a grupos químicos relacionados a diversos atributos de calidad, como se muestra en el Cuadro 1 (Arruda de Brito et al., 2022; Goke, Serra, & Musacchi, 2018; J. Li et al., 2019; Najjar & Abu-Khalaf, 2021; Subedi & Walsh, 2020; H. Wang, Peng, Xie, Bao, & He, 2015).

Cuadro 1. Bandas de absorción típicas a grupos químicos relacionados a atributos de calidad interna

Atributo de calidad	Grupo químico	Longitud de onda (nm)
Azúcar	O-H	1190, 1400
CSS	C-H	910, 963
	O-H	960, 975, 1180, 1450, 2000
	C-H y O-H	950-1075, 1210
	Bandas de combinación de C-H y O-H	2000-2500
	O-H	
	O-H, C-H, N-H	550-1100
Azúcar	C-H y O-H	650-950
Acidez	C-H de COOH	1607
	O-H de ácidos carboxílicos	1127
	C=O de ácido carboxílico	1437
	saturado e insaturado	
	O-H, C-H, N-H	550-1100
	-	675
pH	C-H	768
	O-H	986
CMS o CA	H-O-H	729-975
	O-H	740, 840, 960, 1140
	-	840-960

El CMS representa la suma de carbohidratos solubles (azúcar) e insolubles (almidón), proteínas, minerales y otros compuestos acumulados durante su desarrollo en el árbol y se distingue de otros índices de madurez o calidad, como el color, la firmeza o el CSS porque representa el almidón que se metaboliza en azúcares solubles durante la maduración. Medirlo antes o en el momento de cosecha se puede utilizar para predecir la futura calidad interna después de

largos periodos de almacenamiento, como se ha demostrado en manzanas, mangos y kiwis (Goke et al., 2018).

El CSS se ha podido determinar en frutas intactas como manzana, cítricos, kiwi, mango, melón, cebolla, durazno, papa y tomate, con tecnología Vis-NIR desde antes de 1998, con una correlación de aproximadamente 0.93 (Abbott, 1999). Posteriormente se hicieron contribuciones con la misma tecnología para la evaluación y autenticación de la calidad de nuevas frutas como arándanos, albaricoque, mandarinas, uvas, harina de soja (Srivastava & Sadistap, 2018), peras (Goke et al., 2018), entre otras. Actualmente, el uso de Vis-NIR en la determinación de CMS y CSS en frutas de piel fina está establecido en la práctica comercial (Walsh, Blasco, et al., 2020).

2.3 Modelación de datos espectrales

Para desarrollar un modelo de predicción es necesario escanear muestras con una buena distribución del o los atributos de calidad de interés. El espectrómetro nos permite obtener los datos ópticos en valores de absorbancia (variable X), mientras que por métodos tradicionales de laboratorio podemos obtener los valores de referencia (variable Y), posteriormente estas variables se pueden relacionar usando la quimiometría. Una división de datos en conjuntos de calibración y prueba nos permite desarrollar la ecuación de calibración y comprobar el rendimiento de dicha ecuación, respectivamente. Normalmente los espectros pueden contener cambios en la línea de base, picos superpuestos y otros tipos de ruido que será necesario eliminar con un tratamiento previo de los datos (preprocesamiento) (Sirisomboon, 2018).

2.3.1 Preprocesamiento espectral

La presencia de variaciones no esperadas se puede atribuir a causas como el modo de medición, las derivas instrumentales, el estado de la muestra, entre otros factores físicos, químicos y ambientales externos (Roger, Boulet, Zeaiter, & Rutledge, 2020). Se han desarrollado varios métodos matemáticos para corregir

dichos efectos, el Cuadro 2 presenta los principales métodos y sus efectos correctivos.

Cuadro 2. Principales métodos de preprocesamiento espectral. Recuperado de Sirisomboon (2018).

Preprocesamiento	Efecto
Suavizado	Suprimir el ruido
Primera y segunda derivada	Separación de picos superpuestos y de la línea de base
Normalización	Corrección de variabilidades no deseadas
Eliminación de tendencias (DT)	Eliminar la curvilinealidad y los cambios de línea de base
Transformación por variable normal estándar (SNV)	Eliminar los efectos de dispersión
SNV + DT	Reducir la multicolinealidad, el cambio de línea de base y la curvilinealidad
Corrección de la dispersión multiplicativa (MSC)	Corregir los efectos multiplicativos y aditivos

Otro tipo de preprocesamiento común es la reducción del espacio de variables Y . Se suelen utilizar procedimientos para extraer variables descriptivas (longitudes de onda) características, con mayor capacidad predictiva, eliminando información redundante y reduciendo el tiempo de calibración (H. Wang et al., 2015), o transformar las variables a un nuevo espacio que describa al máximo las variaciones relacionadas o no con el atributo de interés, a través de proyecciones ortogonales (Roger et al., 2020).

No existe una regla base para elegir el tipo de preprocesamiento, por lo que un usuario podría requerir explorar todas las opciones disponibles y aún más, la combinación de múltiples preprocesamientos puede corregir de mejor manera los espectros y mejorar el rendimiento de los modelos (Mishra, Biancolillo, Roger, Marini, & Rutledge, 2020).

2.3.2 Métodos de calibración

El análisis con espectroscopía puntual, a diferencia de las imágenes hiperespectrales o multiespectrales se ha producido con métodos estadísticos.

Inicialmente, se optó por la MLR, sin embargo, en las últimas tres décadas se ha preferido la PLSR, mientras que la búsqueda de solidez en los modelos se ha orientado por lo métodos de aprendizaje automático como las ANN y el DL (Anderson & Walsh, 2022), pudiendo diferenciar dos métodos: los lineales y no lineales.

Cuando se trabaja con grandes conjuntos de datos, del conjunto de calibración se deriva un subconjunto de validación, usado para optimizar factores en el modelo quimiométrico, como el número de componentes principales (PC) en regresión por PC, variables latentes en PLSR, parámetros del kernel en máquinas de vectores de soporte (SVM) (Morais, Santos, & Martin, 2019) y los hiperparámetros en ANN. Cuando no existe este subconjunto, se utiliza la validación cruzada que suele demandar mucho más tiempo, pero permite obtener modelos más robustos.

Regresión lineal

Estos métodos asumen que la relación entre variables X e Y , es lineal. Los más usados son MLR, PCR y PLSR (Saeys, Nguyen Do Trong, Van-Beers, & Nicolai, 2019). La **MLR** asume una función lineal para cada variable predicha Y . El conjunto de coeficientes de regresión $\mathbf{b}$ que minimiza la suma de los errores cuadráticos (SSE) en el conjunto de calibración se determina por:

$$\mathbf{b} = (X_c^T X_c)^{-1} X_c^T y_c$$

Donde X_c y y_c son las matrices y vector de calibración centrados en la media, respectivamente. De esta manera para una nueva muestra, el atributo de calidad se puede obtener como:

$$y_{nueva} = \bar{y} + X_{nueva} * \mathbf{b}$$

El número de muestras n debe ser mayor que el número de variables predictoras k para evitar el problema de la multicolinealidad exacta.

La **PCR** reemplaza las variables originales en X_c por nuevas variables no correlacionadas. Así la matriz X_c se convierte en una matriz T con las puntuaciones de los primeros componentes principales:

$$\mathbf{c} = (T^T T)^{-1} T^T y_c$$

Donde c es un vector con los coeficientes de regresión en el espacio de puntuación de PC. Así, para evaluar una nueva muestra se necesita la nueva matriz de puntuaciones T_{nueva} y multiplicarla por c :

$$y_{new} = \bar{y} + T_{nueva} * c$$

De esta manera se suelen obtener errores más pequeños que MLR y se puede maximizar el rendimiento del modelo incluyendo los suficientes PC que representen la mayor variabilidad de los datos. Para lo cual, se itera con diferente número de PC vigilando el rendimiento en el subconjunto de validación.

Dado que para obtener las PC no se utiliza la variable dependiente, podría haber el caso en el que no sean lo suficientemente informativas para predecir el atributo de calidad. Por ello, el método **PLSR** define las nuevas variables como combinaciones lineales no correlacionadas de las variables originales (variables latentes o LV) que capturan la máxima covarianza entre X e Y. Las LV se obtienen con una descomposición en valor singular de productos cruzados $X_C^T * Y_C$, resultando dos matrices por cada una, de cargas y puntajes. En este caso, se pueden manejar múltiples variables Y al mismo tiempo y definir las LV que capturen la máxima covarianza entre X y todas las variables Y, que a menudo se denomina PLS2. Sin embargo, PLS2 es más recomendado cuando las variables Y están altamente correlacionadas entre sí.

El conjunto de validación igualmente se utiliza para determinar el número óptimo de variables latentes y se espera que PLSR necesite menos LV que PCR CP.

Regresión no lineal

Los métodos anteriores pueden manejar un cierto grado de no linealidad entre las variables X e Y, agregando más CP o LV por ejemplo, pero en no linealidades fuertes no dan un modelo satisfactorio (Gemperline, Long, & Gregoriou, 1991). Se podría forzar una relación lineal aplicando una transformación no lineal en las variables dependientes, sin embargo, es más conveniente usar regresión no lineal, para lo cual se han usado dos métodos: métodos locales y de kernel (Ni, Nørgaard, & Mørup, 2014; Saeys et al., 2019).

Los *métodos locales* suponen que la predicción óptima es constante localmente. La mayoría de métodos desarrolla un modelo específico para cada nueva muestra usando un subconjunto de calibración seleccionado de un conjunto más grande suficiente similar a la nueva muestra que se va a predecir. Normalmente no se obtiene un modelo cerrado y se necesita disponer del conjunto de entrenamiento para hacer predicciones, por lo que también se denominan métodos basados en memoria.

A grandes rasgos se pueden citar los métodos de regresión PLS local (**LPLS**), la regresión ponderada localmente (**LWR**) y el algoritmo central local (**LCA**). Sus diferencias son análogas a su correspondiente método global, pero en todos se suelen elegir las muestras de entrenamiento en función de la distancia de Mohalanobis o euclidiana entre espectros de calibración y de predicción para desarrollar el modelo local.

Por su parte los *métodos de kernel* evitan la comparación de la muestra a predecir con todas las demás muestras al aplicar una transformación de la variable original a un espacio de características con relación lineal a la variable Y. Los kernel más comunes son el kernel lineal y de función de base radial (**RBF**). Sobre esta matriz de características los coeficientes de regresión deben obtenerse con un método de proyección como **Kernel PCR**, **Kernel PLSR** o mediante regularización (regresión **Ridge del Kernel**). Este último equivale a la máquina de vectores de soporte de mínimos cuadrados (LS-SVM).

2.4 Redes neuronales artificiales para análisis espectral

2.4.1 Redes neuronales artificiales

Las ANN son un método atractivo para el análisis espectral por su habilidad para reconocer patrones y manejar las relaciones lineales o no lineales de manera automática. Han sido utilizadas en las investigaciones desde los 1990s y comparadas con los métodos tradicionales (Abbott, 1999). Se descubrió que pueden modelar relaciones lineales con rendimiento similar a PCR o PLSR y relaciones no lineales con mejor rendimiento (Gemperline et al., 1991).

Típicamente una ANN de alimentación hacia adelante (feed-forward ANN) se forma de tres capas: de entrada, oculta y de salida. Para el análisis espectral una ANN es expresada matemáticamente como (Ni et al., 2014):

$$y_i = f \left[\sum_{t=1}^T w_t g \left(\sum_{j=1}^J w_{tj} x_{ij} + \phi_t \right) + \varepsilon \right] + e_i$$

Donde T es la cantidad de neuronas en la capa oculta, w_{tj} indica los pesos entre capa de entrada-oculta, w_t expresa el peso de las neuronas entre capa oculta-salida, ϕ_t y ε representan los sesgos en capas oculta y salida, respectivamente y e_i es el ruido. Las funciones de transferencia f y g denotan funciones lineales y sigmoideas, respectivamente. Para ajustar los parámetros de la red (pesos) y minimizar el error de salida se usa el algoritmo de propagación hacia atrás. Las ANN pueden tener rendimientos similares con entradas de datos originales, transformados, reducidos, preprocesados y no preprocesados.

El estudio exhaustivo de Ni et al. (2014) donde se compararon los principales métodos no lineales, entre ellos ANN, demostró que generalmente las ANN no son el método más adecuado para la modelación espectral, debido a su tendencia al sobreajuste, especialmente cuando se entrenan con conjuntos de datos pequeños o cuando se aplican en problemas lineales, por lo que a menudo es criticada la compleja estructura del modelo. Los métodos bayesianos pueden enfrentar el ajuste excesivo que combinados con ANN puede producir modelos más sólidos y precisos.

Sin embargo, estudios más actuales y más especializados con ANN como el de Basile, Marsico, & Perniola (2022) demuestran que con la configuración óptima, las ANN pueden producir modelos cualitativos y cuantitativos más precisos que PLSR, por lo que la configuración de los parámetros, que afectan el rendimiento final del modelo, es un límite para la aplicación de las ANN en el análisis espectral. Para que esta técnica tenga una aceptación generalizada se deben desarrollar métodos automáticos de configuración óptima.

2.4.2 Aprendizaje profundo para el análisis espectral

Funcionamiento

Los datos espectrales contienen características informativas y no informativas, los cambios en las no informativas ocurren bajo muchas condiciones variables, como diferentes estructuras físicas, espectrómetros de diferentes fabricantes, o diferentes condiciones ambientales, lo que limita la generalización de los enfoques quimiométricos convencionales (X. Zhang et al., 2021). Los modelos de aprendizaje profundo tienen la capacidad para manejar características lineales y no lineales de los datos espectrales sin procesar y, por tanto, sin necesidad de conocimientos humanos previos. En consecuencia, este método de modelado de extremo a extremo ha presentado mejores rendimientos de generalización que los métodos quimiométricos convencionales.

Existen varios métodos de modelado profundo, como las redes neuronales residuales (ResNet), las redes neuronales de codificador automático profundo (DAE), las redes neuronales recurrentes (RNN) y las redes neuronales convolucionales (CNN), de las cuales las CNN se encuentran a la vanguardia en el análisis espectral unidimensional (X. Zhang et al., 2021).

A partir del 2017 se despertó el interés por las CNN en el análisis espectral con la adaptación de las CNN usadas con imágenes a los datos unidimensionales y comprobado un rendimiento superior al de PLSR (Acquarelli et al., 2017). Son redes neuronales de propagación hacia atrás que incluyen al menos una capa convolucional y una capa totalmente conectada. La capa convolucional convoluciona los espectros de entrada con filtros convolucionales que se deslizan sobre los datos, adoptando el producto punto entre la entrada y el filtro, de esta manera los pesos de la capa se asignan al filtro, que son compartidos y se recolectan localmente (Figura 4). Debido a esta conexión local y los pesos compartidos, las CNN tienen menos parámetros entrenables que las redes completamente conectadas, reduciendo el riesgo de sobreajuste (X. Zhang, Lin, Xu, Luo, & Ying, 2019).

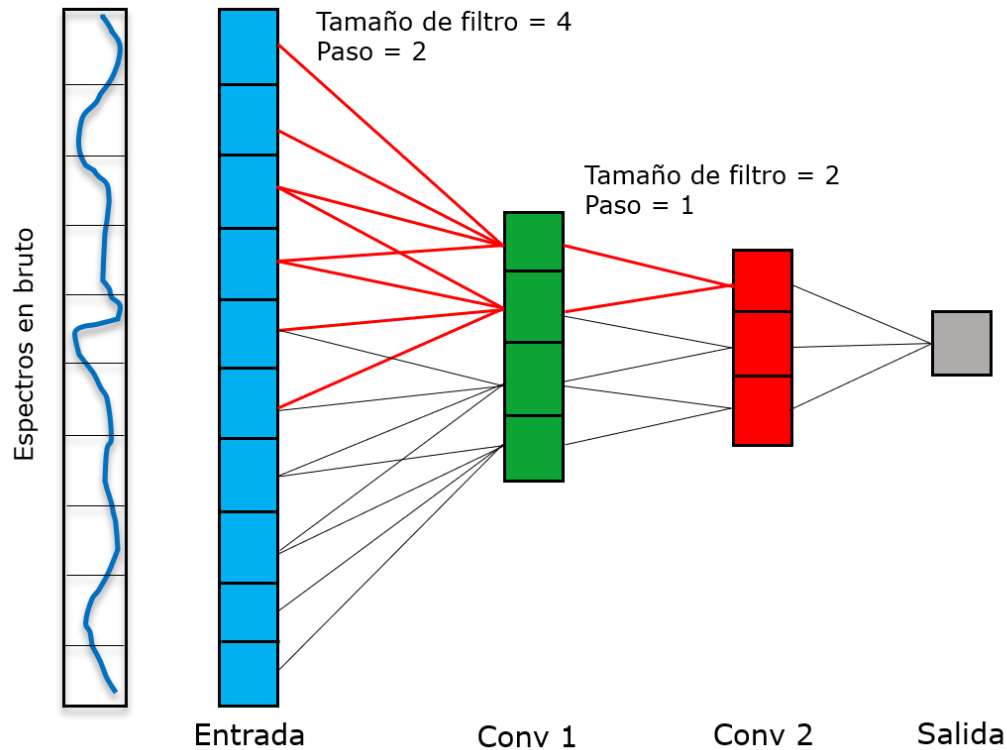


Figura 4. Arquitectura de una 1D-CNN para datos espectrales. Una neurona en la capa Conv2 cubre un campo receptivo de seis variables en el espectro original. Una vez que se activa la neurona el patrón del campo receptivo asociado se aprende como característica. Modificado de X. Zhang et al. (2019, 2021).

El diseño de una arquitectura de CNN se basa en los datos de entrada, número de muestras y cantidad de longitudes de onda principalmente, que determinan el número de capas, la cantidad de filtros y los pasos. Se ha observado una tendencia a usar más capas y más filtros cuando se dispone de más muestras, pudiendo contener de 1 hasta 10 capas, con más de 10 mil muestras para el caso extremo, mientras que el tamaño de los filtros van desde 3 hasta 100, aunque algunos investigadores prefieren usar los más pequeños para disminuir la cantidad de parámetros entrenables (X. Zhang et al., 2021).

Hiperparámetros

La mayoría de los modelos de aprendizaje profundo tienen varios hiperparámetros que controlan el comportamiento del algoritmo y que tiene un efecto en el tiempo, costo de memoria y calidad del modelo cuando se implementa en nuevas muestras (GoodFellow et al., 2016). El Cuadro 3 presenta

algunos hiperparámetros que se pueden configurar razonando su efecto en la capacidad del modelo.

Cuadro 3. Efecto de algunos hiperparámetros en la capacidad del modelo. Recuperado de GoodFellow et al. (2016).

Hiper-parámetro	Mayor capacidad cuando:	Razón	Advertencias
Número de unidades ocultas	Aumenta	Aumentar el número de unidades ocultas incrementa la capacidad de representación del modelo	Esto aumenta tanto el tiempo como el costo de la memoria de prácticamente todas las operaciones en el modelo
Tasa de aprendizaje	Ajustada de manera óptima	Una tasa de aprendizaje demasiado alta o demasiado baja genera un modelo con capacidad efectiva baja debido a una falla en la optimización	
Tamaño del filtro	Aumenta	Aumentar el ancho del filtro aumenta el número de parámetros del modelo	Un filtro más grande genera una dimensión de salida más estrecha, lo que reduce la capacidad del modelo a menos que se use un relleno de ceros. Asimismo, se requiere más memoria de almacenamiento y tiempo de ejecución, pero una salida más estrecha reduce el costo de memoria
Relleno cero implícito	Aumenta	Agregar ceros implícitos antes de la convolución mantiene el tamaño de representación grande	Aumenta el tiempo y el costo de la memoria de la mayoría de las operaciones
Coefficiente de caída de pesos	Disminuye	Este coeficiente libera los parámetros del modelo para que sean más grandes	
Tasa de abandono	Disminuye	Apagar unidades con menos frecuencia les da	

a las unidades más
oportunidades de
“conspirar” entre sí para
adaptarse al conjunto de
entrenamiento

El modelado con DL ya ha superado los métodos tradicionales de quimiometría y los enfoques clásicos de aprendizaje automático (Passos & Mishra, 2021). Se está explotando su uso en el análisis espectral cualitativo, por ejemplo para la identificación de variedades, detección de origen geográfico, reconocimiento de adulteración y detección de moretones, así como en el análisis cuantitativo de frutas, granos y cultivos (X. Zhang et al., 2021). Sin embargo, las ganancias obtenidas por esta simplificación del análisis con DL pueden opacarse debido al alto número de hiperparámetros que se requieren configurar en los modelos (Passos & Mishra, 2021).

Optimización de hiperparámetros

Básicamente los hiperparámetros en una CNN se pueden escoger de manera manual o automática. La configuración manual requiere de conocimiento sobre el efecto de cada uno en el rendimiento del modelo en términos de error de entrenamiento, capacidad de generalización y los recursos computacionales. Si bien las CNN pueden funcionar bien con unos pocos hiperparámetros óptimos, a menudo se benefician más de una optimización más completa. Los métodos automáticos pueden elegir sus hiperparámetros en base a una función objetivo (por ejemplo el error de validación) y sus propios hiperparámetros con más fáciles de elegir, por ejemplo, el rango de valores a explorar (GoodFellow et al., 2016).

La búsqueda de cuadrícula, búsqueda aleatoria y los métodos bayesianos son los principales métodos de optimización automática de hiperparámetros (HPO) (Yu & Zhu, 2020). La búsqueda de cuadrícula es una búsqueda en paralelo donde cada ensayo se ejecuta individualmente sin la influencia de la secuencia de tiempo. Se necesita un conocimiento previo sobre los hiperparámetros candidatos, además, su costo computacional crece exponencialmente con el número de hiperparámetros. La búsqueda aleatoria es una mejora básica a la

anterior, su patrón de búsqueda la hace comparativamente más probable de que encuentre la configuración óptima que la búsqueda de cuadrícula, porque puede extenderse hasta un presupuesto predeterminado o hasta que alcanza la precisión deseada. Esta búsqueda aleatoria se suele usar como línea de base para otros HPO. Los métodos bayesianos usan modelos de regresión bayesiana para estimar tanto el valor esperado en la función objetivo como la incertidumbre alrededor de esta expectativa, por lo que se puede determinar el siguiente punto de muestreo haciendo más eficiente el gasto computacional, requiere de menos ensayos y no requiere de conocimientos previos sobre la distribución de hiperparámetros. La Figura 5 ilustra gráficamente la diferencia entre cada método.

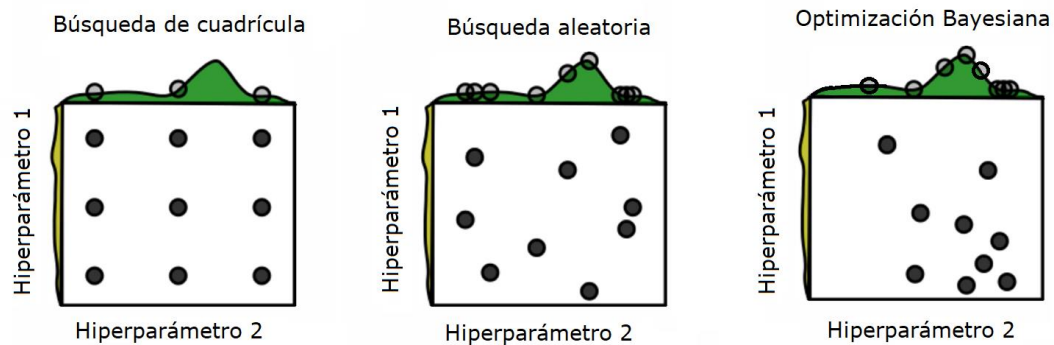


Figura 5. Comparación de tres métodos de ajuste automático de hiperparámetros. Las regiones coloreadas representan la pérdida. Modificado de Yu & Zhu (2020).

Se ha utilizado la optimización bayesiana con otras técnicas como hiperbanda y se ha comprobado que mejora el rendimiento de modelos con arquitectura e hiperparámetros elegidos manualmente, superando además, los mejores resultados conocidos e incluso conduciendo a modelos más simples, todo bajo un costo computacional asequible (Passos & Mishra, 2021).

2.5 Estado del problema

Para garantizar una consistencia en el uso de la técnica no destructiva Vis-NIR se requiere atención e investigación adicional. En primer lugar se requiere que todos los investigadores reporten sus resultados de manera uniforme, en cuanto al rendimiento de los modelos a menudo se usa el error cuadrático medio de

calibración, validación y prueba (RMSEC, RMSEV, RMSEP, respectivamente), así como la relación entre valores predichos y reales (R^2) (Walsh, Mcglone, et al., 2020), sin embargo, es más confiable usar el coeficiente R^2 en lugar del RMSE para comparar resultados entre investigaciones, ya que tiene unidades adimensionales. En segundo lugar, la robustez de los modelos debe abordarse a través de la modelación y validación multi-temporada, multi-cultivar y multi-origen (Cavaco, Passos, Pires, Antunes, & Guerra, 2021), principalmente. Estos modelos globales hacen frente a la intrigante calibración de modelos locales debido a las múltiples temporadas de cosecha, cultivares y huertos de origen de las frutas.

Los estudios de modelos globales multi-cultivar son los que más se encuentran en la literatura, con frutas como manzana, pera, melón y tomate (Kusumiyati, Hadiwijaya, Sutari, & Munawar, 2022; Najjar & Abu-Khalaf, 2021; J. Wang, Wang, Chen, & Han, 2017; Y. Zhang, Nock, Shoffe, & Watkins, 2019). Otros estudios usan muestras con varias fuentes de variabilidad. Teh, Coggins, Kostick, & Evans (2020) desarrollaron un modelo global multi-anual, multi-origen y multi-cultivar para predicción de CMS en manzanas y descubrieron que el modelado global lineal PLSR tiene altas precisiones de predicción en un conjunto externo, sin embargo, la modelación local por cultivar aumentó la precisión de la predicción. Anderson, Walsh, Subedi, & Hayes (2020) desarrollaron un modelo global similar al anterior, pero con mango, y fue validado con muestras de una temporada externa. El CMS fue predicho globalmente en mangos de una cuarta temporada con RMSEP de 1.01% mientras que modelos locales disminuyeron el error a 0.88%, además una ANN obtuvo un RMSEP global de 0.89%, lo que indica que gran variabilidad en muestras globales agrega no linealidad a la relación entre variables. Posteriormente dichos datos se modelaron con regresión PLS local y otros modelos no lineales (Anderson, Walsh, Flynn, & Walsh, 2021) y se encontró un modelo global que combinó ANN, regresión de proceso gaussiano (GPR) y LPLS-S, el cual fue más preciso que modelos locales (RMSEP de 0.839% y 0.881%, respectivamente). La misma base de datos fue retomada por Mishra & Passos (2021a) y por primera vez se modeló con la técnica de DL. Para esto no se

investigó un preprocesamiento específico, sino que se combinaron SNV y filtrado Savitzky-Golay para aumentar la cantidad de variables predictoras, logrando un modelo global con el mejor rendimiento de validación externa conocido (RMSEP=0.79%).

La diferencia entre rendimientos de modelos globales y locales también se observa en otros estudios. Por ejemplo, un estudio con cinco cultivares de pera y métodos quimiométricos clásicos, generó un modelo de predicción de CSS con un $R^2 = 0.85$ y RMSEP = 0.46 (J. Wang et al., 2017), mientras un estudio con una sola variedad de pera (que igualmente se encontraba en el estudio anterior) logró un RMSEP de 0.54 con un modelo de DL de salida única así como un RMSEP=0.56 y $R^2=81$ con un modelo DL de salida múltiple (Mishra & Passos, 2022).

La predicción multi-objetivo (MT) es otro enfoque de investigación para modelos globales. La motivación se debe a dos razones importantes, la primera es que permite modelar de manera efectiva los conjuntos de datos al considerar también las relaciones entre los objetivos, garantizando así mayor interpretabilidad real del modelo, mientras que la segunda razón es que se pueden desarrollar modelos más simples y con una mejor eficiencia computacional (Borchani, Varando, Bielza, & Larrañaga, 2015). Ante esto, Barbon-Junior et al. (2020) compararon la predicción MT con la predicción de un solo objetivo (ST), para la predicción de seis parámetros de calidad de trigo, usando el enfoque de transformación de problemas (o métodos locales), los cuales transforman el problema MT en problemas independientes ST, y finalmente concatenan las predicciones. La regresión se realizó con SVM, bosque aleatorio (RF) y regresión lineal (LR), encontrando una mejora del 7% con MT en comparación con ST. Por otro lado, el problema también se ha abordado a través de la adaptación de algoritmos (o métodos globales), que se basan en la adaptación del método ST para manejar directamente las múltiples salidas. Yang, Sun, Pan, Huang, & Zhu (2023) propusieron un algoritmo genético multitarea (EM4GPO) y lo compararon con métodos ST, entre ellos de DL, así como con otros métodos MT, para la

predicción de CCS y firmeza en manzana, obtenido mejores resultados con el método propuesto. Si embargo no se comparó con un método MT basado en DL, aunque poco antes ya se había demostrado que las redes CNN se podían adaptar a salidas múltiples para la predicción de CSS y CMS en peras y, aunque no se encontró mejora significativa en la precisión con la adaptación MT, sí fue más eficiente computacionalmente, en comparación con la predicción ST (Mishra & Passos, 2022).

El análisis cuantitativo global también ha seguido el enfoque multi-producto, que involucra el desarrollo de modelos para diferentes tipos de productos con gran variabilidad. Rambo, Ferreira, & Amorim (2015) desarrollaron un modelo de predicción de composición química de biomásas de plátano, café y coco con PLSR. Aunque el rendimiento del modelo global puede disminuir en comparación con modelos locales, el modelo multi-producto fue robusto y evitó los costos elevados de la recolección de muestras de productos específicos. También se ha predicho la AT y vitamina C de nueve tipos de jugo con un solo modelo para cada objetivo, usando PLSR (Dos Santos et al., 2019), lo que permitió ahorrar tiempo, mejoró la robustez y la practicidad. Así mismo, se ha investigado la modelación de calidad en diversos frutos de una misma familia vegetal. Kusumiyati, Hadiwijaya, Putri, & Munawar (2021a) desarrollaron dos modelos multi-producto para predecir contenido de agua y CSS en seis cucurbitáceas con PLSR, descubriendo que la variabilidad de múltiples productos mejoró el rendimiento de predicción en comparación con estudios previos de un solo producto.

Si bien se han demostrado las ventajas que presentan los modelos globales, aún falta investigar modelos que combinen todas las fuentes de variabilidad descritas, por ejemplo, modelos multicultivar-multisalida, multitemporada-multisalida, multiorigen-multisalida, multiproducto-multisalida, e incluso, un modelo que combinen todos estos aspectos, lo cual aumentará la complejidad de las relaciones entre variables y se necesitarán de los métodos más novedosos de regresión no lineal. A la mejor búsqueda, hasta la fecha no se ha realizado una investigación de este tipo.

2.6 Literatura citada

- Abbott, J. A. (1999). Quality measurement of fruits and vegetables. *Postharvest Biology and Technology*, 15(3), 207–225. [https://doi.org/10.1016/S0925-5214\(98\)00086-6](https://doi.org/10.1016/S0925-5214(98)00086-6)
- Acquarelli, J., Laarhoven, T. Van, Gerretzen, J., Tran, T. N., Buydens, L. M. ., & Marchiori, E. (2017). Convolutional neural networks for vibrational spectroscopic data analysis. *Analytica Chimica Acta*, 954, 22–31. <https://doi.org/10.1016/j.aca.2016.12.010>
- Anderson, N. T., & Walsh, K. B. (2022). Review: The evolution of chemometrics coupled with near infrared spectroscopy for fruit quality evaluation. *Journal of Near Infrared Spectroscopy*, 30(1). <https://doi.org/https://doi.org/10.1177/09670335211057235>
- Anderson, N. T., Walsh, K. B., Flynn, J. R., & Walsh, J. P. (2021). Achieving robustness across season , location and cultivar for a NIRS model for intact mango fruit dry matter content. II. Local PLS and nonlinear models. *Postharvest Biology and Technology*, 171, 111358. <https://doi.org/10.1016/j.postharvbio.2020.111358>
- Anderson, N. T., Walsh, K. B., Subedi, P. P., & Hayes, C. H. (2020). Achieving robustness across season, location and cultivar for a NIRS model for intact mango fruit dry matter content. *Postharvest Biology and Technology*, 168, 111202. <https://doi.org/10.1016/j.postharvbio.2020.111202>
- Arendse, E., Fawole, O. A., Magwaza, L. S., & Opara, U. L. (2017). Non-destructive prediction of internal and external quality attributes of fruit with thick rind: A review. *Journal of Food Engineering*, 217, 11–23. <https://doi.org/10.1016/j.jfoodeng.2017.08.009>
- Arruda de Brito, A., Campos, F., dos Reis Nascimento, A., Damiani, C., Alves da Silva, F., de Almeida Teixeira, G. H., & Cunha Júnior, L. C. (2022). Non-destructive determination of color, titratable acidity, and dry matter in intact tomatoes using a portable Vis-NIR spectrometer. *Journal of Food Composition and Analysis*, 107, 104288. <https://doi.org/10.1016/j.jfca.2021.104288>
- Barbon-Junior, S., Mastelini, S. M., Barbon, A. P. A. C., Barbin, D. F., Calvini, R., Lopes, J. F., & Ulrici, A. (2020). Multi-target prediction of wheat flour quality parameters with near infrared spectroscopy. *Information Processing in Agriculture*, 7(2), 342–354. <https://doi.org/10.1016/j.inpa.2019.07.001>
- Basile, T., Marsico, A. D., & Perniola, R. (2022). Use of Artificial Neural Networks and NIR Spectroscopy for Non-Destructive Grape Texture Prediction. *Foods*, 11(3), 281. <https://doi.org/10.3390/foods11030281>
- Borchani, H., Varando, G., Bielza, C., & Larrañaga, P. (2015). A survey on multi-output regression. *WIREs Data Mining and Knowledge Discovery*, 5(5), 216–233. <https://doi.org/10.1002/widm.1157>

- Brasil, I. M., & Siddiqui, M. W. (2018). Postharvest Quality of Fruits and Vegetables: An Overview. En M. Siddiqui (Ed.), *Preharvest Modulation of Postharvest Fruit and Vegetable Quality* (1st ed., pp. 1–40). Elsevier Inc. <https://doi.org/10.1016/B978-0-12-809807-3.00001-9>
- Butz, P., Hofmann, C., & Tauscher, B. (2005). Recent Developments in Noninvasive Techniques for Fresh Fruit and Vegetable Internal Quality Analysis. *Jornal of food science*, 70(9), 131–141. <https://doi.org/https://doi.org/10.1111/j.1365-2621.2005.tb08328.x>
- Cavaco, A. M., Passos, D., Pires, R. M., Antunes, M. D., & Guerra, R. (2021). Nondestructive Assessment of Citrus Fruit Quality and Ripening by Visible–Near Infrared Reflectance Spectroscopy. En M. S. Khan & I. A. Khan (Eds.), *Citrus - Research, Development and Biotechnology* (3rd ed., pp. 251–281). London, United Kingdom: IntechOpen. <https://doi.org/10.5772/intechopen.95970>
- Dos Santos, D. A., de Lima, K. P., Consolin, M. F. B., Filho, N. C., Março, P. H., & Valderrama, P. (2019). Multi-product multivariate calibration: Determination of quality parameters in soybean industrialized juices. *Acta Scientiarum - Technology*, 41(1), e37382. <https://doi.org/10.4025/actascitechnol.v41i2.37382>
- Gemperline, P. J., Long, J. R., & Gregoriou, V. G. (1991). Nonlinear multivariate calibration using principal components regression and artificial neural networks. *Analytical Chemistry*, 63(20), 2313–2323. <https://doi.org/10.1021/ac00020a022>
- Goke, A., Serra, S., & Musacchi, S. (2018). Postharvest Dry Matter and Soluble Solids Content Prediction in d’Anjou and Bartlett Pear Using Near-infrared Spectroscopy. *HortScience*, 53(5), 669–680. <https://doi.org/10.21273/HORTSCI12843-17>
- GoodFellow, I., Bengio, J., & Courville, A. (2016). *Deep Learning*. London, England: MIT Press.
- Gutiérrez-Rico, V. (2017). *Cosecha, postcosecha y comercialización de la manzana*. Villa Alcalá.
- Kusumiyati, Hadiwijaya, Y., Putri, I. E., & Munawar, A. A. (2021a). Multi-product calibration model for soluble solids and water content quantification in Cucurbitaceae family, using visible/near-infrared spectroscopy. *Heliyon*, 7(8), e07677. <https://doi.org/10.1016/j.heliyon.2021.e07677>
- Kusumiyati, K., Hadiwijaya, Y., Putri, I. E., & Munawar, A. A. (2021b). Dataset of Vis-NIR spectra for intact cucurbitaceae fruits. En *Mendeley Data*, V1. <https://doi.org/doi:10.17632/k55b8mvs84.1>
- Kusumiyati, K., Hadiwijaya, Y., Sutari, W., & Munawar, A. A. (2022). Global model for in-field monitoring of sugar content and color of melon pulp with comparative regression approach. *AIMS Agriculture and Food*, 7(2), 312–325. <https://doi.org/10.3934/agrfood.2022020>

- Li, J., Wang, Q., Xu, L., Tian, X., Xia, Y., & Fan, S. (2019). Comparison and Optimization of Models for Determination of Sugar Content in Pear by Portable Vis-NIR Spectroscopy Coupled with Wavelength Selection Algorithm. *Food Analytical Methods*, 12(1), 12–22. <https://doi.org/10.1007/s12161-018-1326-7>
- Magwaza, L. S., Opara, U. L., Nieuwoudt, H., Cronje, P. J. R., Saeys, W., & Nicolaï, B. (2012). NIR Spectroscopy Applications for Internal and External Quality Analysis of Citrus Fruit-A Review. *Food and Bioprocess Technology*, 5(2), 425–444. <https://doi.org/10.1007/s11947-011-0697-1>
- Malek, S., Melgani, F., & Bazi, Y. (2017). One-dimensional convolutional neural networks for spectroscopic signal regression. *Journal of Chemometrics*, 32(5), e2977. <https://doi.org/10.1002/cem.2977>
- Meritxell-Gago, J. (2011). *Determinaciones cuantitativas de la composición de yogures por tecnología NIRS*. Universidad Pública de Navarra.
- Mishra, P., Biancolillo, A., Roger, J. M., Marini, F., & Rutledge, D. N. (2020). New data preprocessing trends based on ensemble of multiple preprocessing techniques. *Trends in Analytical Chemistry*, 132, 116045. <https://doi.org/10.1016/j.trac.2020.116045>
- Mishra, P., & Passos, D. (2021). A synergistic use of chemometrics and deep learning improved the predictive performance of near-infrared spectroscopy models for dry matter prediction in mango fruit. *Chemometrics and Intelligent Laboratory Systems*, 212, 104287. <https://doi.org/10.1016/j.chemolab.2021.104287>
- Mishra, P., & Passos, D. (2022). Multi-output 1-dimensional convolutional neural networks for simultaneous prediction of different traits of fruit based on near-infrared spectroscopy. *Postharvest Biology and Technology*, 183, 111741. <https://doi.org/10.1016/j.postharvbio.2021.111741>
- Mishra, P., Woltering, E., Brouwer, B., & Echtelt, E. H. (2021). Improving moisture and soluble solids content prediction in pear fruit using near-infrared spectroscopy with variable selection and model updating approach. *Postharvest Biology and Technology*, 171, 111348. <https://doi.org/10.1016/j.postharvbio.2020.111348>
- Morais, C. L. M., Santos, M. C. D., & Martin, F. L. (2019). Improving data splitting for classification applications in spectrochemical analyses employing a random-mutation Kennard-Stone algorithm approach. *Bioinformatics*, 35(4), 1–7. <https://doi.org/10.1093/bioinformatics/btz421>
- Munawar, A. A. (2014). *Multivariate analysis and artificial neural network approaches of near infrared spectroscopic data for non-destructive quality attributes prediction of Mango (Mangifera indica L.)*. Georg-August-Universität Göttingen.
- Murray, I. (2004). Scattered information: philosophy and practice of near infrared spectroscopy. En A. M. C. Davis & A. Garrido-Varo (Eds.), *Near Infrared*

- Spectroscopy: Proceedings of the 11th International Conference* (p. 12). Cordoba, Spain: NIR Publications. Recuperado de www.nirpublications.com
- Najjar, K., & Abu-Khalaf, N. (2021). Non-destructive quality measurement for three varieties of tomato using vis/nir spectroscopy. *Sustainability (Switzerland)*, 13(19), 10747. <https://doi.org/10.3390/su131910747>
- Ni, W., Nørgaard, L., & Mørup, M. (2014). Non-linear calibration models for near infrared spectroscopy. *Analytica Chimica Acta*, 813, 1–14. <https://doi.org/10.1016/j.aca.2013.12.002>
- Ozaki, Y., & Huck, C. (2021). Introduction. En Y. Ozaki, C. Huck, S. Tsuchikawa, & S. B. Engelsen (Eds.), *Near-Infrared Spectroscopy* (pp. 3–10). Singapore: Springer. <https://doi.org/https://doi.org/10.1007/978-981-15-8648-4>
- Ozaki, Y., & Morisawa, Y. (2021). Principles and Characteristics of NIR Spectroscopy. En Y. Ozaki, C. Huck, S. Tsuchikawa, & S. B. Engelsen (Eds.), *Near-Infrared Spectroscopy* (pp. 11–35). Singapore: Springer.
- Passos, D., & Mishra, P. (2021). An automated deep learning pipeline based on advanced optimisations for leveraging spectral classification modelling. *Chemometrics and Intelligent Laboratory Systems*, 215, 104354. <https://doi.org/10.1016/j.chemolab.2021.104354>
- Piccolo, M., Aceto, M., & Vitorino, T. (2019). UV-Vis spectroscopy. *Physical Sciences Reviews*, 4(4), 1–14. <https://doi.org/10.1515/psr-2018-0008>
- Rambo, M. K. D., Ferreira, M. M. C., & Amorim, E. P. (2015). Multi-product calibration models using NIR spectroscopy. *Chemometrics and Intelligent Laboratory Systems*, 151, 108–114. <https://doi.org/10.1016/j.chemolab.2015.12.013>
- Roger, J.-M., Boulet, J.-C., Zeaiter, M., & Rutledge, D. N. (2020). Pre-processing methods. En S. Brown, R. Tauler, & B. Walczak (Eds.), *Comprehensive Chemometrics Chemical and Biochemical Data Analysis* (2nd ed., p. 75). Elsevier. <https://doi.org/10.1016/b978-0-12-409547-2.14878-4>
- Saeys, W., Nguyen Do Trong, N., Van-Beers, R., & Nicolai, B. M. (2019). Multivariate calibration of spectroscopic sensors for postharvest quality evaluation: A review. *Postharvest Biology and Technology*, 158, 110981. <https://doi.org/10.1016/j.postharvbio.2019.110981>
- Sirisomboon, P. (2018). NIR Spectroscopy for Quality Evaluation of Fruits and Vegetables. *Materials Today: Proceedings*, 5(10), 22481–22486. <https://doi.org/10.1016/j.matpr.2018.06.619>
- Srivastava, S., & Sadistap, S. (2018). Non-destructive sensing methods for quality assessment of on-tree fruits: a review. *Journal of Food Measurement and Characterization*, 12(1), 497–526. <https://doi.org/10.1007/s11694-017-9663-6>
- Subedi, P. P., & Walsh, K. B. (2020). Assessment of avocado fruit dry matter content using portable near infrared spectroscopy: Method and

- instrumentation optimisation. *Postharvest Biology and Technology*, 161, 111078. <https://doi.org/10.1016/j.postharvbio.2019.111078>
- Teh, S. L., Coggins, J. L., Kostick, S. A., & Evans, K. M. (2020). Location, year, and tree age impact NIR-based postharvest prediction of dry matter concentration for 58 apple accessions. *Postharvest Biology and Technology*, 166, 111125. <https://doi.org/10.1016/j.postharvbio.2020.111125>
- Vidal, G. I. (2019). *Fenotipado rápido de germoplasma de tomate por características de calidad organoléptica para su uso en programas de mejora*. Universitat Jaume I.
- Walsh, K. B., Blasco, J., Zude-Sasse, M., & Sun, X. (2020). Visible-NIR “point” spectroscopy in postharvest fruit and vegetable assessment: The science behind three decades of commercial use. *Postharvest Biology and Technology*, 168, 111246. <https://doi.org/10.1016/j.postharvbio.2020.111246>
- Walsh, K. B., Mcglone, V. A., & Han, D. H. (2020). The uses of near infrared spectroscopy in postharvest decision support: A review. *Postharvest Biology and Technology*, 163, 111139. <https://doi.org/10.1016/j.postharvbio.2020.111139>
- Wang, H., Peng, J., Xie, C., Bao, Y., & He, Y. (2015). Fruit Quality Evaluation Using Spectroscopy Technology: A Review. *Sensors*, 15(5), 11889–11927. <https://doi.org/10.3390/s150511889>
- Wang, J., Wang, J., Chen, Z., & Han, D. (2017). Development of multi-cultivar models for predicting the soluble solid content and firmness of European pear (*Pyrus communis* L.) using portable vis–NIR spectroscopy. *Postharvest Biology and Technology*, 129, 143–151. <https://doi.org/10.1016/j.postharvbio.2017.03.012>
- Xie, L. J., Wang, A. C., Xu, H. R., Fu, X. P., & Ying, Y. B. (2016). Applications of Near-infrared systems for quality evaluation of fruits: A review. *Transactions of the ASABE*, 59(2), 399–419. <https://doi.org/10.13031/trans.59.10655>
- Yang, Y., Sun, S., Pan, L., Huang, M., & Zhu, Q. (2023). Predictions of multiple food quality parameters using near-infrared spectroscopy with a novel multi-task genetic programming approach. *Food Control*, 144, 109389. <https://doi.org/10.1016/j.foodcont.2022.109389>
- Yu, T., & Zhu, H. (2020). Hyper-Parameter Optimization: A Review of Algorithms and Applications. *arXiv*, 1–56. Recuperado de <http://arxiv.org/abs/2003.05689>
- Zhang, X., Lin, T., Xu, J., Luo, X., & Ying, Y. (2019). DeepSpectra: An end-to-end deep learning approach for quantitative spectral analysis. *Analytica Chimica Acta*, 1058, 48–57. <https://doi.org/10.1016/j.aca.2019.01.002>
- Zhang, X., Yang, J., Lin, T., & Ying, Y. (2021). Food and agro-product quality evaluation based on spectroscopy and deep learning: A review. *Trends in Food Science and Technology*, 112, 431–441.

<https://doi.org/10.1016/j.tifs.2021.04.008>

- Zhang, Y., Nock, J. F., Shoffe, Y. Al, & Watkins, C. B. (2019). Non-destructive prediction of soluble solids and dry matter contents in eight apple cultivars using near-infrared spectroscopy. *Postharvest Biology and Technology*, 151, 111–118. <https://doi.org/10.1016/j.postharvbio.2019.01.009>
- Zhou, L., Zhang, C., Liu, F., Qiu, Z., & He, Y. (2019). Application of Deep Learning in Food: A Review. *Comprehensive Reviews in Food Science and Food Safety*, 18(6), 1793–1811. <https://doi.org/10.1111/1541-4337.12492>

CAPÍTULO 3. ARTÍCULO CIENTÍFICO

PREDICCIÓN SIMULTÁNEA DE CALIDAD INTERNA EN MÚLTIPLES FRUTAS CON ESPECTROSCOPIA VIS-NIR Y APRENDIZAJE PROFUNDO

SIMULTANEOUS PREDICTION OF INTERNAL QUALITY IN MULTIPLE FRUITS WITH VIS-NIR SPECTROSCOPY AND DEEP LEARNING

Carlos Juárez González¹, Carlos Alberto Villaseñor Perea², Gilberto de Jesús López Canteñs³, Artemio Pérez López³ y Juan Carlos Olguín Rojas⁴

3.1 Resumen

Acceder a la calidad interna de la fruta fresca de manera rápida y no destructiva es de gran interés en toda la cadena de producción, desde el momento de cosecha hasta el destino del producto. La espectroscopía de punto Vis-NIR permite relacionar atributos de calidad con absorciones vibracionales, pero se necesitan calibrar modelos específicos para frutas o atributos de interés. El aprendizaje profundo (DL) tiene el potencial para el modelado global, por ello, el objetivo de este estudio fue investigar el potencial del DL en la calibración de modelos multiproducto-multisalida. Se usó una base de datos de acceso libre con 300 espectros de absorbancia Vis-NIR de seis frutas de la familia Cucurbitaceae con mediciones de referencia de contenido de agua (CA) y contenido de sólidos solubles (CSS). Se realizó un aumento de datos y con ellos se optimizaron automáticamente tres redes neuronales convolucionales unidimensionales (1D-CNN) de diferente profundidad y de salida múltiple para comparar su rendimiento con modelos de salida única y regresión por mínimos cuadrados parciales (PLSR). Al comparar los modelos base y óptimos, la raíz del error cuadrático medio combinado (RMSE) mejoró hasta 13.5%. El modelo de una capa convolucional obtuvo el mejor rendimiento en comparación con los modelos de dos y tres capas convolucionales, logró un RMSE ~ 3% menor que un estudio previo con PLSR de salida única y espectros preprocesados y hasta un 38% menor que con espectros en bruto. La optimización de modelos de salida única solo mejoró menos del 1% el RMSE del mejor modelo de salida múltiple y demoró el doble de tiempo. Aunque los tiempos de optimización superan por mucho la simplicidad de PLSR, el DL permite desarrollar modelos globales de manera más fácil y precisa.

Palabras clave: calibración, CNN unidimensional, modelo global, multiproducto-multisalida.

Tesis de maestría en ingeniería, Ingeniería Agrícola y Uso integral del Agua, Universidad Autónoma Chapingo. ¹Tesista, ²Director, ³Asesores, ⁴Colaborador

Abstract

Accessing the internal quality of fresh fruit quickly and non-destructively is of great interest throughout the production chain, from the moment of harvest to the destination of the product. Vis-NIR point spectroscopy allows quality attributes to be related to vibrational absorptions, but specific models need to be calibrated for fruits or attributes of interest. Deep learning (DL) has the potential for global modeling, therefore, the objective of this study was to investigate the potential of DL in the calibration of multi-product-multi-output models. A free access database with 300 Vis-NIR absorbance spectra of six fruits of the *Cucurbitaceae* family with reference measurements of water content (WC) and soluble solids content (SSC) was used. Data augmentation was performed and three one-dimensional convolutional neural networks (1D-CNN) of different depth and multiple output were automatically optimized to compare their performance with single output models and partial least squares regression (PLSR). When comparing the base and optimal models, the combined root mean square error (RMSE) improved until 13.5%. The one-convolutional layer model got the best performance compared to the two- and three-convolutional layer models, achieved ~3% lower RMSE than a previous study with single-output PLSR and preprocessed spectra and up to 38% lower than with raw spectra. Optimizing the single output models only improved the RMSE of the best multiple output model by less than 1% and took twice as long. Although optimization times far outweigh the simplicity of PLSR, DL makes global model development easier and more accurate.

Keywords: calibration, one-dimensional CNN, global model, multi-product-multi-output.

3.2 Introducción

La calidad de frutas y verduras depende de atributos externos e internos que determinan la aceptabilidad del consumidor, el precio de venta y su control juega un papel importante para su exportación (Jha et al., 2012; Pathmanaban, Gnanavel, & Anandan, 2019). Acceder a la calidad interna de manera rápida y no destructiva es de interés para los gestores de la cadena de producción que toman decisiones desde el momento de cosecha hasta el destino de los productos (Walsh, Mcglone, et al., 2020). Generalmente esta calidad se determina por métodos destructivos que requieren tiempo, son costosos, sesgados y no son adecuados para la industria (Lakshmi et al., 2017). La espectroscopía puntual Vis-NIR permite relacionar atributos de calidad interna con información espectral sobre absorciones vibracionales de enlaces como OH, CH y NH (H. Wang et al., 2015), lo que la convierte en una técnica rápida, no destructiva y que no requiere el uso de químicos (Ferrari, Foca, Calvinini, & Ulrici, 2015). Sin embargo, la aplicabilidad de esta técnica se ve limitada por la necesidad de crear modelos con gran variedad de muestras y principalmente porque deben calibrarse modelos específicos (Vidal, 2019).

La técnica Vis-NIR exige primeramente la recolección de los espectros de las frutas y posteriormente la calibración de modelos que transformen estos datos en los atributos de interés. Generalmente se requiere de un paso intermedio de corrección de espectros a través del preprocesamiento y la eliminación de información redundante, para poder calibrar un modelo preciso. Actualmente existe una gran cantidad de métodos/estrategias de preprocesamiento de datos (Mishra et al., 2020) y técnicas de calibración multivariada (Saeys et al., 2019), sin embargo, diferentes conjuntos de datos pueden requerir diferentes preprocesamientos y la elección puede afectar fuertemente el rendimiento del modelo (Mishra, Rutledge, Roger, Wali, & Khan, 2021; X. Zhang et al., 2019) por lo que la mejor combinación se investiga para cada aplicación (Torniainen et al., 2020). La regresión lineal múltiple, regresión por componentes principales y regresión por mínimos cuadrados parciales (MLR, PCR y PLSR, por sus siglas

en inglés) son los métodos de calibración más utilizados debido a que permiten manejar las múltiples variables espectrales y las altas correlaciones entre ellas. Varias revisiones resaltan el avance y potencial de la técnica Vis-NIR como método no destructivo de evaluación de calidad interna de una amplia variedad de frutas (Anderson & Walsh, 2022; Arendse et al., 2017; Walsh, Mcglone, et al., 2020; H. Wang et al., 2015), desarrollando modelos para predecir atributos como el contenido de sólidos solubles (CSS), contenido de materia seca (CMS), acidez titulable (AT) y firmeza principalmente. El uso práctico de un modelo requerirá que pueda generalizar para poder ser aplicado en muestras futuras, por lo que se debe calibrar con un conjunto de datos representativo de la población en la que se utilizará (Saeys et al., 2019), incluir datos con las múltiples variaciones posibles y reportar su rendimiento en muestras independientes al conjunto de calibración (Walsh, Mcglone, et al., 2020).

Desarrollar un modelo para un fruto y un atributo de calidad en particular resulta costoso y tardado, por lo que las investigaciones se han orientado al desarrollo de modelos globales. Se ha investigado la calibración de modelos multiproducto (Kusumiyati et al., 2021a; Masithoh, Lohumi, Yoon, Amanah, & Cho, 2020; Rambo et al., 2015) y multisalida (Mishra & Passos, 2022), modelos robustos con múltiples cultivares, orígenes geográficos o de árbol, grados de madurez y temporadas de cosecha, usando frutos ejemplo como pera (J. Wang et al., 2017), manzana (Bobelyn et al., 2010; Teh et al., 2020) y mango (Anderson et al., 2021, 2020), pero no se han modelado múltiples propiedades y múltiples productos en un solo modelo de regresión. En ocasiones los conjuntos de datos con gran variabilidad son puestos a disposición general para remodelarse con nuevas técnicas y enfoques quimiométricos con el objetivo de mejorar la precisión, la rapidez y la facilidad con la que se calibran los modelos.

El DL fue adaptado para la regresión multivariante a través de las redes 1D-CNN y se expandió al análisis espectroscópico (Cui & Fearn, 2018) aprovechando sus múltiples ventajas, principalmente que requiere de la mínima ingeniería de datos (Lecun, Bengio, & Hinton, 2015), manejo de relaciones no lineales, es escalable

y por su capacidad de actualización. Varias investigaciones demuestran la superioridad del DL sobre la regresión lineal (Mishra & Passos, 2021a, 2021c, 2021b), pero aún es un desafío configurar la gran cantidad de parámetros propios de las redes neuronales artificiales (ANN). El DL tiene el potencial para la modelación multiproducto-multisalida y comprobar su rendimiento contribuiría al desarrollo de modelos globales más precisos y de manera más simple, especialmente con métodos de optimización automática.

La base de datos de Kusumiyati et al. (2021b) contiene espectros en la región Vis-NIR de seis frutas de la familia *Cucurbitaceae* y dos referencias de calidad interna, CSS y contenido de agua (CA). Se identificó una oportunidad para investigar el potencial del DL en el modelado multiproducto-multisalida. El objetivo de este trabajo fue usar dicha base de datos para comprobar el rendimiento del modelado global multiproducto para la predicción simultánea de las dos referencias de calidad y compararlo con la regresión lineal de salida única reportada en estudios previos.

3.3 Materiales y métodos

3.3.1 Descripción de los datos

Para este estudio se usaron 300 espectros de muestras de seis productos de la familia *Cucurbitaceae*, los cuales son, calabacín, calabaza amarga, calabaza de cresta, melón, chayote y pepino, disponibles para hacer predicción de dos propiedades de calidad interna; CA y CSS (Kusumiyati et al., 2021b). El conjunto de datos contiene valores espectrales de absorbancia en bruto de la región 381 a 1065 nm con un paso de 3 nm. Asimismo, contiene valores de referencia de CSS en un rango de 3.17 - 8.48 %Brix y CA de 89.8 – 95.76 % de agua y están divididos en dos subconjuntos, 200 muestras para calibración y 100 para prueba. Los datos están relacionados a dos publicaciones anteriores (Kusumiyati et al., 2021a; Kusumiyati, Hadiwijaya, Putri, & Munawar, 2021c), cuyos resultados se usan como referencia de comparación con los del presente trabajo.

3.3.2 Aumento y separación de datos

Los 300 espectros disponibles son un conjunto de datos pequeño para el modelado con DL por lo que fue necesario un aumento de datos. Se puede aumentar el número de muestras generando copias un poco diferentes del espectro original manteniendo los mismos valores de las referencias. Se trata de simular formas esperadas de ruido como el desplazamiento, la intensidad y el cambio de pendiente de línea de base (Bjerrum, Glahder, & Skov, 2017). Se ha descubierto que la compensación y multiplicación funciona bien con el $\pm 10\%$ de la desviación estándar del conjunto de datos, mientras que el cambio de pendiente se puede establecer en ± 0.05 (Bjerrum et al., 2017), por lo que se aplicaron dichos cambios aleatorios al presente conjunto de datos. No existe un límite conocido para la cantidad de copias que se pueden crear pero se ha comprobado que 10 es una buena cantidad (Bjerrum et al., 2017; Mishra & Passos, 2022) e incluso se ha ejecutado por más de 100 veces (Nallan-Chakravartula, Moscetti, Bedini, Nardella, & Massantini, 2022) pero en este caso fue suficiente con aumentar los datos de entrenamiento en factor de 10. El aumento de datos es usado comúnmente en redes neuronales para lograr un entrenamiento eficaz y evitar el sobreajuste.

Se conservó la separación de los datos para entrenamiento y prueba del modelo de la fuente original, ya que garantiza una representación completa de la variabilidad en ambos conjuntos, así como para poder hacer una comparación justa de los resultados. El conjunto de entrenamiento se dividió en dos subconjuntos, para calibración y validación (o ajuste) de la red en proporciones 67/33% respectivamente, de dos maneras diferentes: separación aleatoria (RS) después del aumento de datos y separación por el método Kennard-Stone (K-S) (Kennard & Stone, 1969) seguida del aumento de datos (Figura 1), con el objetivo de saber si la división de datos tiene impacto en el rendimiento del modelo.

Para alimentar a la red los espectros se normalizaron por columnas usando la media y la desviación estándar del conjunto de entrenamiento como escala para todos los datos. Además, considerando que las variables objetivo tienen

diferentes rangos de valores y que se están modelando de manera conjunta, los valores CA y CSS fueron normalizados de 0 a 1 considerando un rango 80 a 100 y de 0 a 15, respectivamente. De esta manera se evita que la minimización de la pérdida combinada durante el entrenamiento no tenga preferencia para alguna de las predicciones.

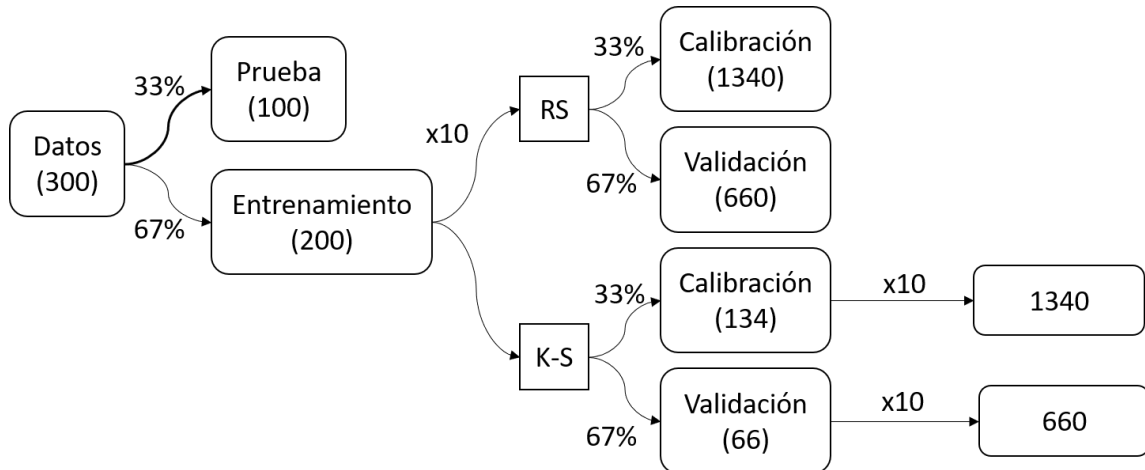


Figura 1. Separación y aumento de datos

3.3.3 Modelos 1D-CNN multisalida

En este trabajo se comparó el rendimiento de tres estructuras 1D-CNN usadas para resolver problemas de predicción similares, las cuales tienen diferente grado de profundidad. Estas redes se forman de tres capas principales, una de entrada, una de extracción de características y una de regresión. La red 1 (Figura 2a) está basada en la estructura propuesta por Cui & Fearn (2018). Su capa de entrada convierte los espectros en vectores para pasar a la siguiente, donde una capa convolucional (Conv1) extrae un mapa de características con un solo filtro y una capa plana los vuelve a vectorizar para pasar a la parte final con capas completamente conectadas (FC). La regresión la realiza un bloque de tres capas densas seguidas de capas de abandono y una capa de salida con dos neuronas, correspondientes a las dos referencias, CA y CSS. Las tres capas densas contienen 128, 64 y 32 neuronas, respectivamente. El tamaño y paso del filtro y las tasas de abandono fueron parte del grupo de hiperpámetros a optimizar utilizando el método que se menciona más adelante.

La red 2 fue adaptada de la estructura DeepSpectra propuesta por X. Zhang et al. (2019), con mayor profundidad tanto en ancho como en largo (Figura 2c). Su capa de extracción de características se forma de tres capas convolucionales (Conv1, Conv2 y Conv3) con varios filtros de diferentes tamaños y operaciones paralelas entre capas, lo que le permite tener mayor capacidad de aprendizaje de patrones y al mismo tiempo mantener constante la complejidad de cálculos al usar capas de convoluciones 1x1 y de agrupación máxima. Las salidas paralelas son concatenadas y vectorizadas por una capa plana antes de pasar a una capa FC. La regresión se realiza con una capa densa seguida de una capa de abandono, además, una capa de normalización por lotes antes y después de la capa densa acelera el entrenamiento y mejora la precisión. La salida consta de dos neuronas correspondientes a las dos referencias. El tamaño y paso de los filtros, el número de neuronas y la tasa de abandono se añadieron como hiperparámetros durante la optimización. Se ha comprobado que esta red tiene mejores resultados de predicción que redes con estructuras de menor y mayor profundidad (X. Zhang et al., 2019), pero no se ha probado su rendimiento en la regresión multisalida.

La red 3 fue adaptada a este problema de la estructura de varias ramas propuesta por Padarian, Minasny, & Mcbratney (2018). La extracción de características se divide en dos partes: una capa Conv1 para las generales y dos capas Conv2 para las específicas a cada salida (Figura 2b). Cada rama realiza la regresión de una referencia con un bloque de capas FC que contiene pares de capas densas, de abandono y de normalización por lotes para finalizar en una capa de salida con una neurona. Las capas Conv1 y Conv2 tienen diferentes filtros de diferentes tamaños y las capas densas contienen 64 y 16 neuronas, respectivamente. El tamaño y paso de los filtros y la tasa de abandono fueron hiperparámetros a optimizar. La motivación para usar esta estructura de red fue que las dos referencias de calidad podrían compartir características en los espectros y que una red con una sola salida podría desempeñarse mejor que su equivalente de dos salidas.

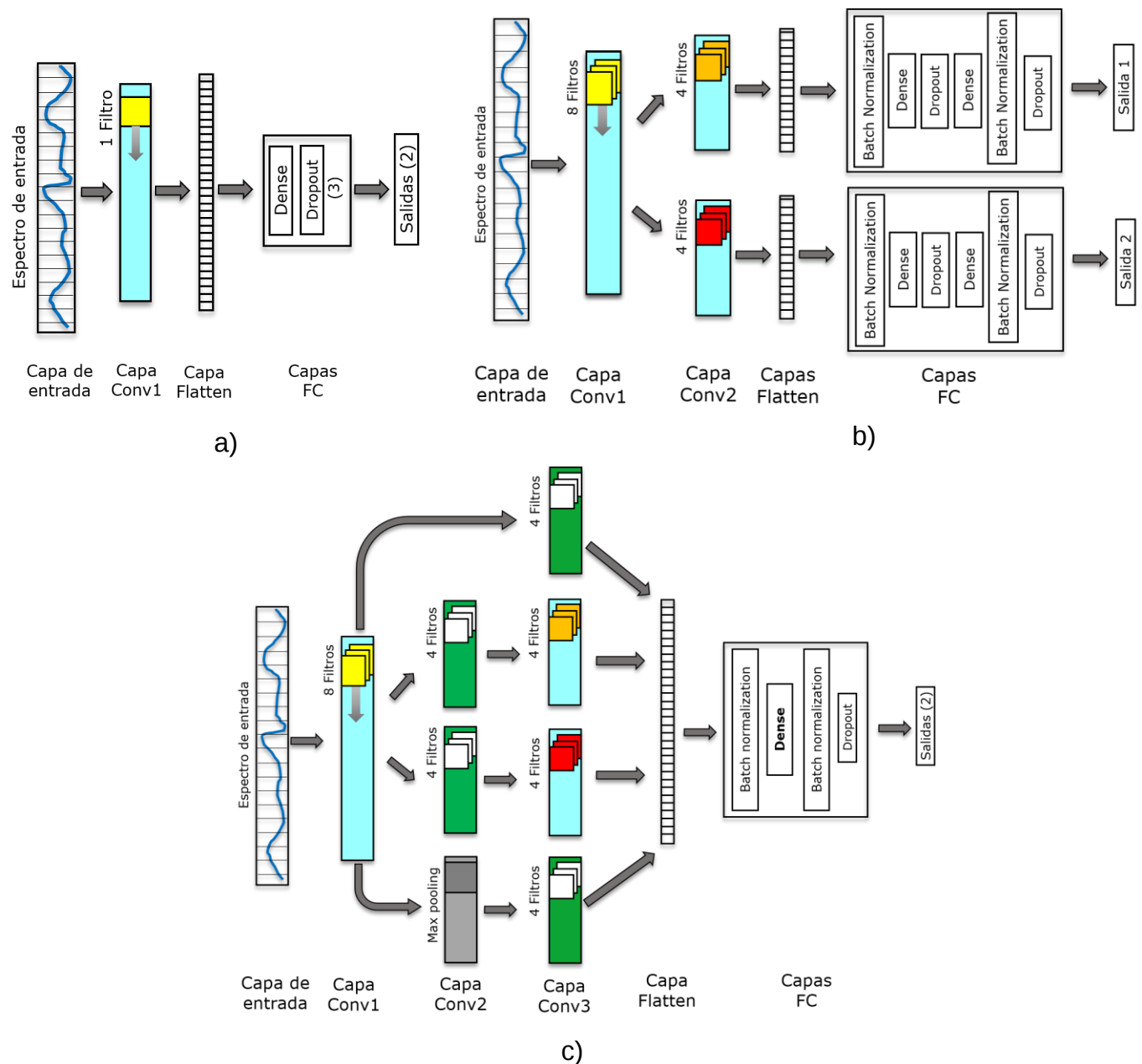


Figura 2. Estructuras 1D-CNN de salida múltiple: (a) red 1; (b) red 3; (c) red 2. Los cuadros pequeños dentro de los rectángulos representan tensores, los módulos cian representan convolución general, los verdes convolución 1x1 y el gris agrupación máxima. Los colores diferente en los filtros representa su tamaño y las flechas en múltiples direcciones representan conexión entre capas.

Dado que el rendimiento de un modelo de DL se ve fuertemente afectado por la inicialización de los pesos (GoodFellow et al., 2016), se usó la función "he_normal" con semilla 42 para inicializar los pesos en cada red durante el

proceso de optimización y garantizar mayor reproducibilidad. Los pesos se optimizaron con el algoritmo Adam, cuya tasa de aprendizaje (LR) se redujo automáticamente durante el entrenamiento a un mínimo de 1×10^{-6} con la función “ReduceLROnPlateau” en un factor 0.5, al no haber mejora en la pérdida de validación durante 25 épocas. La función de pérdida a optimizar fue el error cuadrático medio (MSE) con regularización L2 a los pesos de cada capa (Ecuación 1) excepto la de salida, para evitar el sobreajuste, además de la función “EarlyStopping” para detener el entrenamiento antes de las épocas definidas en caso de no haber un cambio mínimo de 1×10^{-4} en la pérdida de validación, con una paciencia de 50 épocas. Las capas de abandono también ayudan a evitar el sobreajuste al apagar estocásticamente una fracción de las neuronas durante el entrenamiento y no hacerlo durante la inferencia. En la red 1 se usó la función de activación de unidad lineal exponencial (ELU) para las capas Conv y densas, en las redes 2 y 3 se usó la función de activación de unidad lineal rectificada en su versión con fugas (LeakyReLU) y en las capas de salida siempre se usó una función lineal, debido al problema de regresión.

3.3.4 Optimización de hiperparámetros

Con las estructuras base se optimizaron otros hiperparámetros clave usando la canalización de optimización automática descrita por Passos & Mishra (2022) que hace uso de la técnica de optimización Bayesiana (BO) y que se puede aplicar en problemas de clasificación o regresión espectral. Se basa en probar diferentes configuraciones de hiperparámetros elegidos de manera informada sobre una función objetivo. Para ello usa un estimador Parzen de estructura de árbol (TPE) que estima densidades de probabilidad asociados a configuraciones de hiperparámetros “buenos” y “malos” (Bergstra, Bardenet, Bengio, & Kégl, 2011). En complemento, el algoritmo hiperbanda (L. Li, Jamieson, DeSalvo, Rostamizadeh, & Talwalkar, 2017) optimiza los recursos computacionales como el tiempo de entrenamiento o el número de épocas predefinidas, en relación a una función objetivo, para asignar más recursos a las configuraciones más prometedoras, por lo que actúa como una poda en los modelos que van por un mal camino y se evita el gasto de recursos en configuraciones “malas” de

hiperparámetros. La función objetivo para el estimador TPE fue la raíz del error cuadrático medio (RMSE) del conjunto de validación y para el algoritmo hiperbanda la pérdida de validación.

La OB debe partir de un número inicial de ensayos de manera aleatoria, es decir, sin usar el estimador TPE, para posicionar la exploración en el espacio de hiperparámetros y después explotar dicho espacio con una cantidad mayor de ensayos. Se usaron mínimo 10 ensayos iniciales y 150 adicionales por cada hiperparámetro a optimizar como lo recomiendan Passos & Mishra (2022) y se configuraron 450 épocas como máximo para cada ensayo.

En teoría esta técnica permite la optimización de todos los hiperparámetros de cualquier modelo, así como de búsqueda de arquitectura neuronal (NAS), pero es más práctico partir de una estructura base y optimizar hiperparámetros clave y/o cambiar restringidamente su arquitectura, por ejemplo, con la adición o no de capas de abandono (una capa con tasa cero equivale a no haber capa), con la modificación del número de neuronas en capas FC o con el uso o no de regularización L2 en los pesos. En este caso se partió de una configuración base hallada en la literatura (Cui & Fearn, 2018; Passos & Mishra, 2022; X. Zhang et al., 2019), modificada con experiencia personal y un procedimiento de prueba y error, lo que permitió tener un rendimiento inicial de comparación para el resultado de la BO. Los valores base, los rangos y los pasos de cada hiperparámetro y arquitectura se resumen en el Cuadro 1 para cada una de las redes usadas.

En cada ensayo (con una configuración de hiperparámetros) la optimización automática calibra un modelo con el subconjunto de calibración, lo ajusta con el de validación y guarda el mejor modelo con la función “ModelCheckpoint”. Posteriormente carga el modelo para calcular sobre el subconjunto de validación el RMSE combinado para las referencias CA y CSS y lo usa como función objetivo para la optimización.

Cuadro 1. Intervalo de valores usados para la optimización de hiperparámetros y de arquitectura 1D-CNN.

Hiperparámetro o arquitectura	(Valor base) [Rango] / Paso		
	Red 1	Red 2	Red 3
Tamaño de filtro amarillo	(11) [3 – 49] / 2	(7) [3 – 29] / 2	(11) [3 – 29] / 2
Paso de filtro amarillo	(1) [1 – 10] / 1	(3) [1 – 10] / 1	(3) [1 – 10] / 1
Regularización L2	(0.0005) [0 – 0.05] ^a / 0.0005	(0.001) [0.0005 – 0.05] ^a / 0.0005	(0.002) [0.003 – 0.05] ^a / 0.0005
Tamaño de batch	(64) [32 – 128] / 32	(32) [64 – 128] / 64	(64) [32 – 128] / 32
Tasa de abandono	(0) ^b [0 – 0.6] / 0.005	(0.2) [0 – 0.6] / 0.005	(0.1) ^b [0 – 0.6] / 0.005
Tasa de aprendizaje inicial	() ^c [0.0025 – 0.02] / 0.0025	(0.01) [1x10-8 – 0.07]	() ^c [0.0025 – 0.02] / 0.0025
Tamaño de filtro naranja		(3) [3 – 29] / 2	(9) [3 – 29] / 2
Paso de filtro naranja		(3) [1 – 10] / 1	(2) [1 – 10] / 1
Tamaño de filtro rojo		(5) [3 – 29] / 2	(7) [3 – 29] / 2
Paso de filtro rojo		(3) [1 – 10] / 1	(2) [1 – 10] / 1
Paso de filtros en Conv2		(2) [3 – 29] / 2	
Número de neuronas		(16) [8 – 128] / 8	
Ensayos iniciales	100	100	130
Ensayos totales	1300	1600	2100

^a Para datos con división K-S el límite inferior fue 0. ^b Tasa de abandono específica para cada capa. ^c Tasa de aprendizaje inicial determinada por la heurística LR=0.01*Bach/256

El problema de sobreajuste o desajuste que sufren las arquitecturas de ANN puede ser fuertemente expuesto durante este proceso de optimización automática, especialmente si incluye NAS que permita una arquitectura poco profunda (Passos & Mishra, 2022), o si los conjuntos de datos no están bien equilibrados. Para evitar lo primero se limita a la NAS restringida mencionada anteriormente y para lo segundo se usan métodos de regularización. Añadir capas de abandono es la manera más simple de prevenir el sobreajuste y la regularización L2 es la forma más común de regularizar porque afecta directamente a los pesos w del modelo en un factor λ , cómo se observa en la Ecuación 1 (Cui & Fearn, 2018).

$$Loss = MSE + \frac{1}{2}\lambda \sum w_i^2 \quad (1)$$

Un modelo más profundo puede tener el mejor rendimiento si está bien regularizado (Lecun et al., 2015) o sobre ajustarse en caso contrario. Para evitar

este problema con las redes más profundas de esta investigación se acotó un límite inferior diferente de cero para el rango de valores de regularización L2 en la optimización con datos RS.

3.3.5 Evaluación de resultados

Para la evaluación del desempeño de los modelos en la predicción se usaron las métricas RMSE y coeficiente de determinación R^2 (Ecuación 2 y 3). En los subconjuntos de calibración y validación se definió el RMSEC y RMSEV, respectivamente, mientras que para el conjunto de prueba se usó el RMSEP y R^2 . Las unidades del error están dadas en valores normalizados de CA y CSS, pero también se desnormalizaron las predicciones para poder calcular los errores en valores originales de las variables y compararlos con estudios anteriores.

$$RMSE = \sqrt{\frac{1}{N} \sum_{n=1}^N (y_n - \hat{y}_n)^2} \quad (2)$$

$$R^2 = 1 - \frac{\sum_{n=1}^N (y_n - \hat{y}_n)^2}{\sum_{n=1}^N (y_n - \bar{y})^2} \quad (3)$$

Donde N es el número de muestras sometidas a predicción, y_n , $\hat{y}_n$ y $\bar{y}$ son el valor objetivo real, predicho y promedio, respectivamente. El mayor rendimiento se definió con el menor RMSE y mayor R^2 .

3.3.6 Implementación

El diseño y entrenamiento de los modelos se realizó con las librerías de código abierto TensorFlow y Keras (2.8). Para la optimización de los modelos se usó la biblioteca Optuna (2.10.1). Todo se programó con lenguaje Python (3.7.14) en un entorno de ejecución de Google Colaboratory con ambiente de GPU.

3.4 Resultados y discusión

3.4.1 Análisis de los datos

Las referencias de calidad cubren un amplio rango de valores que se equilibran entre los conjuntos de entrenamiento y prueba (Cuadro 2). El pepino tiene el valor promedio más bajo de CSS y el chayote de CA mientras que el melón tiene los valores más altos para ambas referencias, el calabacín presenta la mayor

variación en ambas propiedades seguido del melón para CSS, la variabilidad de CSS es mucho mayor que CA en todos los productos. Localmente los productos con mayor variación tuvieron un menor R^2 y viceversa, con el mejor modelo, pero la variación general mejoró el rendimiento total. Los espectros de los productos tienen un patrón similar, pero presentan múltiples variaciones (Figura 3). La región más diferenciada es la visible (400-700 nm), que está relacionada al color de la muestra, la absorbancia de clorofila en 680 nm se observa como un pico para los productos verdes. Las bandas de absorción de enlaces H-O y C-H están entre 700-800 nm y entre 900-1000 nm (Murray, 2004), respectivamente, y se observan como valles y picos en dichas regiones mientras que otras características de absorción se superponen y hacen complejo el espectro. Los atributos de calidad pueden ser relacionados con la presencia de estos enlaces químicos, incluso con el contenido de pigmentos en la muestra (H. Wang et al., 2015).

Cuadro 2. Referencias de calidad interna en seis frutas de la familia *Cucurbitaceae*.

Atributo	Producto	Completo (M $\pm$ DS*)	Entrenamiento (M $\pm$ DS*)	Prueba (M $\pm$ DS*)	R ²
CA	Calabacín	94.00 $\pm$ 1.3	93.99 $\pm$ 1.3	94.03 $\pm$ 1.3	0.50
	Calabaza amarga	90.28 $\pm$ 1	90.31 $\pm$ 1.1	91.36 $\pm$ 0.9	0.14
	Calabaza de cresta	92.16 $\pm$ 0.7	92.00 $\pm$ 0.8	92.49 $\pm$ 0.5	0.27
	Melón	95.77 $\pm$ 0.5	95.83 $\pm$ 0.5	95.63 $\pm$ 0.6	0.54
	Chayote	89.80 $\pm$ 0.7	89.78 $\pm$ 0.6	89.85 $\pm$ 0.8	0.96
	Pepino	91.76 $\pm$ 0.4	91.67 $\pm$ 0.4	91.93 $\pm$ 0.5	0.83
	<i>General</i>	92.3 $\pm$ 2.4	92.29 $\pm$ 2.4	92.31 $\pm$ 2.3	0.93
CSS	Calabacín	3.77 $\pm$ 14.2	3.74 $\pm$ 15.4	3.83 $\pm$ 11.6	0.28
	Calabaza amarga	4.06 $\pm$ 9.2	4.08 $\pm$ 9.2	3.62 $\pm$ 10.2	0.34
	Calabaza de cresta	3.23 $\pm$ 6.1	3.24 $\pm$ 6.2	3.22 $\pm$ 6.1	0.33
	Melón	8.49 $\pm$ 13.3	8.54 $\pm$ 13.4	8.39 $\pm$ 13.4	0.82
	Chayote	3.82 $\pm$ 5.3	3.80 $\pm$ 5.4	3.84 $\pm$ 5.1	0.70
	Pepino	3.18 $\pm$ 5.5	3.20 $\pm$ 5.9	3.12 $\pm$ 4.2	0.54
	<i>General</i>	4.42 $\pm$ 43.6	4.45 $\pm$ 44.1	4.37 $\pm$ 42.6	0.97

*DS se presenta como el porcentaje de desviación estándar respecto a la media (M) de valores en cada producto. Cada R² corresponde a ~17 productos de prueba.

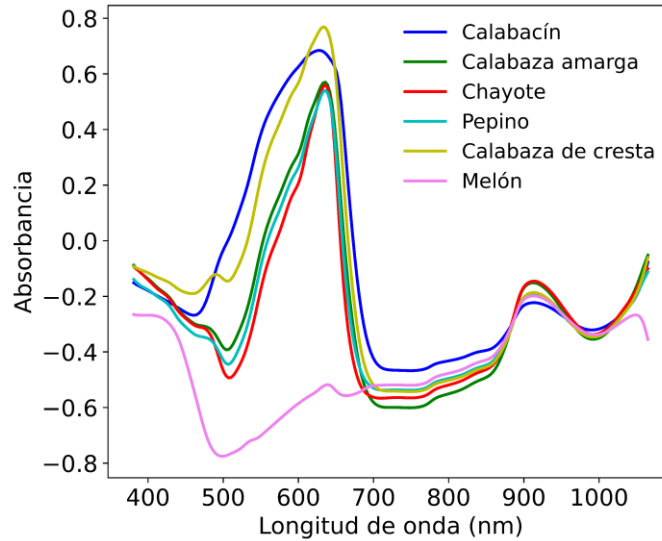


Figura 3. Espectros promedio de seis productos de la familia *Cucurbitaceae*

Trabajar con los espectros en bruto es difícil, pero no depender del correcto preprocesamiento ahorraría tiempo y contribuiría en la implementación generalizada de esta técnica. Por ello, los únicos preprocesamientos a los datos fueron el aumento de datos y la estandarización por predictores (columnas), de esta manera las muestras aumentadas de la Figura 4A, se transformaron en los datos de la Figura 4B.

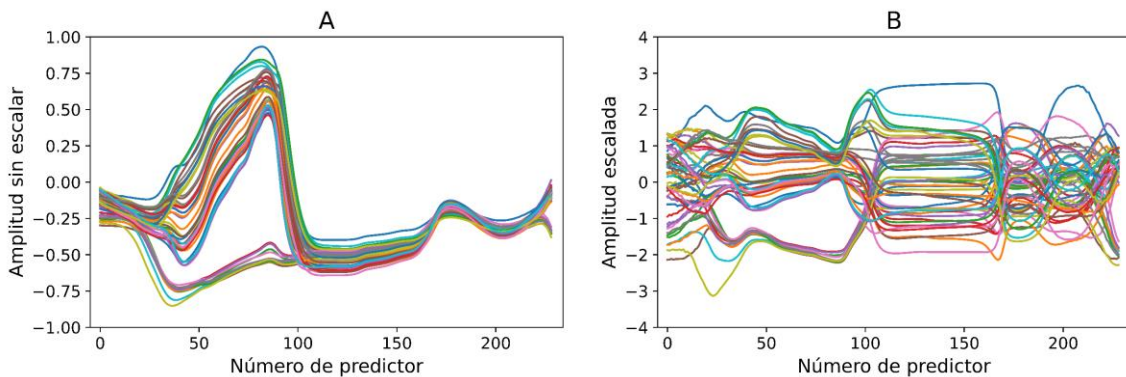


Figura 4. Muestra de 50 espectros aumentados y asignados al conjunto de validación, antes (A) y después (B) de la estandarización.

3.4.2 Regularización de modelos

Con la configuración base de cada red la primera búsqueda de hiperparámetros fue el valor mínimo de regularización para evitar que los modelos “aprendan” las salidas en vez de generalizar los patrones. Los resultados de esta exploración se

muestran en la Figura 5 con el comportamiento del RMSE en cada subconjunto de datos en función de la regularización L2 (β). Se observa que para datos RS, en cada red el RMSEC y RMSEV es muy similar (Figura 5a-c), debido a la similitud en la variabilidad de los datos en ambos subconjuntos. En cambio, la división K-S seguida del aumento de datos creó dos subconjuntos con diferente grado de complejidad, por lo que los RMSEC y RMSEV son muy diferentes y el segundo es mayor (Figura 5d-e). El conjunto de prueba siempre es el mismo y se usó para detectar sobreajuste cuando su error es mayor al de entrenamiento. Con una tasa de abandono del 0.1 en la red 1 fue suficiente para generalizar incluso sin regularización L2 (Figura 5a), pero la red 2 y 3 necesitan adicionalmente a su tasa de abandono un valor β mínimo de 0.0005 y 0.003, respectivamente (Figura 5b, c). Por su parte, el RMSEP de los modelos calibrados con datos K-S siempre se mantiene menor o entre el RMSEC y RMSEV (Figura 5d-f), lo que indica que las redes pueden generalizar incluso sin regularización L2. Por ello se eligen estos valores de β como límite mínimo para la optimización de cada red (línea discontinua verde). Al ejecutar una optimización sin este límite, para la red 2 se encontró una tasa de abandono de 0.13 y $\beta=0$ como hiperparámetros óptimos y se obtuvieron valores de RMSEC=0.0212, RMSEV=0.0225 y RMSEP=0.0279, que muestran un claro ejemplo de sobreajuste. Por ello, una tasa de abandono no es suficiente para evitar esta tendencia.

El hecho de que con mayor fuerza de regularización aumente el error se debe a la disminución de la capacidad de las redes, pero se espera que la optimización disminuya el error de entrenamiento sin disminuir la brecha observada entre error de entrenamiento y prueba, a mayor velocidad de la que disminuye el error de entrenamiento. La similitud del RMSEV con el RMSEP justifica el uso de dicho error como función objetivo de la optimización. Adicionalmente, se puede observar que los errores de predicción en el conjunto de prueba son más bajos con el modelo de la red 3, lo que significa que es más fácil hallar manualmente una configuración de hiperparámetros en la red de dos ramas para calibrar un modelo más preciso.

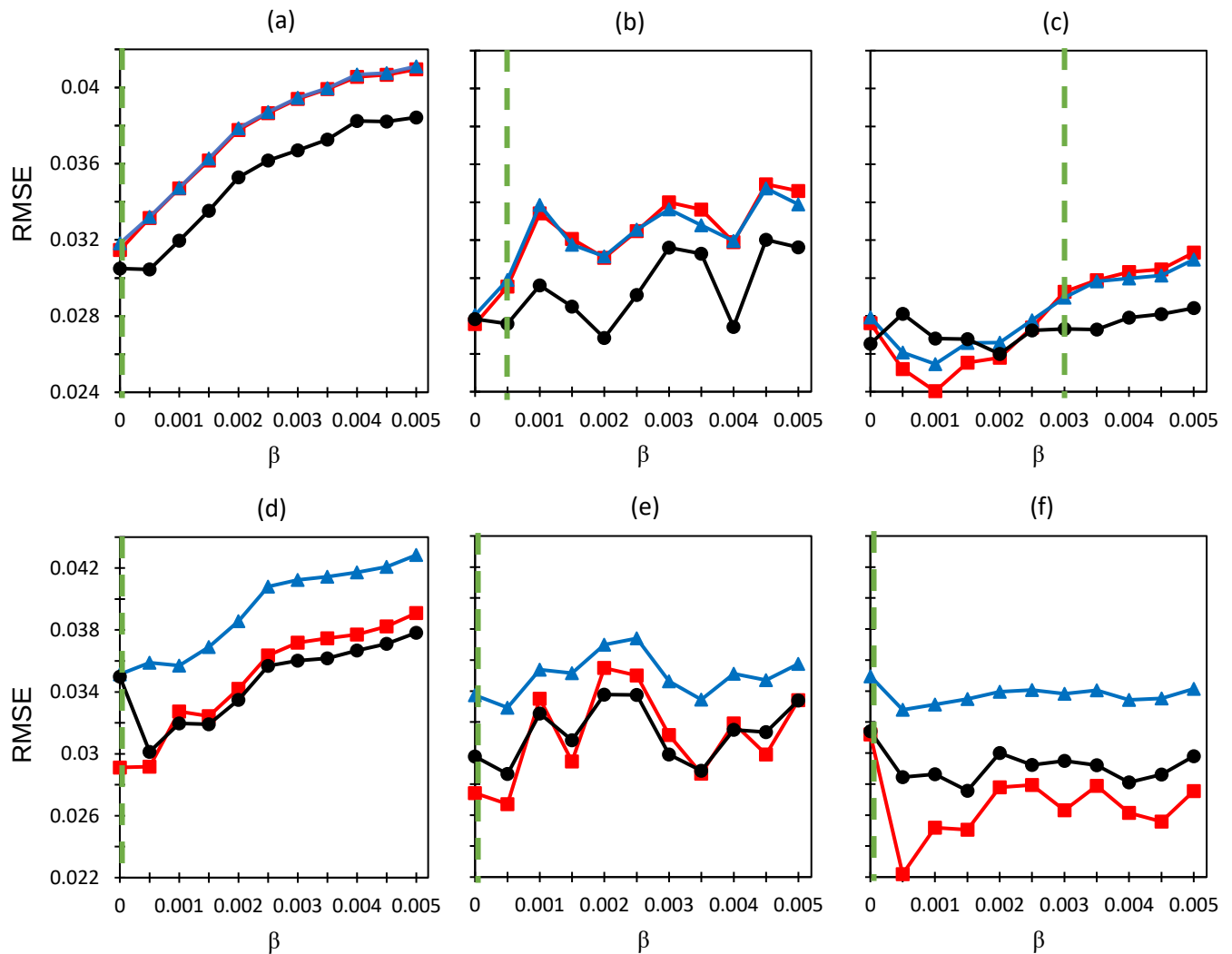


Figura 5. Criterios de selección de mínima fuerza de regularización de pesos para la optimización con datos RS (a, b, c) y K-S (d, e, f), de las redes 1 (a, d), 2 (b, e) y 3 (c, f). En rojo, el RMSE de calibración, azul RMSE de ajuste y negro RMSE de prueba. La línea discontinua verde representa el punto donde las redes comienzan a generalizar.

3.4.3 Optimización de los modelos

En la estación de trabajo utilizada con una RAM asignada de 16GB, el proceso de optimización de los 1300 modelos con la red de una capa Conv tomó un tiempo de 10h 48min y 7h 48min con datos RS y K-S, respectivamente. Para el modelo de tres capas Conv tomó 18h 20min y 11h 50min con datos RS y K-S, respectivamente, en optimizar los 1600 modelos. Por su parte, la optimización de los 2100 modelos con la red de dos ramas demoró 21h 50min y 23h 50min con

datos RS y K-S, respectivamente. Los hiperparámetros optimizados se muestran en el Cuadro 3 y en la última fila la mejor prueba para cada caso.

Cuadro 3. Hiperparámetros optimizados para tres modelos 1D-CNN

Hiperparámetro o arquitectura	Valor (Datos RS / datos K-S)		
	Red 1	Red 2	Red 3
Tamaño de filtro amarillo	23 / 33	19 / 27	23 / 15
Paso de filtro amarillo	3 / 7	2 / 4	10 / 5
Regularización L2	0 / 0	0.0005 / 0.006	0.003 / 0.002
Tamaño de Batch	64 / 64	64 / 64	64 / 64
Tasa de abandono por capa	[0.135, 0.31, 0] / [0.225, 0.02, 0.115]	0.205 / 0.105	[0.075, 0.075, 0.395, 0.195]* / [0.045, 0.36, 0.15, 0.175]*
Tasa de aprendizaje inicial	0.0075 / 0.015	0.0075 / 0.005	0.0025 / 0.0075
Tamaño de filtro naranja		25 / 19	17 / 15
Paso de filtro naranja		2 / 1	1 / 10
Tamaño de filtro rojo		19 / 19	21 / 7
Paso de filtro rojo		2 / 1	2 / 1
Pasos de filtros Conv2		2 / 1	
Número de neuronas		24 / 72	
Mejor ensayo	605 / 832	1482 / 1595	895 / 1640

* Las dos primeras tasas corresponden a la rama superior y las dos restantes a la rama inferior.

Con los hiperparámetros óptimos se creó una nueva instancia de 1D-CNN, se entrenó por 50 ejecuciones repetidas durante un máximo de 400 épocas y se registró el promedio y la desviación estándar del RMSE en cada subconjunto, para tener una medida más robusta del rendimiento de cada red. Además, la tasa de aprendizaje inicial al ser considerada el hiperparámetro más importante (GoodFellow et al., 2016) fue reajustada variándola a un valor cercano al encontrado por la optimización, para mejorar aún más el error de validación promedio de las 50 repeticiones. De esta manera la LR fue de 0.006 y 0.015 para la red 1, 0.007 y 0.004 para la red 2 y para la red 3 de 0.0025 y 0.007.

Para mejor comprensión el Cuadro 4 muestra los RMSE de predicción combinados para CA y CSS en valores normalizados. Se observa que la optimización Bayesiana mejoró, en relación al valor base, el RMSEV en las redes 1 y 2 pero no en la 3, sin embargo sí mejoró en las tres redes el RMSEP, lo que concuerda con los resultados de los desarrolladores del método (Passos & Mishra, 2022). El aumento del RMSEV en la red 3 se debe a que es un modelo

más regularizado que el modelo base (mayor tasa de abandono y fuerza de regularización) y se refleja en la mejora del RMSEP. La separación de datos tuvo un efecto significativo en el rendimiento de todos los modelos, con una ventaja para RS (Cuadro 4 y Cuadro 5), contrario a los resultados de Ferreira, Teixeira, & Peternelli (2022), donde obtienen errores más bajos con división K-S en comparación con RS. Sin embargo, se debe considerar que el presente problema trata con muestras multiproducto y una selección aleatoria tiene mayor probabilidad de separar el mismo porcentaje de datos para cada producto, lo que podría resultar en un conjunto de validación más similar al conjunto de prueba. La red de tres capas Conv presenta la mayor variación en sus resultados con desviación estándar de RMSEP de hasta $\pm 3.8\%$ el promedio, pero es menor que la reportada por X. Zhang et al. (2019) de hasta $\pm 17\%$, debido a que en este caso se está configurando una semilla para la inicialización de los pesos. Por su parte la red de dos ramas es más robusta que la red de tres capas Conv y la red de una sola capa Conv tiene el rendimiento más robusto, esto indica que la profundidad de una 1D-CNN aumenta la aleatoriedad propia de la modelación estocástica, en especial cuando se utiliza GPU. Esta variación permitió obtener de las repeticiones, modelos con mayor y menor precisión que el promedio y en la columna “Mejor” del Cuadro 4 se indica el RMSEP más bajo en cada caso.

Cuadro 4. Errores de predicción de tres modelos 1D-CNN antes y después de la optimización Bayesiana. Los resultados se presentan como la media y desviación estándar de los errores combinados para las dos predicciones, en unidades normalizadas.

Datos	Red	Base (Media $\pm$ DS)			Optimizada (Media $\pm$ DS)			Mejora		Mejor
		RMSEC	RMSEV	RMSEP	RMSEC	RMSEV	RMSEP	RMSEV	RMSEP	
RS	N1	0.0307 $\pm 1.4\text{E-}9$	0.0305 $\pm 1.7\text{E-}9$	0.0286 $\pm 2.2\text{E-}9$	0.0270 $\pm 1.9\text{E-}9$	0.0274 $\pm 1.9\text{E-}9$	0.0247 $\pm 3.3\text{E-}9$	10.3%	13.5%	0.0247
	N2	0.0322 $\pm 1.5\text{E-}3$	0.0321 $\pm 1.5\text{E-}3$	0.0287 $\pm 1.7\text{E-}3$	0.0275 $\pm 1.4\text{E-}3$	0.0279 $\pm 1.2\text{E-}3$	0.0265 $\pm 1.0\text{E-}3$	13.2%	7.9%	0.0247
	N3	0.0264 $\pm 6.4\text{E-}4$	0.0270 $\pm 4.5\text{E-}4$	0.0266 $\pm 3.6\text{E-}4$	0.0279 $\pm 3.3\text{E-}4$	0.0283 $\pm 2.6\text{E-}4$	0.0256 $\pm 3.9\text{E-}4$	-4.55%	3.5%	0.0254

K-S	N1	0.0239 $\pm 1.2\text{E-}8$	0.0338 $\pm 5.8\text{E-}9$	0.0294 $\pm 6.9\text{E-}9$	0.0294 $\pm 2.4\text{E-}9$	0.0334 $\pm 1.5\text{E-}9$	0.0290 $\pm 3.2\text{E-}9$	1.0%	1.4%	0.0290
	N2	0.0304 $\pm 1.8\text{E-}3$	0.0347 $\pm 1.1\text{E-}3$	0.0303 $\pm 1.5\text{E-}3$	0.0291 $\pm 1.2\text{E-}3$	0.0343 $\pm 5.6\text{E-}4$	0.0297 $\pm 6.9\text{E-}4$	1.1%	1.9%	0.0289
	N3	0.0219 $\pm 2.0\text{E-}3$	0.0327 $\pm 4.9\text{E-}4$	0.0289 $\pm 6.1\text{E-}4$	0.0270 $\pm 1.4\text{E-}3$	0.0334 $\pm 5.1\text{E-}4$	0.0280 $\pm 8.8\text{E-}4$	-2.4%	3.4%	0.0276

El número de iteraciones para convergencia antes y después de la optimización cambió sustancialmente para la red 1 y se mantuvo similar para la red 2 y 3 (Figura 6). Durante los entrenamientos repetidos, la red 1 siempre dejó de entrenar a las 154 épocas y nunca converge en una pérdida mayor que las demás redes. Se confirma que el modelo propuesto por Cui & Fearn (2018) a pesar de ser simple es un modelo promisorio para el modelado de datos espectrales de manera precisa, rápida y robusta (Mishra & Passos, 2021b, 2021a, 2022). El modelo de la red 2 es más profundo que la red 3, sin embargo, su tiempo de optimización fue menor gracias a que el módulo Inception en las capas Conv2 y Conv3 añaden profundidad a la red sin aumentar el costo computacional (Szegedy et al., 2015), su pérdida de convergencia promedio es mayor a la red 1 (Figura 6), lo cual se debe a su varianza, sin embargo, puede lograr el mismo rendimiento en el conjunto de prueba que la red 1 (Figura 7).

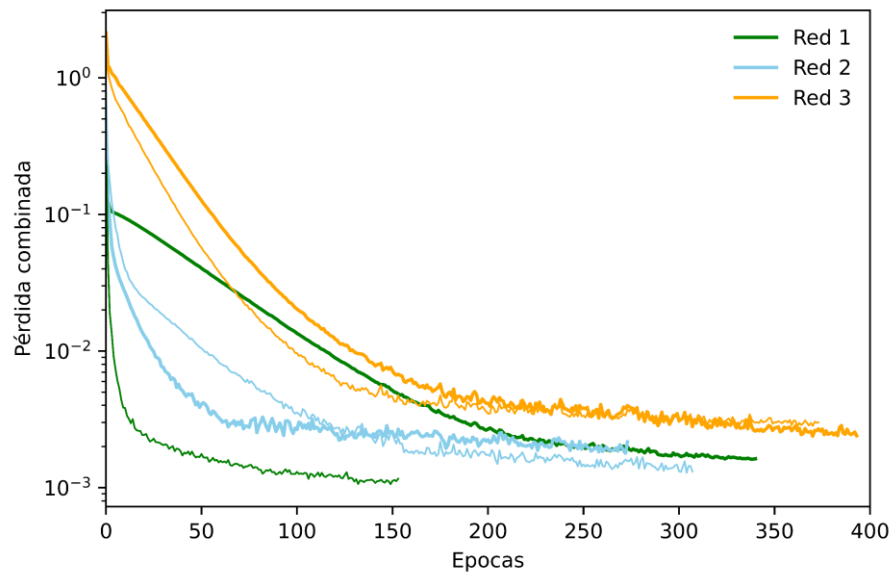


Figura 6. Evolución de la pérdida de validación combinada para CA y CSS durante un entrenamiento de modelos 1D-CNN antes (línea gruesa) y después (línea fina) de la optimización.

Durante la optimización de la red 3 el algoritmo hiperbanda detuvo menor número de ensayos, lo que indica que muchas configuraciones de hiperparámetros fueron buenas (530 en comparación con 423 en la red 2), considerando además que su configuración base tuvo el mejor rendimiento de las tres, se puede inferir que es más probable encontrar una configuración óptima al azar en la red de dos ramas.

Las 50 predicciones de los modelos óptimos se desnormalizaron para calcular el coeficiente de correlación entre valores reales y predichos, así como el RMSEP para cada referencia en las mismas unidades de los atributos de calidad. El modelo de red 1 obtuvo el RMSEP promedio más bajo de 0.566 y 0.309 para CA y CCS, respectivamente, y la mejor correlación con un R^2 medio de 0.93 y 0.972 para CA y SSC, respectivamente, que es menor que el mejor reportado por Kusumiyati et al. (2021a) usando regresión PLS con espectros preprocesados y corrección de la señal ortogonal, asimismo es mejor que el mejor reportado por Kusumiyati et al. (2021c) usando regresión por componentes principales (Cuadro 5). Con este modelo de DL el RMSE combinado para ambas referencias fue aproximadamente 3% menor que PLSR con espectros preprocesados, pero al considerar el RMSE con espectros en bruto (0.78 y 0.43 para CA y CSS, respectivamente), la mejora fue de más del 38%.

El modelo de una sola capa Conv tiene un R^2 significativamente mayor para ambos atributos de calidad que los otros modelos lineales y no lineales con nivel de confianza del 95%. Con el mismo nivel de confianza, los modelos 2 y 3 tienen un R^2 de predicción de SSC significativamente más alto que los modelos PLSR, pero no hay diferencia significativa en la predicción de CA. Por su parte, todos los modelos 1D-CNN superaron los resultados de PCR.

Cuadro 5. Rendimiento de predicción de tres modelos 1D-CNN en comparación con PLSR y PCR. Los resultados se expresan como el promedio y la desviación estándar de 50 repeticiones del entrenamiento.

Datos	Métrica	Red 1 (Media $\pm$ DS)	Red 2 (Media $\pm$ DS)	Red 3 (Media $\pm$ DS)	PLSR*	PCR*
RS	Contenido de agua (%)					
	RMSEP	0.566 $\pm 3.4\text{E-}16$	0.612 ± 0.023	0.599 ± 0.006	0.58	0.85
	R ²	0.930 $\pm 9.0\text{E-}16$	0.918 ± 0.006	0.921 ± 0.002	0.92	0.84
	Contenido de sólidos solubles (%Brix)					
	RMSEP	0.309 $\pm 3.4\text{E-}16$	0.323 ± 0.018	0.313 ± 0.005	0.32	0.38
	R ²	0.972 $\pm 1.0\text{E-}15$	0.97 ± 0.004	0.971 ± 0.001	0.96	0.95
K-S	Contenido de agua (%)					
	RMSEP	0.636 $\pm 2.1\text{E-}7$	0.677 ± 0.016	0.644 ± 0.022		
	R ²	0.911 $\pm 5.5\text{E-}8$	0.90 ± 0.005	0.909 ± 0.006		
	Contenido de sólidos solubles (%Brix)					
	RMSEP	0.388 $\pm 9.4\text{E-}8$	0.373 ± 0.012	0.354 ± 0.012		
	R ²	0.956 $\pm 2.2\text{E-}8$	0.959 ± 0.003	0.964 ± 0.002		

*Resultados de estudios previos. La división de datos es para calibración/ajuste de las redes y no aplica para estas columnas. En negritas los mejores valores.

La Figura 7 muestra un resumen de los mejores modelos en cada red. Como consecuencia de la variación de la red 2, esta pudo generar un modelo con rendimiento similar al de la red 1 y la red 3 el siguiente mejor modelo. Aunque los resultados de la red 2 y 3 dependen más de la aleatoriedad en las repeticiones, en la práctica se preferirá el modelo con el menor RMSE y mayor R², que podría ser encontrado con las repeticiones apropiadas. Los coeficientes de correlación más altos encontrados con el modelo multiproducto-multisalida fueron mejores que los obtenidos en la predicción de CA en pepino (Kavdir, Lu, Ariana, & Ngouajio, 2007), predicción de CSS en melón (Ito, Morimoto, & Yamauchi, 2001; M. Li, Han, & Liu, 2019) y predicción de CA en melón (Hadiwijaya, Kusumiyati, & Munawar, 2020). La variabilidad global mejoró el rendimiento en la modelación no lineal, similar al estudio de Anderson et al. (2021), mientras que en modelos lineales se han encontrado mejores resultados con modelos locales (Anderson et al., 2020; Teh et al., 2020; Y. Zhang et al., 2019).

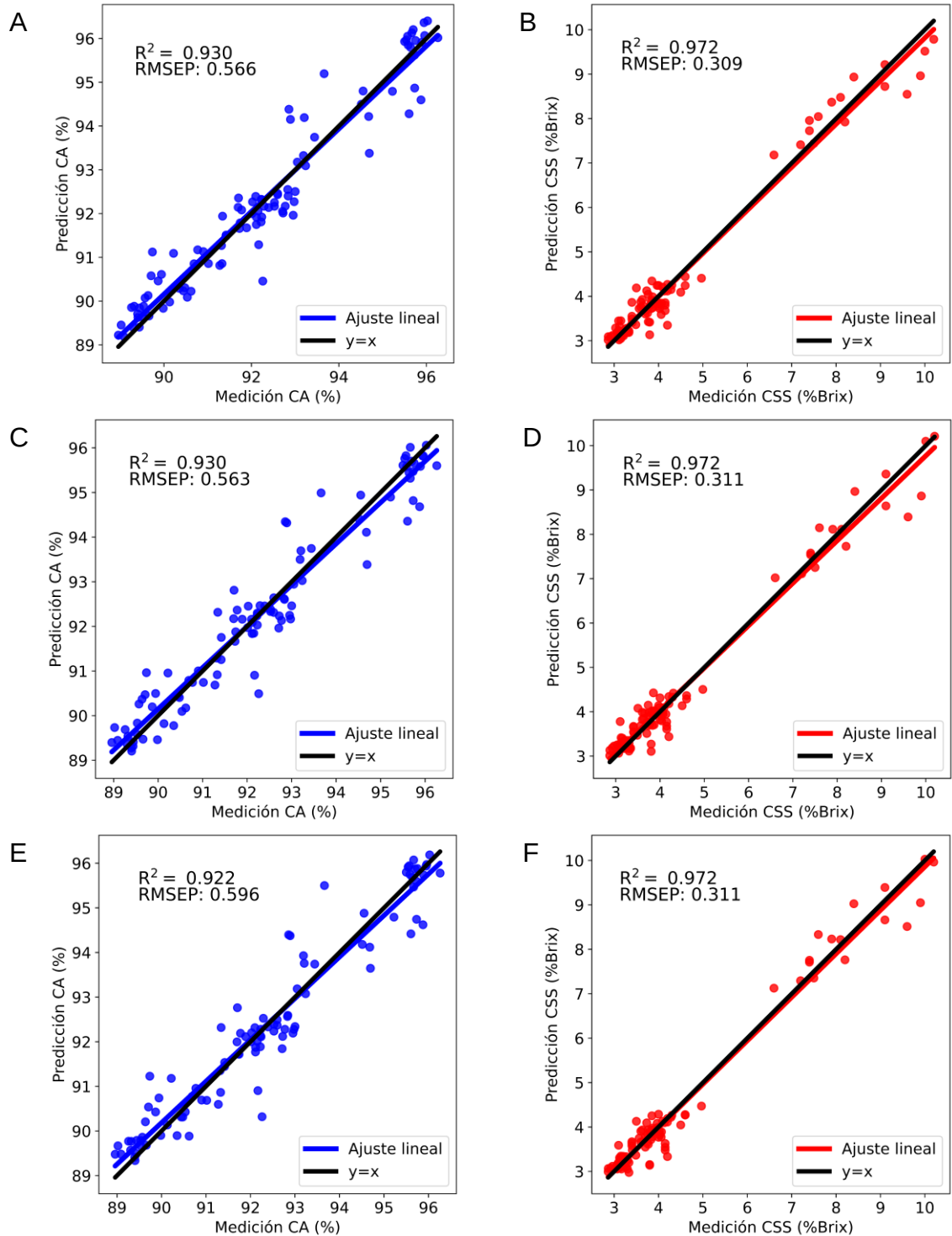


Figura 7. Mejores rendimientos de modelos de DL para predecir simultáneamente el contenido de agua (azul) y contenido de sólidos solubles (rojo) en seis productos de la familia *Cucurbitaceae* con espectros Vis-NIR. Resultados con modelo 1D-CNN de una capa convolucional (A, B), de 3 capas convolucionales (C, D) y de dos ramas (E, F).

Los modelos de salida múltiple se comportaron igual o mejor que los modelos lineales de salida única. Para comprobar si un modelo 1D-CNN de salida múltiple es más conveniente que sus equivalentes de salida única, se optimizaron dos modelos de la red 1 con datos RS para salidas individuales de CA y CSS con el mismo rango de hiperparámetros del Cuadro 1. El proceso demoró un tiempo similar para ambos modelos, 11h 10min para CA y 11h 20min para SSC. Los hiperparámetros óptimos se muestran en el Cuadro 6 y con ellos se reentrenó cada modelo y se aplicó al conjunto de prueba, los resultados finales para CA fueron un RMSEP=0.568 y $R^2=0.929$, mientras que para CSS fueron RMSEP=0.30 y $R^2=0.974$. En el caso de CA el rendimiento fue prácticamente el mismo en comparación con el modelo de salida múltiple y para SSC hubo una ligera mejora en el error. Esta mejora también se observó en el estudio de Mishra & Passos (2022) para predecir CA y SSC en peras con modelos 1D-CNN de salida múltiple y única. Si bien se podría esperar una mejora para ambos atributos con redes de salida única, especialmente si uno es más difícil de modelar que otro, en este caso el RMSEP y R^2 más bajos para CA se pueden considerar un límite para su modelación con los espectros y el método utilizado.

Considerando que para optimizar dos modelos se requiere el doble de tiempo que para optimizar uno, no hubo una ventaja práctica en la optimización de modelos de salida única en comparación con modelos de salida múltiple.

Cuadro 6. Hiperparámetros óptimos para la red de una capa convolucional de salida única para predicción de CA y CSS.

Hiperparámetro o arquitectura	Valor para CA	Valor para CSS
Tamaño de filtro	23	39
Paso de filtro	6	10
Regularización L2	0	0
Tamaño de Batch	64	64
Tasa de deserción por capa	[0.1, 0.305, 0.075]	[0.03, 0.495, 0.025]
Tasa de aprendizaje inicial	0.0155	0.02

Como nota final, mientras que el preprocesamiento es un requisito previo para aplicar los espectros en el modelado convencional, en el modelado con DL las capas convolucionales lo realizaron automáticamente. Dada la similitud y mejora de los resultados de esta investigación con respecto a estudios anteriores, se puede inferir que las capas convolucionales corrigieron los espectros de manera similar a la corrección de la señal ortogonal y en ocasiones de mejor manera. La importancia no solo radica en el ahorro de pasos para el modelado espectral sino que además, se evita elegir el preprocesado erróneo y empeorar el rendimiento como descubrió Mishra, Rutledge, et al. (2021).

3.5 Conclusiones

Los resultados de esta investigación demostraron que el DL es útil para la predicción simultánea de dos atributos de calidad interna en seis frutas diferentes de una familia vegetal, utilizando tan solo 50 espectros Vis-NIR de cada fruta y una técnica de aumento de datos. La optimización automática de hiperparámetros permitió encontrar una estructura 1D-CNN simple pero eficiente en la modelación multiproducto-multisalida, haciendo más atractivo el uso del DL para el modelado espectral, pero se debe vigilar el espacio de búsqueda de hiperparámetros para evitar la tendencia al sobreajuste. Los modelos más complejos requirieron más tiempo de optimización y presentaron mayor tendencia al sobreajuste por lo que su capacidad de generalización se vio comprometida, asimismo, un modelo de salida múltiple tuvo un rendimiento similar a dos modelos de salida única, con la ventaja de requerir la mitad del tiempo de optimización.

El rendimiento del modelado de extremo a extremo tuvo una precisión ligeramente mejor al modelado clásico con PLSR en quimiometría, aunque la mejora puede no ser de gran relevancia, en la práctica siempre será conveniente usar el modelo con la mejor precisión y debe considerarse que se requirió de la mínima interferencia humana. Por lo tanto, se comprobó que el DL puede facilitar el desarrollo y mejorar el rendimiento de modelos globales a costa de poder computacional, el cual puede ser compensado por el rápido desarrollo

tecnológico de nuestra época, contribuyendo al respaldo de la adopción de la técnica Vis-NIR para la evaluación de calidad interna en frutas en todo el proceso postcosecha. Se anima a otros investigadores a aplicar las técnicas DL en nuevos problemas de regresión y clasificación espectral, especialmente añadiendo cada vez más frutas o verduras y más atributos de interés como acidez titulable y firmeza en un mismo modelo global.

3.6 Literatura citada

- Anderson, N. T., & Walsh, K. B. (2022). Review: The evolution of chemometrics coupled with near infrared spectroscopy for fruit quality evaluation. *Journal of Near Infrared Spectroscopy*, 30(1). <https://doi.org/https://doi.org/10.1177/09670335211057235>
- Anderson, N. T., Walsh, K. B., Flynn, J. R., & Walsh, J. P. (2021). Achieving robustness across season , location and cultivar for a NIRS model for intact mango fruit dry matter content. II. Local PLS and nonlinear models. *Postharvest Biology and Technology*, 171, 111358. <https://doi.org/10.1016/j.postharvbio.2020.111358>
- Anderson, N. T., Walsh, K. B., Subedi, P. P., & Hayes, C. H. (2020). Achieving robustness across season, location and cultivar for a NIRS model for intact mango fruit dry matter content. *Postharvest Biology and Technology*, 168, 111202. <https://doi.org/10.1016/j.postharvbio.2020.111202>
- Arendse, E., Fawole, O. A., Magwaza, L. S., & Opara, U. L. (2017). Non-destructive prediction of internal and external quality attributes of fruit with thick rind: A review. *Journal of Food Engineering*, 217, 11–23. <https://doi.org/10.1016/j.jfoodeng.2017.08.009>
- Bergstra, J., Bardenet, R., Bengio, Y., & Kégl, B. (2011). Algorithms for hyper-parameter optimization. En *Adv. Neural Inf. Process. Syst.* (pp. 1–9).
- Bjerrum, E. J., Glahder, M., & Skov, T. (2017). Data Augmentation of Spectral Data for Convolutional Neural Network (CNN) Based Deep Chemometrics. *arXiv:1710.01927 [cs.LG]*, 1–10. Recuperado de <https://arxiv.org/abs/1710.01927>
- Bobelyn, E., Serban, A., Nicu, M., Lammertyn, J., Nicolai, B. M., & Saeys, W. (2010). Postharvest quality of apple predicted by NIR-spectroscopy: Study of the effect of biological variability on spectra and model performance. *Postharvest Biology and Technology*, 55, 133–143. <https://doi.org/10.1016/j.postharvbio.2009.09.006>
- Cui, C., & Fearn, T. (2018). Modern practical convolutional neural networks for multivariate regression: applications to NIR calibration. *Chemometrics and Intelligent Laboratory Systems*, 182, 9–20. <https://doi.org/10.1016/j.chemolab.2018.07.008>

- Ferrari, C., Foca, G., Calvini, R., & Ulrici, A. (2015). Fast exploration and classification of large hyperspectral image datasets for early bruise detection on apples. *Chemometrics and Intelligent Laboratory Systems*, 146, 108–119. <https://doi.org/http://dx.doi.org/10.1016/j.chemolab.2015.05.016>
- Ferreira, R. D. A., Teixeira, G., & Peternelli, L. A. (2022). Kennard-Stone method outperforms the Random Sampling in the selection of calibration samples in SNPs and NIR data. *Ciencia Rural*, 52(5), e20201072.
- GoodFellow, I., Bengio, J., & Courville, A. (2016). *Deep Learning*. London, England: MIT Press.
- Hadiwijaya, Y., Kusumiyati, K., & Munawar, A. A. (2020). Penerapan Teknologi Visible-Near Infrared Spectroscopy untuk Prediksi Cepat dan Simultan Kadar Air Buah Melon (Cucumis melo L.) Golden. *Agroteknika*, 3(2), 67–74. <https://doi.org/10.32530/agroteknika.v3i2.83>
- Ito, H., Morimoto, S., & Yamauchi, R. (2001). Potential of near infrared spectroscopy for non-destructive estimation of soluble solids in growing melons. *Acta Horticulturae*, 566, 483–486. <https://doi.org/10.17660/ActaHortic.2002.588.57>
- Jha, S. N., Jaiswal, P., Narsaiah, K., Gupta, M., Bhardwaj, R., & Singh, A. K. (2012). Non-destructive prediction of sweetness of intact mango using near infrared spectroscopy. *Scientia Horticulturae*, 138(2012), 171–175. <https://doi.org/10.1016/j.scienta.2012.02.031>
- Kavdir, I., Lu, R., Ariana, D., & Ngouajio, M. (2007). Visible and near-infrared spectroscopy for nondestructive quality assessment of pickling cucumbers. *Poster Proceedings of the Sixth European Conference on Precision Agriculture (6ECPA)*, 44, 165–174. <https://doi.org/10.1016/j.postharvbio.2006.09.002>
- Kennard, R. W., & Stone, L. A. (1969). Computer aided design of experiments. *Technometrics*, 11(1), 137–148.
- Kusumiyati, Hadiwijaya, Y., Putri, I. E., & Munawar, A. A. (2021a). Multi-product calibration model for soluble solids and water content quantification in Cucurbitaceae family, using visible/near-infrared spectroscopy. *Heliyon*, 7(8), e07677. <https://doi.org/10.1016/j.heliyon.2021.e07677>
- Kusumiyati, K., Hadiwijaya, Y., Putri, I. E., & Munawar, A. A. (2021b). Dataset of Vis-NIR spectra for intact cucurbitaceae fruits. *En Mendeley Data*, V1. <https://doi.org/doi:10.17632/k55b8mvs84.1>
- Kusumiyati, K., Hadiwijaya, Y., Putri, I. E., & Munawar, A. A. (2021c). Enhanced visible/near-infrared spectroscopic data for prediction of quality attributes in Cucurbitaceae commodities. *Data in Brief*, 39, 107458. <https://doi.org/https://doi.org/10.1016/j.dib.2021.107458>
- Lakshmi, S., Pandey, A. K., Ravi, N., Chauhan, O. P., Gopalan, N., & Sharma, R. K. (2017). Non-destructive quality monitoring of fresh fruits and vegetables. *Defense Life Science Journal*, 2(2), 103–110.

<https://doi.org/10.14429/dlsj.2.11379>

- Lecun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521, 436–444. <https://doi.org/10.1038/nature14539>
- Li, L., Jamieson, K., DeSalvo, G., Rostamizadeh, A., & Talwalkar, A. (2017). Hyperband: bandit-based configuration evaluation for hyperparameter optimization. En *ICLR* (pp. 1–15).
- Li, M., Han, D., & Liu, W. (2019). Non-destructive measurement of soluble solids content of three melon cultivars using portable visible/near infrared spectroscopy. *Biosystems Engineering*, 188, 31–39. <https://doi.org/10.1016/j.biosystemseng.2019.10.003>
- Masithoh, R. E., Lohumi, S., Yoon, W.-S., Amanah, H. Z., & Cho, B. (2020). Development of multi-product calibration models of various root and tuber powders by fourier transform near infra-red (FT-NIR) spectroscopy for the quanti fi cation of polysaccharide contents. *Heliyon*, 6, e05099. <https://doi.org/10.1016/j.heliyon.2020.e05099>
- Mishra, P., Biancolillo, A., Roger, J. M., Marini, F., & Rutledge, D. N. (2020). New data preprocessing trends based on ensemble of multiple preprocessing techniques. *Trends in Analytical Chemistry*, 132, 116045. <https://doi.org/10.1016/j.trac.2020.116045>
- Mishra, P., & Passos, D. (2021a). A synergistic use of chemometrics and deep learning improved the predictive performance of near-infrared spectroscopy models for dry matter prediction in mango fruit. *Chemometrics and Intelligent Laboratory Systems*, 212, 104287. <https://doi.org/10.1016/j.chemolab.2021.104287>
- Mishra, P., & Passos, D. (2021b). Deep chemometrics: Validation and transfer of a global deep near-infrared fruit model to use it on a new portable instrument. *Journal of Chemometrics*, 35(10), 1–12. <https://doi.org/10.1002/cem.3367>
- Mishra, P., & Passos, D. (2021c). Realizing transfer learning for updating deep learning models of spectral data to be used in a new scenario. *Chemometrics and Intelligent Laboratory Systems*, 104248. <https://doi.org/https://doi.org/10.1016/j.chemolab.2021.104283>
- Mishra, P., & Passos, D. (2022). Multi-output 1-dimensional convolutional neural networks for simultaneous prediction of different traits of fruit based on near-infrared spectroscopy. *Postharvest Biology and Technology*, 183, 111741. <https://doi.org/10.1016/j.postharvbio.2021.111741>
- Mishra, P., Rutledge, D. N., Roger, J. M., Wali, K., & Khan, H. A. (2021). Chemometric pre-processing can negatively affect the performance of near-infrared spectroscopy models for fruit quality prediction. *Talanta*, 229, 122303. <https://doi.org/10.1016/j.talanta.2021.122303>
- Murray, I. (2004). Scattered information: philosophy and practice of near infrared spectroscopy. En A. M. C. Davis & A. Garrido-Varo (Eds.), *Near Infrared Spectroscopy: Proceedings of the 11th International Conference* (p. 12).

- Cordoba, Spain: NIR Publications. Recuperado de www.nirpublications.com
- Nallan-Chakravartula, S. S., Moschetti, R., Bedini, G., Nardella, M., & Massantini, R. (2022). Use of convolutional neural network (CNN) combined with FT-NIR spectroscopy to predict food adulteration: A case study on coffee. *Food Control*, 135, 108816. <https://doi.org/10.1016/j.foodcont.2022.108816>
- Padarian, J., Minasny, B., & Mcbratney, A. B. (2018). Using deep learning to predict soil properties from regional spectral data. *Geoderma Regional*, 15, e00198. <https://doi.org/10.1016/j.geodrs.2018.e00198>
- Passos, D., & Mishra, P. (2022). A tutorial on automatic hyperparameter tuning of deep spectral modelling for regression and classification tasks. *Chemometrics and Intelligent Laboratory Systems*, 223, 104520. <https://doi.org/10.1016/j.chemolab.2022.104520>
- Pathmanaban, P., Gnanavel, B. K., & Anandan, S. S. (2019). Recent application of imaging techniques for fruit quality assessment. *Trends in Food Science & Technology*, 94, 32–42. <https://doi.org/10.1016/j.tifs.2019.10.004>
- Rambo, M. K. D., Ferreira, M. M. C., & Amorim, E. P. (2015). Multi-product calibration models using NIR spectroscopy. *Chemometrics and Intelligent Laboratory Systems*, 151, 108–114. <https://doi.org/10.1016/j.chemolab.2015.12.013>
- Saeys, W., Nguyen Do Trong, N., Van-Beers, R., & Nicolaï, B. M. (2019). Multivariate calibration of spectroscopic sensors for postharvest quality evaluation: A review. *Postharvest Biology and Technology*, 158, 110981. <https://doi.org/10.1016/j.postharvbio.2019.110981>
- Szegedy, C., Liu, W., Jia, Y., Sermanet, P., Reed, S., Anguelov, D., ... Rabinovich, A. (2015). Going deeper with convolutions. En *IEEE 2015 Conference on Computer Vision and Pattern REcognition (CVPR)* (pp. 1–9). IEEE. <https://doi.org/https://doi.org/10.1109/CVPR.2015.7298594>
- Teh, S. L., Coggins, J. L., Kostick, S. A., & Evans, K. M. (2020). Location, year, and tree age impact NIR-based postharvest prediction of dry matter concentration for 58 apple accessions. *Postharvest Biology and Technology*, 166, 111125. <https://doi.org/10.1016/j.postharvbio.2020.111125>
- Torniainen, J., Afara, I. O., Prakash, M., Sarin, J. K., Stenroth, L., & Töyräs, J. (2020). Open-source python module for automated preprocessing of Near Infrared Spectroscopic data. *Analytica Chimica Acta*, 1108, 1–9. <https://doi.org/10.1016/j.aca.2020.02.030>
- Vidal, G. I. (2019). *Fenotipado rápido de germoplasma de tomate por características de calidad organoléptica para su uso en programas de mejora*. Universitat Jaume I.
- Walsh, K. B., Mcglone, V. A., & Han, D. H. (2020). The uses of near infrared spectroscopy in postharvest decision support: A review. *Postharvest Biology and Technology*, 163, 111139. <https://doi.org/10.1016/j.postharvbio.2020.111139>

- Wang, H., Peng, J., Xie, C., Bao, Y., & He, Y. (2015). Fruit Quality Evaluation Using Spectroscopy Technology: A Review. *Sensors*, *15*(5), 11889–11927. <https://doi.org/10.3390/s150511889>
- Wang, J., Wang, J., Chen, Z., & Han, D. (2017). Development of multi-cultivar models for predicting the soluble solid content and firmness of European pear (*Pyrus communis* L.) using portable vis–NIR spectroscopy. *Postharvest Biology and Technology*, *129*, 143–151. <https://doi.org/10.1016/j.postharvbio.2017.03.012>
- Zhang, X., Lin, T., Xu, J., Luo, X., & Ying, Y. (2019). DeepSpectra: An end-to-end deep learning approach for quantitative spectral analysis. *Analytica Chimica Acta*, *1058*, 48–57. <https://doi.org/10.1016/j.aca.2019.01.002>
- Zhang, Y., Nock, J. F., Shoffe, Y. Al, & Watkins, C. B. (2019). Non-destructive prediction of soluble solids and dry matter contents in eight apple cultivars using near-infrared spectroscopy. *Postharvest Biology and Technology*, *151*, 111–118. <https://doi.org/10.1016/j.postharvbio.2019.01.009>