



Clubes de Ciencia
México

Rodolfo Ferro

ferro@cimat.mx

<https://rodolfoferro.xyz>

Clubes de Ciencia México 2025

"Satélites y Neuronas: Explorando lo Invisible con AI"

Julio, 2025

Introducción a CV

Día 4

Rodolfo Ferro (ferro@ciimat.mx)

- › Sr. SWE (Data Engineer) @ Bisonic México
- › Tutor de Ciencia de Datos en Código Facilito y en el Diplomado en Ciencia de Datos de la ENES UNAM León
- › Profesor de AI & Coordinador del TomorrowLab @ EdgeHub School of Innovation (Ags.)
- › **Formación:** BMath @ UG, CSysEng @ UVEG, StatMethods Specialist @ CIMAT
- › **Experiencia:** ML Engineer @ Vindoo.ai (España), Sherpa Digital en IA @ Microsoft México, AI Research Assistant @ CIMAT & AI Research Intern @ Harvard.



ferro@ciimat.mx

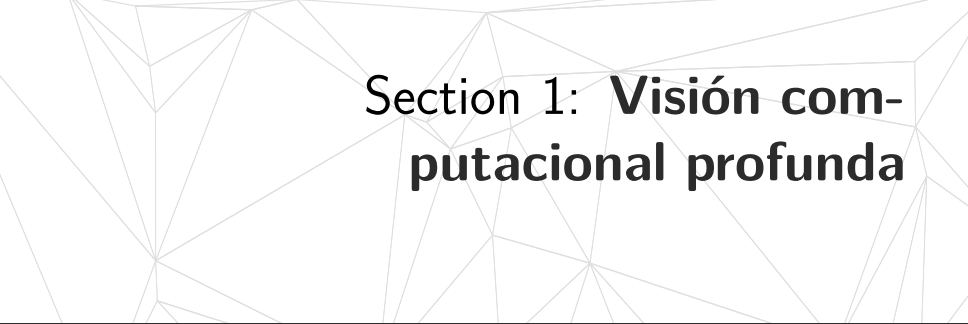


Tabla de contenidos

1 Visión computacional profunda

1 Visión computacional profunda

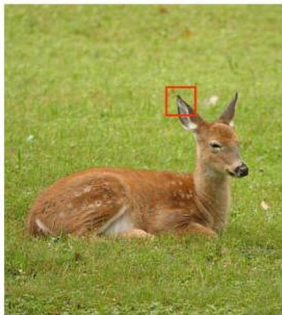
- Introducción a imágenes
- Convoluciones & Pooling
- Redes neuronales convolucionales
- Clasificadores de imágenes (LeNet, VGG16, etc.)
- Trabajos relacionados y avances recientes



Section 1: **Visión computacional profunda**

¿Qué es una imagen?

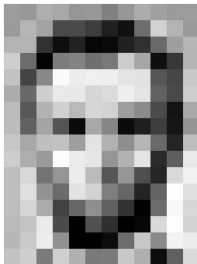
Visión computacional profunda



ferro@ciimat.mx

¿Qué es una imagen?

Visión computacional profunda

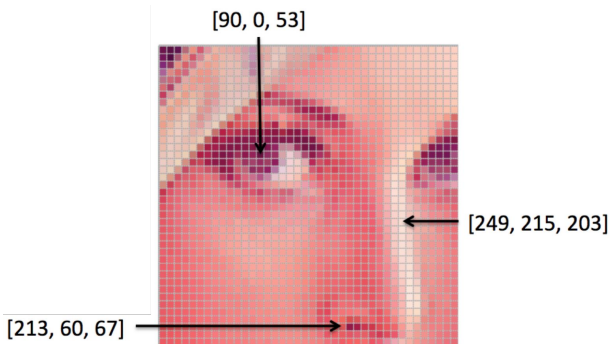


157	153	174	168	150	152	129	161	172	161	155	156
155	182	163	74	75	62	33	17	110	210	180	154
180	180	50	14	54	5	10	33	48	105	159	181
206	109	5	124	131	111	120	204	165	15	55	180
194	68	137	251	237	239	239	228	227	87	71	201
172	105	207	233	233	214	220	239	228	98	74	206
188	66	179	209	185	215	211	158	139	75	20	169
189	97	165	84	10	168	134	11	31	62	22	148
199	168	191	193	158	227	178	143	182	104	36	190
205	174	155	252	236	231	149	178	228	43	95	234
190	216	116	149	236	187	86	150	79	38	218	241
190	224	147	108	227	210	127	102	36	101	255	224
190	214	173	66	103	143	95	50	2	109	249	215
187	196	235	75	1	81	47	0	6	217	255	211
183	202	237	145	0	0	12	108	200	138	243	236
195	206	123	207	177	121	123	200	175	13	96	218

157	153	174	168	150	152	129	161	172	161	155	156
155	182	163	74	75	62	33	17	110	210	180	154
180	180	50	14	34	6	10	33	48	105	159	181
206	109	5	124	131	111	120	204	166	15	55	180
194	68	137	251	237	239	239	228	227	87	71	201
172	105	207	233	233	214	220	239	228	98	74	206
188	66	179	209	185	215	211	158	139	75	20	169
189	97	165	84	10	168	134	11	31	62	22	148
199	168	191	193	158	227	178	143	182	104	36	190
205	174	155	252	236	231	149	178	228	43	95	234
190	216	116	149	236	187	86	150	79	38	218	241
190	224	147	108	227	210	127	102	36	101	255	224
190	214	173	66	100	143	95	50	2	109	249	215
187	196	235	75	1	81	47	0	6	217	255	211
183	202	237	145	0	0	12	108	200	138	243	236
195	206	123	207	177	121	123	200	175	13	96	218

¿Qué es una imagen?

Visión computacional profunda

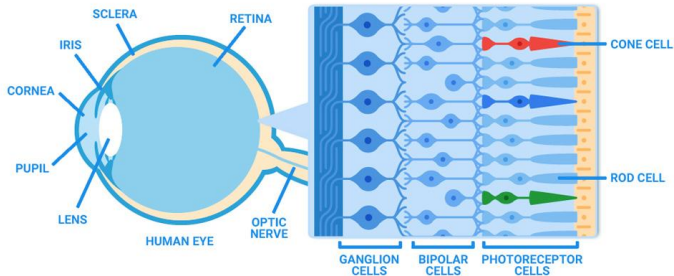


ferro@ciimat.mx

- Una imagen es un arreglo de píxeles, la cual puede tener 1 o más canales de color. Usualmente:
 - » 1 canal de color → Escala de grises
 - » 3 canales de color → Escala RGB
 - » 4 canales de color → Escala RGBA
- Un pixel puede ser visto como un objeto 5-dimensional (x, y, r, g, b) .

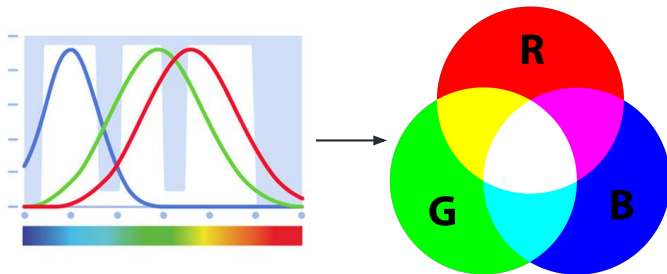
La biología humana

Visión computacional profunda

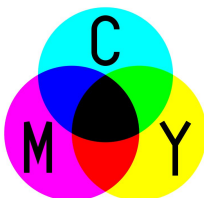
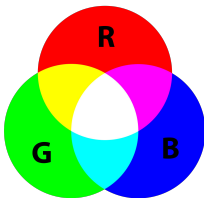


La biología humana

Visión computacional profunda

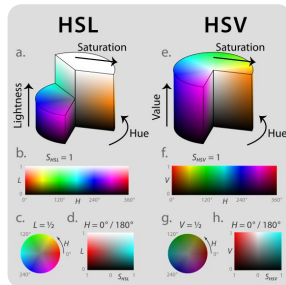


ferro@ciimat.mx



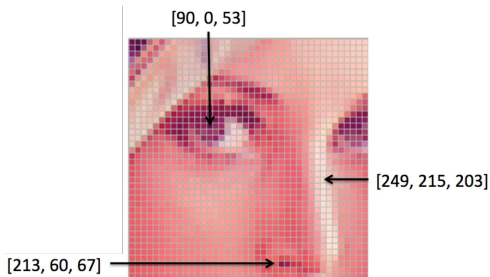
Espacios de color

Visión computacional profunda



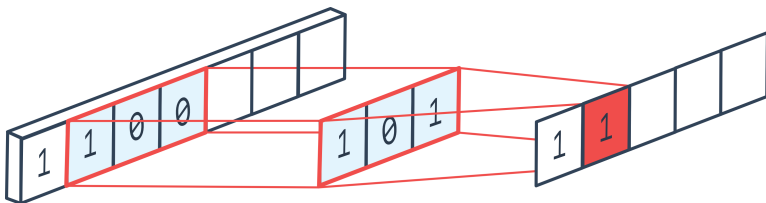
ferro@ciimat.mx

Ejercicio: Introducción a imágenes



¿Qué es una convolución?

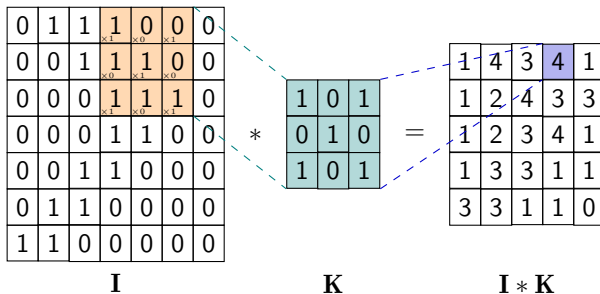
Visión computacional profunda



ferro@ciimat.mx

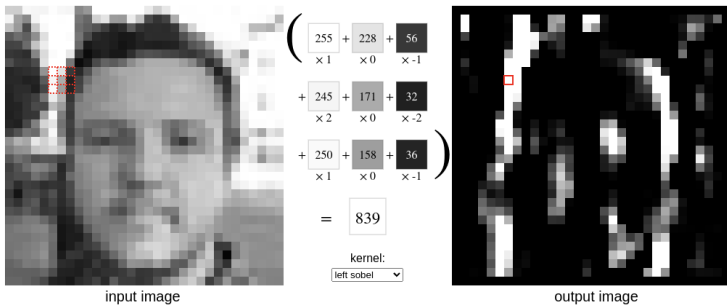
¿Qué es una convolución?

Visión computacional profunda



¿Qué es una convolución?

Visión computacional profunda



ferro@ciimat.mx

Pooling

Visión computacional profunda

Pool size Stride

←						→
-4	0	-2	4	1		
3	1	0	2	1		
1	0	1	1	1		
4	6	5	1	0		
-1	2	0	0	0		

Features

Max Pooling

3	4
6	5

Min Pooling

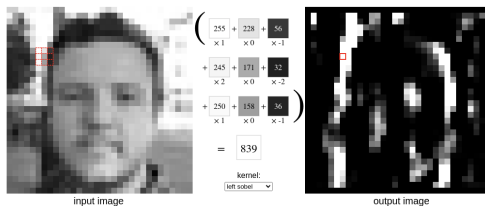
-4	-2
-1	0

Average Pooling

0	1
2	1

Output

Ejercicio: Convoluciones & Pooling



Redes neuronales convolucionales

Visión computacional profunda



Figure: Yann LeCun

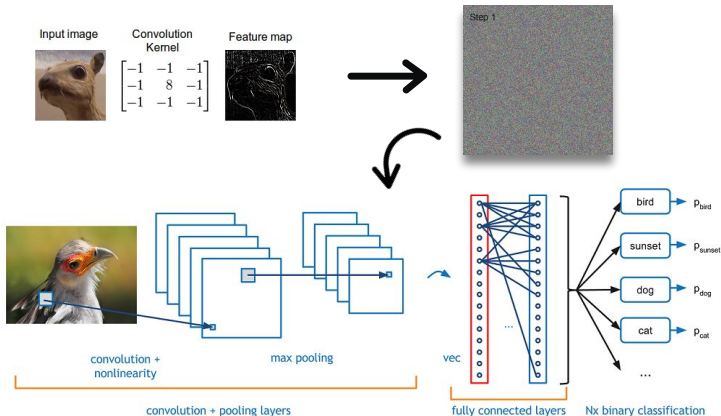
ferro@ciimat.mx



Clubes de Ciencia
México

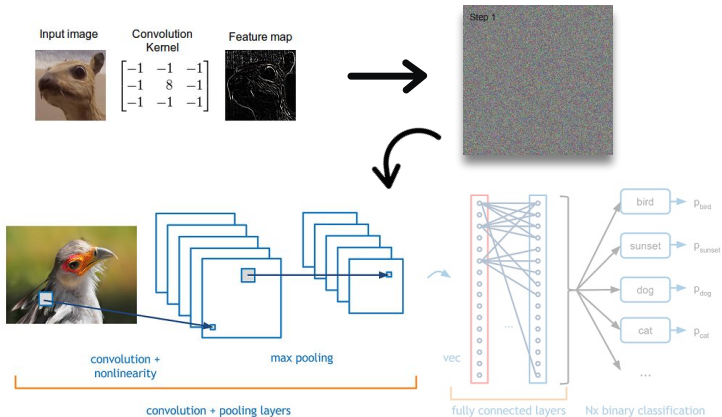
Redes neuronales convolucionales

Visión computacional profunda



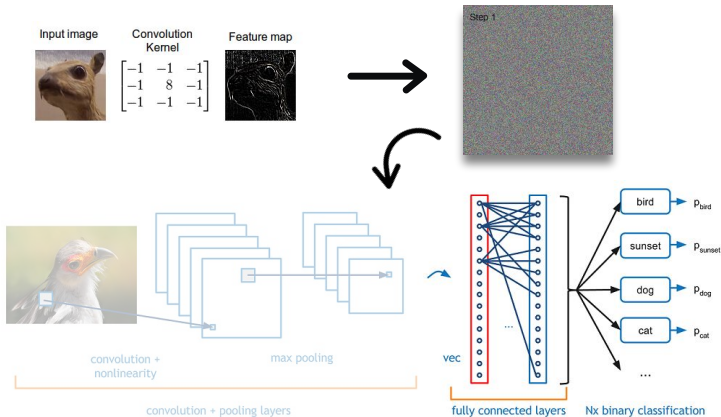
Redes neuronales convolucionales

Visión computacional profunda



Redes neuronales convolucionales

Visión computacional profunda



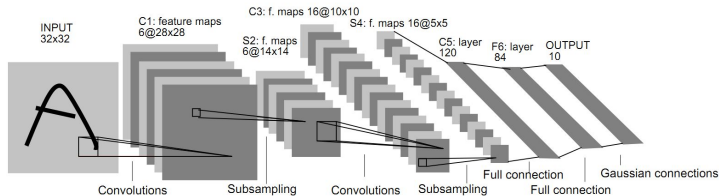
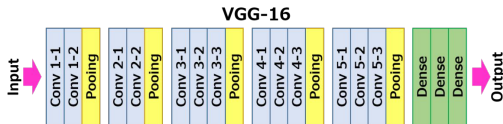
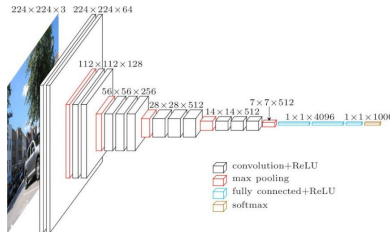


Fig. 2. Architecture of LeNet-5, a Convolutional Neural Network, here for digits recognition. Each plane is a feature map, i.e. a set of units whose weights are constrained to be identical.

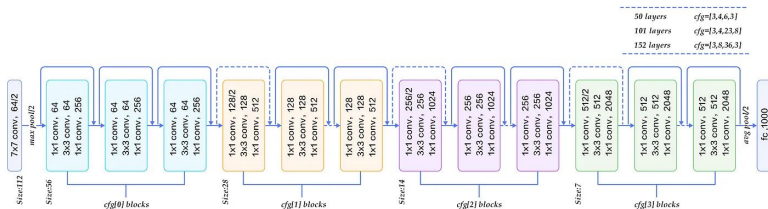
1998



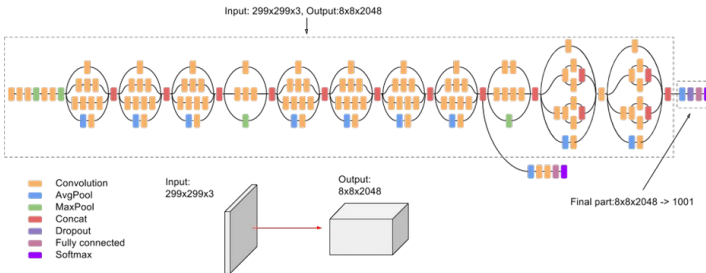
2014

Resnet50

Visión computacional profunda



2015



2015

Ejercicio: Redes neuronales convolucionales

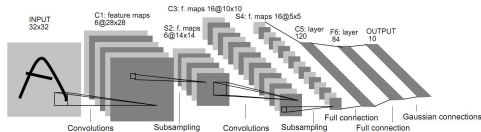


Fig. 2. Architecture of LeNet-5, a Convolutional Neural Network, here for digits recognition. Each plane is a feature map, i.e. a set of units whose weights are constrained to be identical.

1998

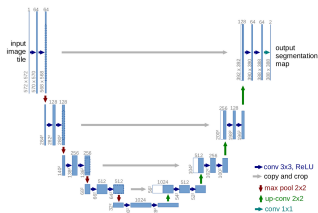
- › Tutorial 1: Image Filtering
- › Image Kernels Explained Visually by Victor Powell
- › Parameterized Pooling Layers by Hao Hao Tan
- › TensorFlow Tutorials: Vision

Trabajos relacionados y avances recientes

Visión computacional profunda

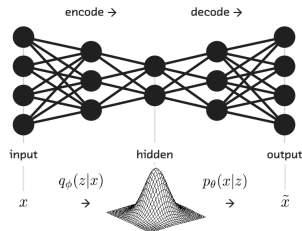
Han habido varios trabajos de investigación y avances recientes que han contribuido al desarrollo de nuevas arquitecturas, técnicas de entrenamiento mejoradas y aplicaciones emergentes.

- **UNet:** Es ampliamente utilizada en el campo de la segmentación de imágenes, pero también se ha aplicado con éxito en tareas de denoising.

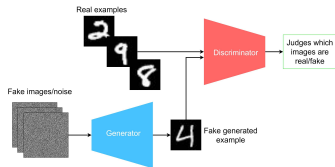


➤ Variational Autoencoders (VAEs):

Los VAEs son una variante de los autoencoders que se utilizan para el aprendizaje de distribuciones latentes. Han demostrado ser efectivos en el denoising de imágenes al aprender representaciones latentes que siguen una distribución probabilística, lo que permite una generación más controlada y realista de imágenes limpias.



- **Generative Adversarial Networks (GANs):** Estos modelos aprovechan la capacidad de los GANs para generar imágenes realistas y para aprender representaciones latentes eficientes. Los GANs han demostrado ser efectivos en el denoising y la generación de imágenes de alta calidad, entre otros.



Tareas en el campo de visión artificial

Visión computacional profunda

- **Clasificación de imágenes:** La tarea de clasificación de imágenes implica asignar una etiqueta o categoría a una imagen de entrada. Esto implica entrenar un modelo para reconocer y distinguir diferentes objetos, personas o escenas en una imagen.
- **Detección de objetos:** La detección de objetos implica localizar y clasificar múltiples objetos en una imagen. El objetivo es detectar la presencia y la ubicación de objetos específicos en una escena, a menudo utilizando cuadros delimitadores para delinear las regiones donde se encuentran los objetos.
- **Denoising o reconstrucción de imágenes:** Consiste en eliminar o reducir el ruido presente en una imagen, obteniendo una versión más limpia y clara. Esta tarea es relevante en áreas como la fotografía, la medicina y la seguridad.

Tareas en el campo de visión artificial

Visión computacional profunda

- **Segmentación semántica:** La segmentación semántica implica asignar una etiqueta a cada píxel de una imagen para identificar y delimitar las diferentes regiones o objetos presentes. El objetivo es comprender la estructura y el contenido de una imagen a nivel de píxel.
- **Detección de rostros:** La detección de rostros es una tarea específica de la visión artificial que implica detectar y localizar los rostros en una imagen. Es ampliamente utilizado en aplicaciones de reconocimiento facial, análisis de emociones y sistemas de seguridad.
- **Reconocimiento y verificación facial:** El reconocimiento facial se refiere a la tarea de identificar y reconocer a una persona específica a partir de una imagen o secuencia de imágenes. La verificación facial se enfoca en verificar si una imagen de rostro coincide con una identidad específica.

Tareas en el campo de visión artificial

Visión computacional profunda

- **Estimación de pose:** La estimación de pose se refiere a la tarea de determinar la posición y orientación de un objeto o persona en una imagen. Esto implica detectar y rastrear las articulaciones o puntos clave en una imagen para comprender la postura y el movimiento.
- **Estimación de profundidad:** La estimación de profundidad implica inferir la información de la distancia o la profundidad de los objetos en una imagen. Es útil en aplicaciones de realidad virtual, conducción autónoma y sistemas de navegación.
- **Super-resolución:** La super-resolución se refiere a aumentar la resolución o la calidad de una imagen de baja resolución. El objetivo es generar una versión de alta resolución que capture más detalles y claridad.

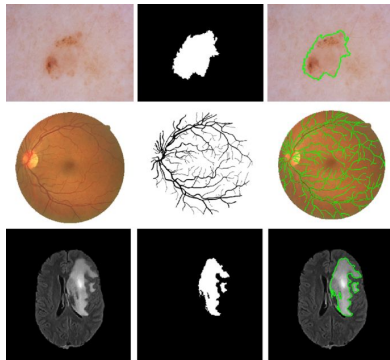
- › From Autoencoder to Beta-VAE
- › Building Autoencoders in Keras
- › Autoencoders: explicación y tutorial en Python
- › Autoencoder For Denoising Images
- › Medical image denoising using convolutional denoising autoencoders

Segmentación de imágenes médicas

Olaf
Ronneberger
et al.,
2015

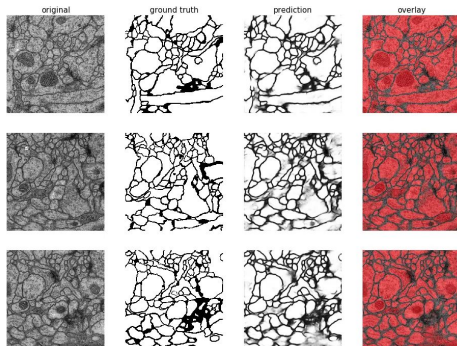


Segmentación de imágenes



<https://www.nature.com/articles/s41598-021-89686-3>

Segmentación con U-Net



<https://medium.com/@venkateshtata9/semantic-segmentation-on-medical-images-3ba8264cda5e>

Segmentación con U-Net

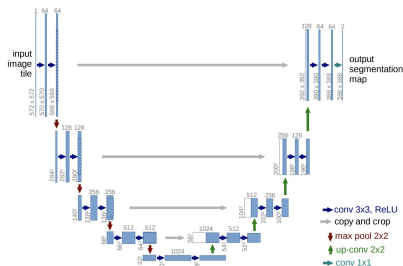
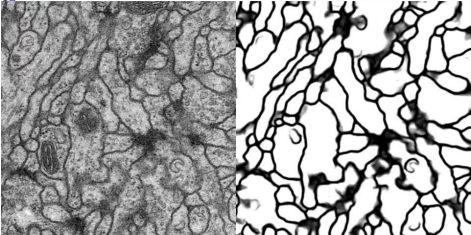


Fig. 1. U-net architecture (example for 32x32 pixels in the lowest resolution). Each blue box corresponds to a multi-channel feature map. The number of channels is denoted on top of the box. The x-y-size is provided at the lower left edge of the box. White boxes represent copied feature maps. The arrows denote the different operations.

<https://lmb.informatik.uni-freiburg.de/people/ronneber/u-net/>

Segmentación con U-Net

Our Results



The image shows a side-by-side comparison of a grayscale electron micrograph of biological tissue and its segmented version. The left image is the 'Input image', showing a complex, textured pattern of cells. The right image is the segmentation result, where the boundaries of the cells are highlighted in black on a white background. Above the images is a blue header with the text 'Our Results'. Below the images, text provides performance metrics: 'Our result: 0.000353 warping error (New best score at submission march 6th, 2015)', 'Sliding-window CNN: 0.000420', and 'Training time: 10h, Application: 1s per image'. At the bottom, a footer identifies the author as 'Olaf Ronneberger, University of Freiburg, Germany, 22.5.2015' and the slide number as '18'.

Input image

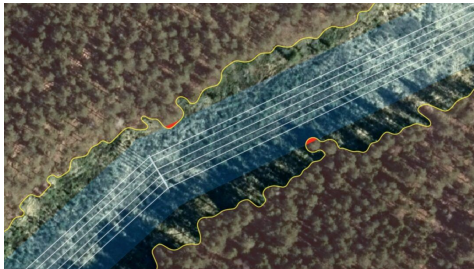
Our result: 0.000353 warping error
(**New best score** at submission march 6th, 2015)
Sliding-window CNN: 0.000420
Training time: 10h, Application: 1s per image

Olaf Ronneberger, University of Freiburg, Germany, 22.5.2015 18

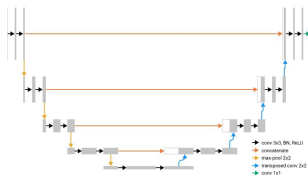
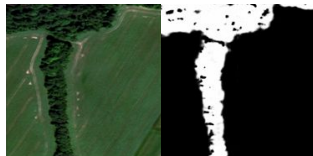
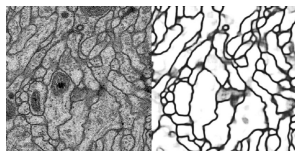
<https://lmb.informatik.uni-freiburg.de/people/ronneber/u-net/>

Caso de estudio

Prevención de incendios



<https://omdena.com/projects/ai-prevent-forest-fires/>

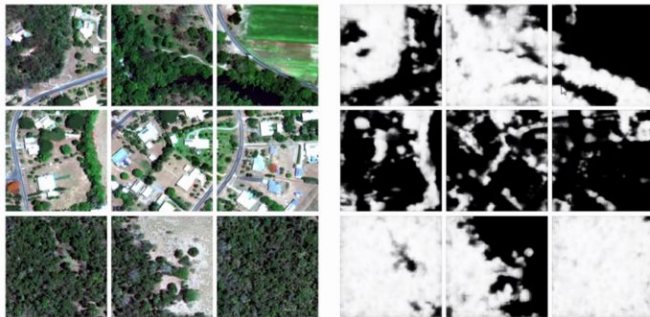


U-Net: Convolutional Networks for Biomedical Image Segmentation.

Prevención de incendios

- **Omdena + Spacept**, una startup sueca
- 36 colaborador@s a nivel global (al menos 3 mexicanos)
- Dataset de 200 imágenes satelitales (Australia)
- Desarrollamos un modelo con 95% de precisión
- Desarrollamos un módulo que preprocesa imágenes
- **iResolvimos el challenge!**

<https://omdena.com/projects/ai-prevent-forest-fires/>



<https://omdena.com/projects/ai-prevent-forest-fires/>

Ejercicio: Entrenamiento de U-Net

