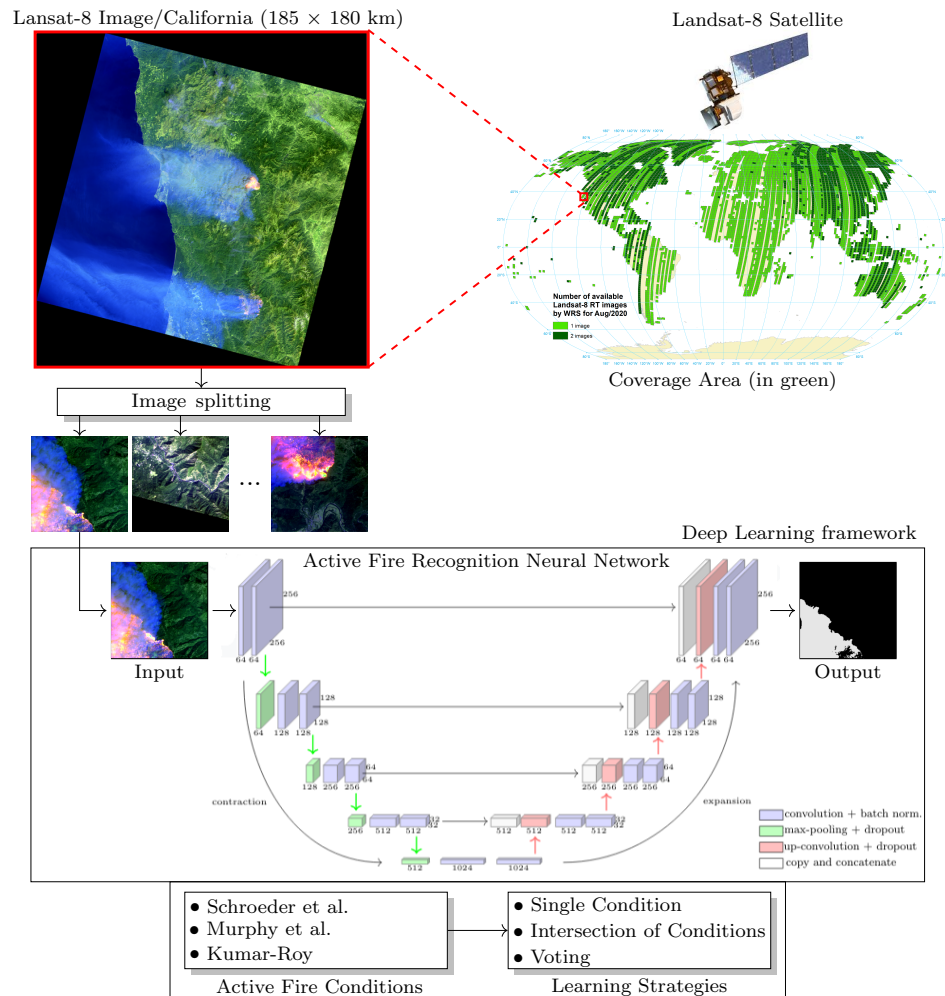


# 1 Graphical Abstract

## 2 Active Fire Detection in Landsat-8 Imagery: a Large-Scale Dataset 3 and a Deep-Learning Study

4 Gabriel Henrique de Almeida Pereira, Andre Minoro Fusioka, Bogdan To-  
5 moyuki Nassu, Rodrigo Minetto,



## 6 Highlights

### 7 **Active Fire Detection in Landsat-8 Imagery: a Large-Scale Dataset** 8 **and a Deep-Learning Study**

9 Gabriel Henrique de Almeida Pereira, Andre Minoro Fusioka, Bogdan To-  
10 moyuki Nassu, Rodrigo Minetto,

- 11     • Proposal of a large-scale dataset for active fire detection in Landsat-8  
12       imagery
- 13     • Deep learning can approximate and improve the results from hand-  
14       crafted algorithms
- 15     • Deep learning networks can approximate combinations of handcrafted  
16       algorithms

17           Active Fire Detection in Landsat-8 Imagery: a  
18           Large-Scale Dataset and a Deep-Learning Study

19           Gabriel Henrique de Almeida Pereira\*, Andre Minoro Fusioka, Bogdan  
20                           Tomoyuki Nassu, Rodrigo Minetto

21                           *Av. Sete de Setembro, 3165, Curitiba, Brazil*

22                           *Department of Informatics - Federal University of Technology - Paraná*  
23

---

24   **Abstract**

Active fire detection in satellite imagery is of critical importance to the management of environmental conservation policies, supporting decision-making and law enforcement. This is a well established field, with many techniques being proposed over the years, usually based on pixel or region-level comparisons involving sensor-specific thresholds and neighborhood statistics. In this paper, we address the problem of active fire detection using deep learning techniques. In recent years, deep learning techniques have been enjoying an enormous success in many fields, but their use for active fire detection is relatively new, with open questions and demand for datasets and architectures for evaluation. This paper addresses these issues by introducing a new large-scale dataset for active fire detection, with over 150,000 image patches (more than 200 GB of data) extracted from Landsat-8 images captured around the world in August and September 2020, containing wildfires in several locations. The dataset was split in two parts, and contains 10-band spectral

---

\*Corresponding author  
Preprint submitted to ISPRS Journal of Photogrammetry and Remote Sensing December 28, 2021  
Email addresses: [pereira.gha@hotmail.com](mailto:pereira.gha@hotmail.com) (Gabriel Henrique de Almeida Pereira), [amfusioka@gmail.com](mailto:amfusioka@gmail.com) (Andre Minoro Fusioka), [btnassu@utfpr.edu.br](mailto:btnassu@utfpr.edu.br) (Bogdan Tomoyuki Nassu), [rminetto@utfpr.edu.br](mailto:rminetto@utfpr.edu.br) (Rodrigo Minetto)

images with associated outputs, produced by three well known handcrafted algorithms for active fire detection in the first part, and manually annotated masks in the second part. We also present a study on how different convolutional neural network architectures can be used to approximate these handcrafted algorithms, and how models trained on automatically segmented patches can be combined to achieve better performance than the original algorithms — with the best combination having 87.2% precision and 92.4% recall on our manually annotated dataset. The proposed dataset, source codes and trained models are available on Github<sup>1</sup>, creating opportunities for further advances in the field.

25 *Keywords:* active fire detection, active fire segmentation, active fire  
26 dataset, convolutional neural network, landsat-8 imagery.

---

## 27 1. Introduction

28 The use of satellite imagery is of critical importance to the management  
29 of environmental conservation policies. Active fire detection is a field of re-  
30 search that extracts relevant knowledge from such images to **support decision-**  
31 **making** (Cardil et al. (2019)). For example, fires can be employed as a means  
32 to clear out an area after trees are cut down, giving place to pastures and  
33 agricultural land, so the presence of active fire can be directly **related to defor-**  
34 **estation** (Portillo-Quintero et al. (2013)) in many important biomes around  
35 the world. Furthermore, there is a direct relation between biomass burn-  
36 ing and changes to the **climate and atmospheric chemistry** (Chuvieco et al.

---

<sup>1</sup><https://github.com/pereira-gha/activefire>

37 (2019)). Active fire detection in satellite imagery is also relevant for other  
38 monitoring tasks, such as damage assessment, prevention and prediction <sup>2</sup>.

39 Research on the field of active fire detection focuses on providing appropri-  
40 ate methods for determining if a given pixel (corresponding to a georeferenced  
41 area) in a multi-spectral image corresponds to an active fire. As the satel-  
42 lites orbit in different altitudes and their systems are equipped with different  
43 sensors, with different wavelengths, usually such methods must be designed  
44 specifically for each system. For example, images from the Landsat-8 satel-  
45 lite — the one considered in this work — are usually processed by series of  
46 conditions such as those proposed by Schroeder et al. (2016), Murphy et al.  
47 (2016) or Kumar and Roy (2018). These conditions are especially tuned for  
48 the Operational Land Imager (OLI) sensor on board Landsat-8, as detailed  
49 in Section 2. Other important sensors and satellites for active fire recognition  
50 are: the MODIS sensor, with 250 m to 1 km of spatial resolution per pixel,  
51 that equips the NASA Terra and Aqua satellites (orbiting 705 km height,  
52 1-2 days of revisit); the AVHRR sensor, with 1 km of spatial resolution per  
53 pixel, on board NOAA-18 (833 km height), NOAA-19 (870 km height) and  
54 METOP-B (827 km height) satellites, the three with a daily revisit time;  
55 the Visible Infrared Imaging Radiometer Suite (VIIRS) sensor, with 375 m  
56 of spatial resolution and on board the joint NASA/NOAA Suomi National  
57 Polar-orbiting Partnership (Suomi NPP) and NOAA-20 satellites (824 km

---

<sup>2</sup><https://effis.jrc.ec.europa.eu/about-effis/technical-background/rapid-damage-assessment>

height), both with daily revisit time; the ABI sensor, that equips the geostationary satellite GOES-16 ( $\approx 36,000$  km height), that is able to take images each 10 minutes with spatial resolution from 500 meters to 2 km. In addition, the Sentinel-2 constellation have also been used for research on this field as they produce images similar to Landsat, with spatial resolution between 10 and 60 meters, with 5 days of revisit.

Studies on automatic methods for detecting active fire regions date back to the 1970s, when the most suitable spectral intervals for detecting forest fires were analyzed (Kondratyev et al. (1972)). Matson and Holben (1987) argued that remote sensing was the only viable alternative to monitor active fire in isolated regions. In their study, they used the AVHRR sensor to detect fire activity in the Amazon area. Other active fire detection algorithms have also been developed for AVHRR (Flannigan and Haar (1986); Lee and Tag (1990)), some of them being used as a basis for the development of algorithms for other satellite sensors. Ji and Stocker (2002) and Giglio et al. (2003) developed algorithms that include contextual analysis for active fire detection for the Visible and Infrared Scanner (VIRS) sensor onboard the Tropical Rainfall Measuring Mission (TRMM) satellite. Kaufman et al. (1998) proposed an approach for active fire detection using the MODIS sensor, which was later improved by Giglio et al. (2003). This method is still used today due to its relevance, and because the MODIS sensor is still active, and has an almost daily global coverage, being able to detect fires with size  $100\text{-}300\text{ m}^2$  (Maier et al. (2013)). Furthermore, the MODIS detection

81 algorithm was improved in many ways (Morisette et al. (2005); Giglio et al.  
82 (2016)) and served a reference for developing equations for other satellites.  
83 As an example, in 2014, Schroeder et al. (2014) proposed an algorithm for  
84 active fire detection for the VIIRS sensor based on the MODIS algorithm.  
85 The active fire detection algorithms mentioned above, in general, are based  
86 on comparisons to fixed threshold in certain bands and statistics from their  
87 surrounding region.

88 In recent years, deep-learning methods (Lecun et al. (2015)) have been  
89 promoting major advances in artificial intelligence, inspiring the design of  
90 novel architectures and strategies for a wide range of applications in remote  
91 sensing, including several related to active fire recognition. That includes ar-  
92 chitectures for tasks such as: distinguish smoke from similar patterns such as  
93 clouds, dust, haze, land and seaside by using images from the Aqua and Terra  
94 satellites (Ba et al. (2019)); super-resolution enhancement for the specific  
95 context of wildfire on images from the Sentinel-2 satellite; use of temporal  
96 information from a same location, before and after wildfires, so as to detect  
97 burned/unburned areas in images of Sentinel-1 or VIIRS data (Ban et al.  
98 (2020), Pinto et al. (2020)); tackling the imbalanced classification problem,  
99 usually present in active fire segmentation, that can negatively impact CNN  
100 performance (Langford et al. (2018)); and the use of generative adversarial  
101 networks (GAN) to synthesize missing or corrupted multispectral optical im-  
102 ages with multitemporal data from satellites (Bermudez et al. (2019)). In  
103 summary, although this is a relatively new field, with no specific architec-

104 ture finding widespread application on remote sensing applications, most of  
105 these works are based on a particular type of deep learning network, known  
106 as Convolutional Neural Network (CNN) (Lecun et al. (2015)), that tries to  
107 discover intricate patterns in massive data by using a sequence of process-  
108 ing layers that allows to learn representations of data with multiple levels of  
109 abstraction.

110 The use of deep learning techniques for active fire recognition is a rela-  
111 tively new field, which still lacks large-scale datasets and architectures for  
112 evaluation. This paper brings several contributions to the field. The first  
113 contribution is that we introduce a new public dataset for active fire recogni-  
114 tion, built from 8,194 Landsat-8 images around the world, covering the period  
115 of August 2020, containing large wildfire events in many areas such as the  
116 Amazon region, Africa, Australia, United States, among others. The dataset  
117 contains 146,214 image patches (around 210 GB uncompressed), including  
118 10-band spectral images and associated outputs produced by three well es-  
119 tablished, handcrafted algorithms for active fire detection (Schroeder et al.  
120 (2016), Murphy et al. (2016) and Kumar and Roy (2018)). The second con-  
121 tribution is a secondary dataset, also public, containing 9,044 image patches  
122 (around 12 GB uncompressed) extracted from 13 Landsat-8 images captured  
123 in September 2020, along with manually annotated fire pixels, which can  
124 be employed for assessing the quality of automatic segmentation compared  
125 to a human specialist. As a third contribution, we present a study on how  
126 convolutional neural networks can be used to approximate the handcrafted

127 fire detection algorithms mentioned above, including the possibility of saving  
128 bandwidth resources by reducing the number of analyzed spectral bands, as  
129 well as means of combining multiple fire detection outputs to produce more  
130 robust detection results. As a fourth contribution, we highlight the compar-  
131 ison between the fire detection algorithms mentioned above and the CNNs  
132 when tested against manual annotations from a human specialist. Finally, as  
133 a fifth contribution, all the source code of the deep learning and handcrafted  
134 algorithms employed for producing this paper are available for free access on  
135 github, all of them coded in Python language, which is a step forward for  
136 other researchers in the field towards a benchmark, opening opportunities for  
137 the study of other Landsat-8 time intervals, or the use of different satellites  
138 — this is a first effort in this direction, as the algorithms we analyzed here,  
139 to the best of our knowledge, do not have open source versions available.

140 The rest of this paper is organized as follows. In section 2, we describe  
141 the proposed datasets, including the procedure for generating them. In sec-  
142 tion 3, we describe the CNN architectures for active fire detection and the  
143 performed experiments. Results are presented and discussed in section 4.  
144 Finally, section 5 presents conclusions and points to future work.

## 145 2. Materials

146  
147 The Landsat program, sponsored by NASA/USGS, provides since 1972  
148 continuous acquisition of high-quality satellite imagery of the Earth, becom-

149 ing a key instrument for the continuous monitoring of the environment. As  
 150 described by Roy et al. (2014), the Landsat-8 satellite, which we consider in  
 151 this research, orbits the Earth at an altitude of 705 km, with data segmented  
 152 into scenes with  $185 \times 180$  km, defined according to the second World-wide  
 153 Reference System (WRS) in a 16-day revisit period. This satellite uses an  
 154 Operational Land Imager (OLI) and Thermal Infrared Sensor (TIRS) sensors  
 155 to acquire eleven channels of multi-spectral data  $\{c_1, \dots, c_{11}\}$ . The images  
 156 are encoded in the TIFF format, with resolution of  $\approx 7,600 \times 7,600$  pixels,  
 157 with 16 bits per pixel per channel. Each multispectral pixel corresponds to  
 158 30 meters of spatial resolution (ground).

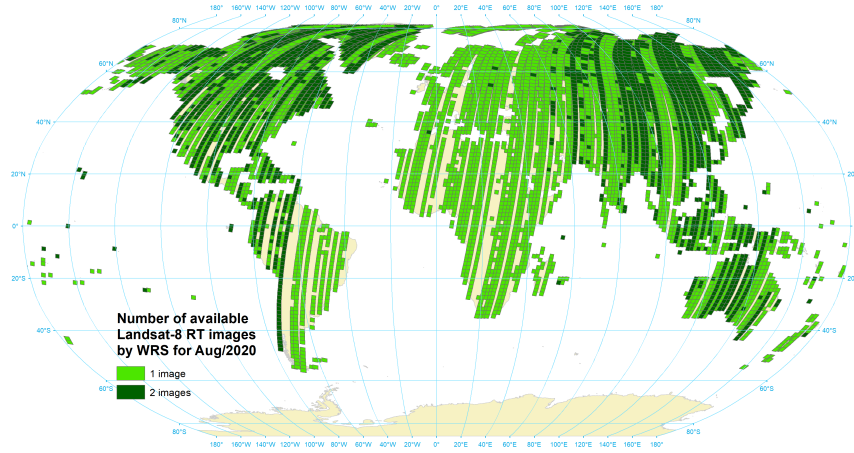


Figure 1: Landsat-8 regions used in our dataset: light green regions have one image for August 2020 and dark green regions have two images. Some regions did not have any real time images available, and are shown in yellow, along with Antarctica, which was not considered in this work. Each rectangle corresponds to an Landsat-8 WRS image scene with  $\approx 7,600 \times 7,600$  pixels covering a  $185 \times 180$  km area.

159 We processed all the Landsat-8 real time images available for August 2020  
160 around the globe, excluding only the Antarctic continent, to a total of 8,194  
161 images and 1.6 TB of data. Figure 1 shows scene locations and the number  
162 of real time images available from WRS.

163 The choice of this particular time slice was not random. Although the  
164 northern and southern hemispheres have different climatic behavior along the  
165 months — especially regarding the seasons, and the occurrence of rain, snow  
166 or droughts — **there is a peak in fire activity around the same time in both**  
167 **hemispheres**. Ferreira et al. (2020) present a study showing the duration of  
168 fire seasons, and how they are spread along the months and across the globe.  
169 They show that August and September are the most critical months, with  
170 fire seasons occurring in all continents, except Antarctica (but including the  
171 Arctic region). In 2020 this situation was not different, with several wildfire  
172 events occurring at the same time in different locations: in mid-August there  
173 were wildfires in Australia, Siberia, in the Parana Delta in Argentina, in  
174 several regions of the United States (Colorado, California, Oregon, Utah and  
175 Washington), and in the Amazon and Pantanal regions in Brazil.

176 Unfortunately, it is unfeasible to manually annotate this massive amount  
177 of data in order to create masks with potential active fire pixels. However,  
178 such a large amount of data is needed for properly training deep convolutional  
179 networks. Therefore, in order to automatically generate segmentation masks,  
180 we relied on three set of conditions, well established in the field, designed  
181 by Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018).

182 Figure 2 shows the regions analyzed by these authors when developing and  
183 testing their sets of conditions — note that our dataset has a more thorough  
184 coverage of the globe, as shown in Figure 1.

185 In the following sections, we detail the sets of conditions for active fire  
186 detection proposed by Schroeder et al. (2016), Murphy et al. (2016) and Ku-  
187 mar and Roy (2018). For the sake of simplicity, let  $\rho_i$  be in the rest of  
188 this section the reflectance of channel  $c_i$ . The equation to convert between  
189 Landsat-8 channels to reflectance are detailed by the USGS <sup>3</sup>. For the Mur-  
190 phy et al. (2016) and Kumar and Roy (2018) sets of conditions, reflectances  
191 are corrected for the solar zenith angle; while for the Schroeder et al. (2016)  
192 conditions, this correction is not performed. Moreover, it is worth noting  
193 that Schroeder et al. (2016) perform a multi-temporal analysis to avoid false  
194 alarms over persistent heat sources. That approach was adapted and used by  
195 Kumar and Roy (2018), and could also bring beneficial results to our work.  
196 However, it was not include here since our focus is on active fire recognition  
197 for single images. After describing these conditions, we detail the procedure  
198 for extracting the segmentation masks, and a secondary dataset, containing  
199 manually annotated images.

## 200 2.1. *Schroeder et al. conditions*

201 Schroeder et al. (2016) proposed a set of conditions that uses seven

---

<sup>3</sup><https://www.usgs.gov/land-resources/nli/landsat/using-usgs-landsat-level-1-data-product>

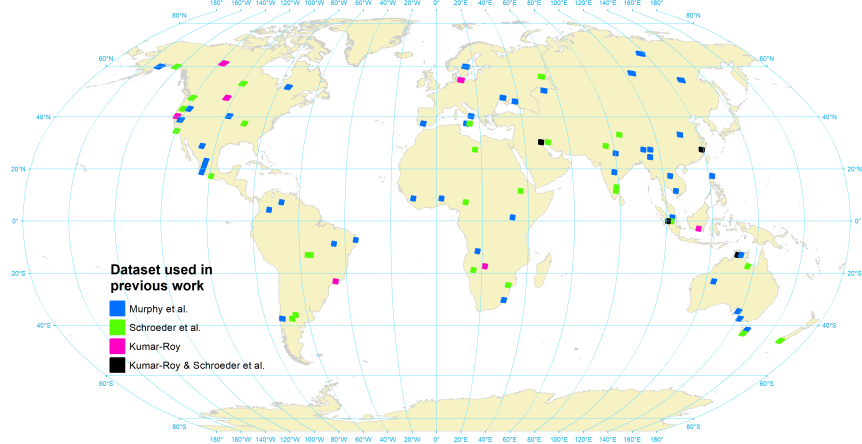


Figure 2: Locations covered in the studies from Murphy et al. (2016), Schroeder et al. (2016) and Kumar and Roy (2018), which observed, respectively, 45 regions, 30 regions and 11 distinct regions. Multiple temporal views were used for some of these regions.

202 Landsat-8 channels,  $c_1 - c_7$ , to identify active fire pixels. If all conditions are  
 203 satisfied the pixel is classified as true (fire), otherwise as false (non-fire).

The first set of conditions identifies unambiguous cases:

$$((R_{75} > 2.5) \text{ and } (\rho_7 - \rho_5 > 0.3) \text{ and } (\rho_7 > 0.5)) \text{ or} \\
((\rho_6 > 0.8) \text{ and } (\rho_1 < 0.2) \text{ and } (\rho_5 > 0.4 \text{ or } \rho_7 < 0.1))$$

204 where  $R_{ij}$  is the ratio between the reflectance in channels  $i$  and  $j$  ( $\rho_i/\rho_j$ ).

Then, by using a neighborhood of  $61 \times 61$  pixels, centered at each fire pixel candidate, the above conditions are relaxed so as to consider the region

context, namely

$$(R_{75} > 1.8) \text{ and } (\rho_7 - \rho_5 > 0.17) \text{ and } (R_{76} > 1.6) \text{ and} \\ (R_{75} > \mu_{R_{75}} + \max(3\sigma_{R_{75}}, 0.8)) \text{ and } (\rho_7 > \mu_{\rho_7} + \max(3\sigma_{\rho_7}, 0.08))$$

where  $\mu$  and  $\sigma$  denote the mean and standard deviation of the pixels in a  $61 \times 61$  window pixel neighborhood. These means and standard deviations are computed excluding fire and water pixels, the latter being classified by:

$$(\rho_4 > \rho_5) \text{ and } (\rho_5 > \rho_6) \text{ and } (\rho_6 > \rho_7) \text{ and } (\rho_1 - \rho_7 < 0.2) \\ \text{and } ((\rho_3 > \rho_2) \text{ or } ((\rho_1 > \rho_2) \text{ and } (\rho_2 > \rho_3) \text{ and } (\rho_3 > \rho_4)))$$

205 For further details about the above conditions refer to the work of Schroeder  
206 et al. (2016).

## 207 2.2. *Murphy et al. conditions*

208 Murphy et al. (2016) proposed a set of conditions that is non-contextual,  
209 in the sense that it does not rely on statistics computed from each pixel's  
210 neighborhood. It is based on channels  $c_5$ ,  $c_6$  and  $c_7$ .

Unambiguous active fire pixels are first identified by:

$$(R_{76} \geq 1.4) \text{ and } (R_{75} \geq 1.4) \text{ and } (\rho_7 \geq 0.15)$$

where  $R_{ij}$  is the ratio between the reflectance in channels  $i$  and  $j$ . Potential

fires that are in the immediate neighborhood of an unambiguous fire (in a  $3 \times 3$  window), are also classified as fires. Potential fires are those that satisfy:

$$((R_{65} \geq 2) \text{ and } (\rho_6 \geq 0.5)) \text{ or } ((\rho_7 \text{ is saturated}) \text{ or } (\rho_6 \text{ is saturated}))$$

211 Criteria for identifying saturated bands are defined by the USGS <sup>4</sup>.

212 *2.3. Kumar-Roy conditions*

The conditions proposed by Kumar and Roy (2018) are based on channels  $c_2 - c_7$ . This work is more recent than the others, and builds upon some of the insights behind previous work. Unambiguous fire pixels are identified as those that satisfy

$$\rho_4 \leq 0.53 \rho_7 - 0.214$$

They also take as unambiguous fire pixels those that are in the immediate 8-pixel vicinity of a fire pixel classified by the above condition, and which satisfy the more relaxed condition

$$\rho_4 \leq 0.35 \rho_6 - 0.044$$

---

<sup>4</sup><https://www.usgs.gov/land-resources/nli/landsat/landsat-collection-1-level-1-quality-assessment-band>

Potential fire pixels are identified by

$$((\rho_4 \leq 0.53 \rho_7 - 0.125) \text{ or } (\rho_6 \leq 1.08 \rho_7 - 0.048))$$

and are only kept if they satisfy the contextual test

$$(R_{75} > \mu_{R_{75}} + \max(3\sigma_{R_{75}}, 0.8)) \text{ and } (\rho_7 > \mu_{\rho_7} + \max(3\sigma_{\rho_7}, 0.08))$$

213 where  $R_{75}$  is the ratio between the reflectance in channels 7 and 5 ( $\rho_7/\rho_5$ ),  
 214 and  $\mu$  and  $\sigma$  denote the mean and standard deviation of the pixels in a  
 215 neighborhood around each candidate fire pixel. This contextual test is the  
 216 same used by the Schroeder et al. (2016) conditions, with two differences.  
 217 The first difference is the condition used to detect water pixels

$$(\rho_2 \geq \rho_3 \geq \rho_4 \geq \rho_5)$$

218 The second difference is that the region size is not fixed: we test progres-  
 219 sively larger neighborhoods ( $5 \times 5$ ,  $7 \times 7$ ,  $9 \times 9$ , and so on, up to maximum  
 220 of  $61 \times 61$  pixels), stopping when the region contains at least 25% of pixels  
 221 not classified as unambiguous or potential fires, or as water — these pixels  
 222 are excluded when computing the means and standard deviations. One point  
 223 that is not addressed in the original paper is how to proceed when the equa-  
 224 tions detect a large area containing potential fire pixels — in these cases,

even with a  $61 \times 61$  neighborhood there are candidate fire pixels mostly surrounded by excluded pixels, and the contextual test becomes unreliable. We considered such cases as non-fires, since doing otherwise made the algorithm highly prone to false detections.

#### 2.4. Segmentation masks

In order to create segmentation masks, we used the conditions from Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) over the original scenes with  $\approx 7,600 \times 7,600$  pixels. Figure 3 shows the locations where fires have occurred in our dataset for each set of conditions.

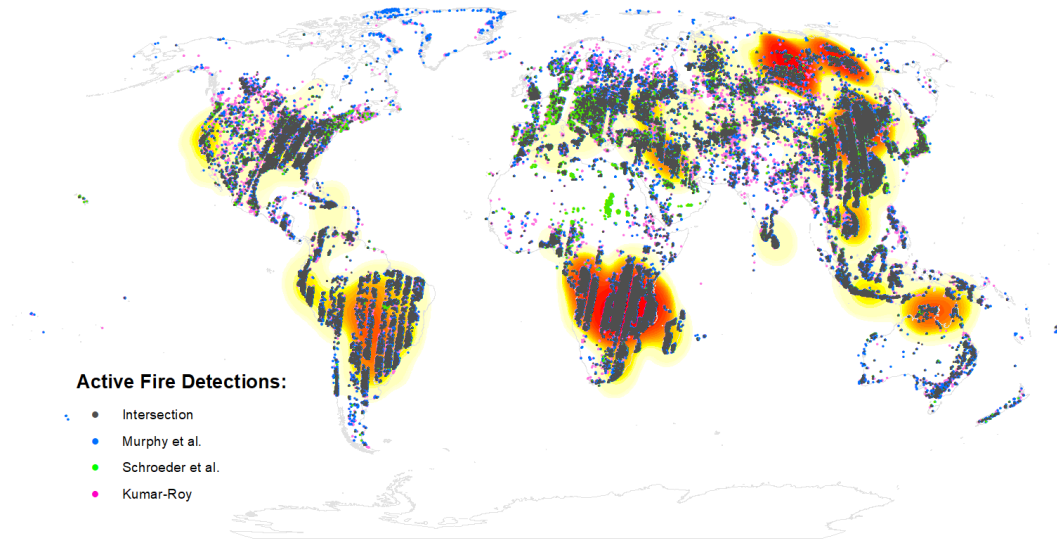


Figure 3: Fire incidence detection and heat map around the globe for August 2020 considering the outputs of Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018). **Dark gray dots** represent wildfire locations simultaneously detected by the three algorithms.

234 This dataset was then cropped into image patches with  $256 \times 256$  pixels  
 235 without overlap. The number of fire pixels per image patch for each set of  
 236 conditions is shown in the histogram in Figure 4. It can be seen that most  
 patches have a small number of fire pixels.

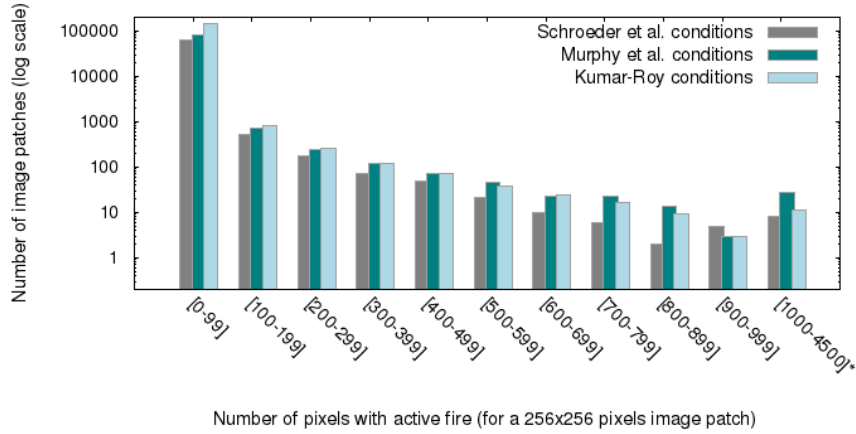


Figure 4: Distribution of image patches per fire pixel count. The total number of image patches (with  $256 \times 256$  pixel size) for the Schroeder et al. conditions is 62,385, for Murphy et al. conditions is 82,036 and 143,133 for the Kumar-Roy conditions. The last bin, marked with \*, has a larger interval because some few image patches had a large number of active fire pixels — keeping the bins with the same interval would result in a very sparse histogram.

237

238 Our dataset is publicly available, and contains 146,214 image patches,  
 239  $256 \times 256$  pixels each, around 180 GB uncompressed, along with their cor-  
 240 responding segmentation masks produced by the three sets of conditions  
 241 discussed above (approximately 28 GB uncompressed). The patches are  
 242 10-band, 16-bit TIFF images, with channels  $c_1, \dots, c_7$  and  $c_9, \dots, c_{11}$ . The  
 243 panchromatic channel  $c_8$ , with 15 meters of spatial resolution, will not be  
 244 considered in this work.

245 Figure 5 shows some image samples with active fire and the correspond-

246 ing segmentation masks for the three set of conditions. It can be seen that,  
 247 although the three algorithms generally agree on the presence of fire on a cer-  
 248 tain level, the pixelwise segmentation tends to differ between them. The Mur-  
 249 phy et al. (2016) conditions seem more sensitive to low level intensities (em-  
 250 ber), while all the algorithms have problems with very high intensities (which  
 251 appear with a greenish tinge in the RGB visualization). These differences  
 252 will be further discussed in Section 4.

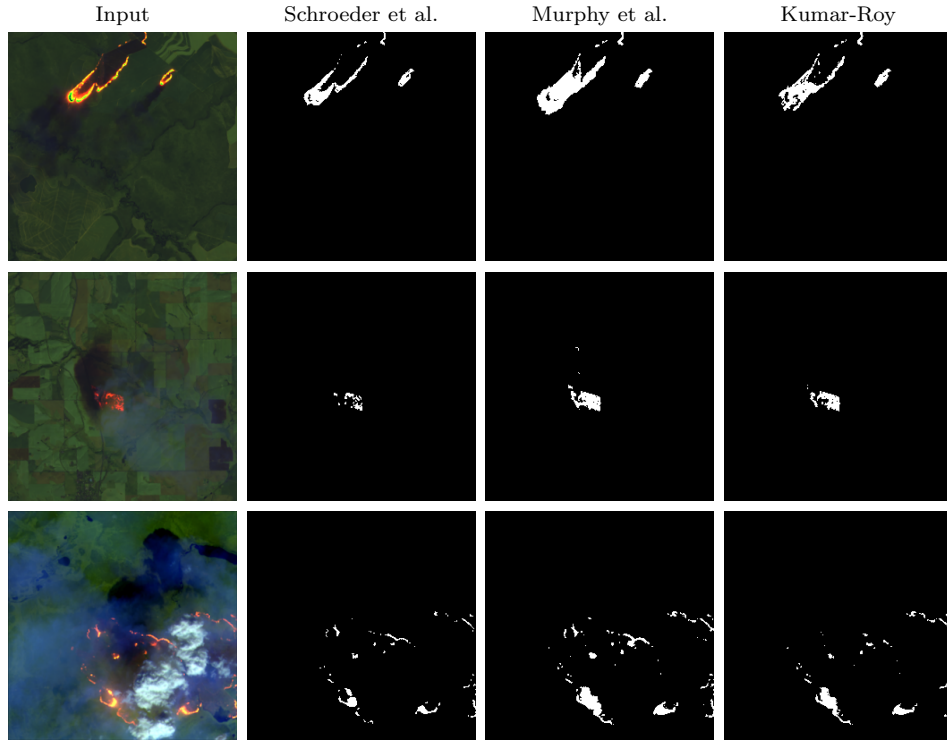


Figure 5: Image patches and the corresponding segmentation masks, extracted from the proposed dataset. For display purposes we used channels  $c_7$ ,  $c_6$  and  $c_2$  to compose the RGB bands, however, the original patches contain 10 bands.

253 We believe that this dataset can prove useful and relevant for other re-

254 searchers working on active fire detection, since it provides the Landsat-8  
255 data in a friendly format that can be directly used by existing machine learn-  
256 ing tools, and provides a challenging target for benchmarking, with samples  
257 covering a large variety of scenarios, including desert, rain forest, agricultural  
258 land, water, snow, clouds, cities, mountains, etc.

### 259 2.5. *Manually annotated dataset*

260 To further validate our trained models, we selected 13 Landsat-8 images,  
261 with  $\approx 7,600 \times 7,600$  pixels, which were split into 9,044 image patches with  
262  $256 \times 256$  pixels each, without overlap, in a total of 12 GB (**uncompressed**),  
263 captured in September 2020, and manually annotated them. Figure 6 shows  
264 the locations of these images.

265 It can be seen that these images are well spread spatially around the  
266 globe, covering the main continents, as well as the different features that  
267 may exist in these regions. There are active fires in 10 of these images,  
268 with different cloud cover and fire properties — some have many fire pixels,  
269 some have only a few; some have large fires, some have small isolated fire  
270 pixels; with the fire intensity also varying. The remaining 3 images have no  
271 active fires, but cover scenarios that proved challenging for at least one of  
272 the sets of conditions, which produced false positives: one from the Sahara  
273 desert, where the Schroeder et al. (2016) conditions detected fires, one from  
274 Greenland, where the Murphy et al. (2016) and Kumar and Roy (2018)  
275 conditions detected fires, and one from a city (Milan), since all the algorithms

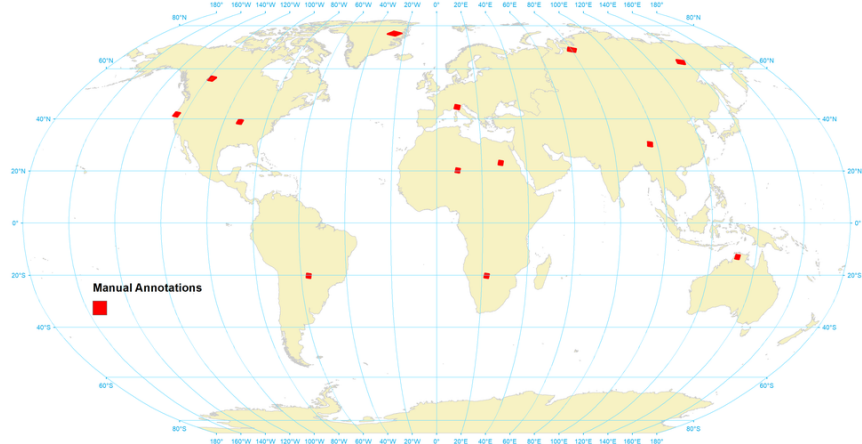


Figure 6: Active fire manual annotations: the active fire present in 13 Landsat-8 images (depicted by red squares) from September of 2020 — with  $\approx 7,600 \times 7,600$  pixels and selected to represent distinct regions of the planet — were manually annotated by the authors. These images and reference masks are also available as a dataset.

276 seem to produce false detections inside large urban settlements. Figure 7  
 277 shows an example of manually annotated image.

278 The purpose of evaluating the learned models against these images is  
 279 having a second reference for validating them, avoiding some potential bias  
 280 that might arise from training them on the outputs from the Schroeder et al.  
 281 (2016), Murphy et al. (2016) and Kumar and Roy (2018) algorithms. For  
 282 this reason, in this work, the manually annotated images will be exclusively  
 283 for testing the models, but not for training them.

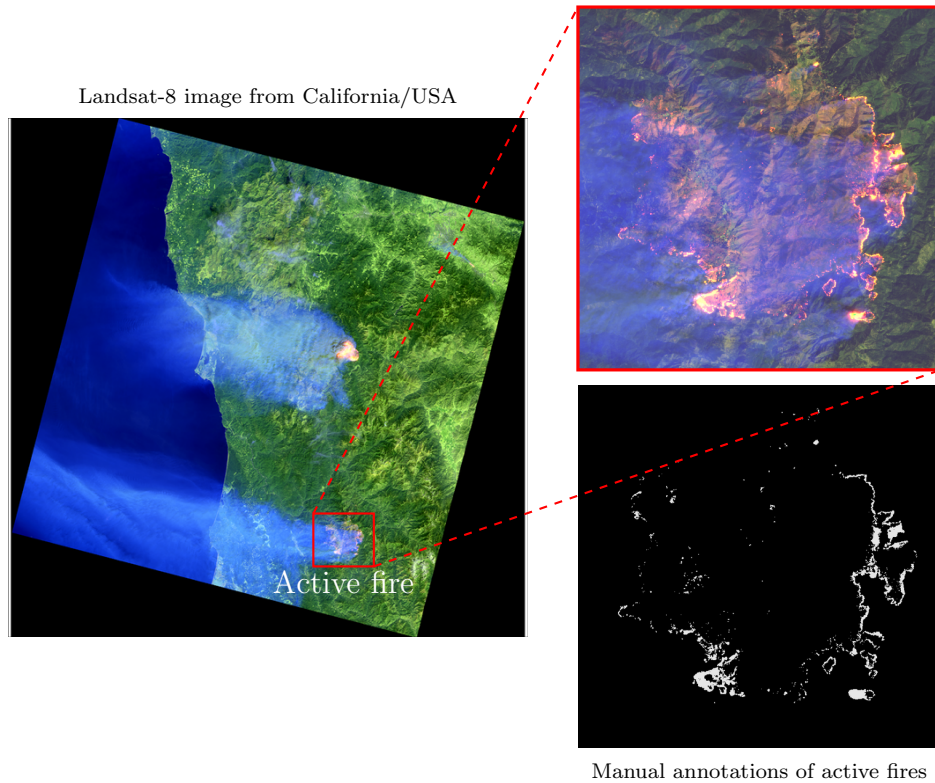


Figure 7: Manual annotation example: on the left, a Landsat-8 image from California, USA — reference code LC08\_L1TP\_046031\_20200908\_20200908\_01\_RT, WRS 046/031, from September 8th, 2020. This image has many wildfire regions, with a total of 55,386 active fire pixels being identified. Highlighted in red, one of the two major wildfire areas (top right) with an extension of  $22 \times 22$  km, with approximately 10,000 active fire pixels; the corresponding manual annotations are also shown (bottom right).

### 284 3. Methods

#### 285 3.1. Convolutional Neural Networks for Active Fire Detection

286

287 Convolutional neural networks (CNNs)(Goodfellow et al. (2016)) are a  
 288 type of feedforward neural network architecture that, in recent years, have  
 289 been enjoying a lot of attention, due to their success in a large variety of

tasks, many of them related to image processing (Garcia et al. (2018); Wang et al. (2020)), analysis (Petersson et al. (2016); Litjens et al. (2017)) and recognition (Paoletti et al. (2019); Yao et al. (2019)). This popularization is mainly because to advances both in processing capabilities — particularly the use of Graphic Processing Units (GPUs) for massively parallel computations — and algorithms and techniques, which led to the ability of learning deeper and more complex models. The reviews by Zhu et al. (2017), Ma et al. (2019a) and Yuan et al. (2020) discuss the major breakthroughs of deep learning for the remote sensing field, including main challenges, and the current state-of-the-art in environmental parameter retrieval, data fusion and downscaling, image registration, scene classification, object detection, land use and land cover classification and region segmentation.

A deep convolutional network is composed by a series of layers, each layer applying to the output of the previous layer a number of filters via the linear convolution operation, followed by some kind of activation function that introduces non-linearities (Goodfellow et al. (2016)). Deeper layers combine features extracted by previous layers, in such a way that the network is capable of encoding progressively more complex concepts (Lecun et al. (2015)). The training procedure for a CNN consists of discovering weights (coefficients and biases for the filters), which allow the input data to be transformed so as to approximate the desired output. In this paper, we consider only the so called supervised learning (Chapelle et al. (2010)), in which the network receives a number of labeled samples — inputs and their associated outputs —

313 and iteratively adjusts its weights via a backpropagation algorithm (Rumel-  
314 hart et al. (1986)) to make its output for a given input more and more similar  
315 to the presented sample output. Although CNNs are the current state of the  
316 art technique for a number of machine learning and computer vision tasks,  
317 there are several concerns that must be kept in mind when working with  
318 them. In this paper, we do not focus on issues directly related to challenges  
319 faced by CNNs in a mathematical or algorithmic level, but we do highlight  
320 that complex architectures demand a high computational power to be trained  
321 (i.e. whenever possible, it is desirable to have smaller and simpler architec-  
322 tures); and that a proper model must be trained with a large amount of  
323 unbiased data, something that raises concerns over both the required infras-  
324 tructure needed for transmission/storage and the availability of good quality  
325 public datasets.

326 The design of robust and highly optimized CNN architectures for active  
327 fire detection is not the main focus of this work. Nevertheless, it is impor-  
328 tant to demonstrate that a CNN trained on the proposed dataset is able to  
329 produce results similar to those obtained by conventional algorithms such  
330 as those from Schroeder et al. (2016), Murphy et al. (2016) and Kumar and  
331 Roy (2018). Our tests were based on networks derived from the U-Net ar-  
332 chitecture from Ronneberger et al. (2015), a very popular architecture for  
333 image segmentation tasks. U-Net is a fully convolutional network, with two  
334 symmetric halves, the first with pooling operations that decrease the data  
335 resolution and the second with upsampling operations that restore the data

336 to its original resolution. The first half extracts basic features and context,  
 337 and the second half allows for a precise pixel-level localization of each fea-  
 338 ture, with skip connections between the two halves being used to combine  
 339 the features from a layer with more abstract features from deeper layers.

340 Figure 8 shows the basic U-Net architecture employed in our tests. It is  
 341 essentially the same concept from the original U-Net, but our implementation  
 342 adds batch normalization after each convolutional layer, and contains addi-  
 343 tional dropout layers, which help avoiding problems with overfitting (Good-  
 344 fellow et al. (2016)) and is related to how the error propagates during training.

345  
 346 We considered 3 variations of the U-Net architecture. The first one, called  
 347 hereafter U-Net (10c), is the basic U-Net architecture, taking as input a 10-  
 348 channel image containing all the Landsat-8 image bands. The second archi-  
 349 tecture, called hereafter U-Net (3c), keeps the same structure, but replacing  
 350 the input with a 3-channel image containing only bands  $c_7$ ,  $c_6$ , and  $c_2$  — the  
 351 rationale behind this decision is checking whether it is possible to obtain a  
 352 good enough approximation of the results while relying on a reduced number  
 353 of bands, something that may be useful for reducing bandwidth, memory  
 354 usage and storage space. **E.g.** it is possible to have a considerable economy  
 355 of bandwidth resources when processing the images in Amazon Web Services  
 356 (AWS), where the original Landsat-8 images are hosted for free download  
 357 and where is possible to download only selected bands.

358 The third architecture, called hereafter U-Net-Light (3c), is a reduced

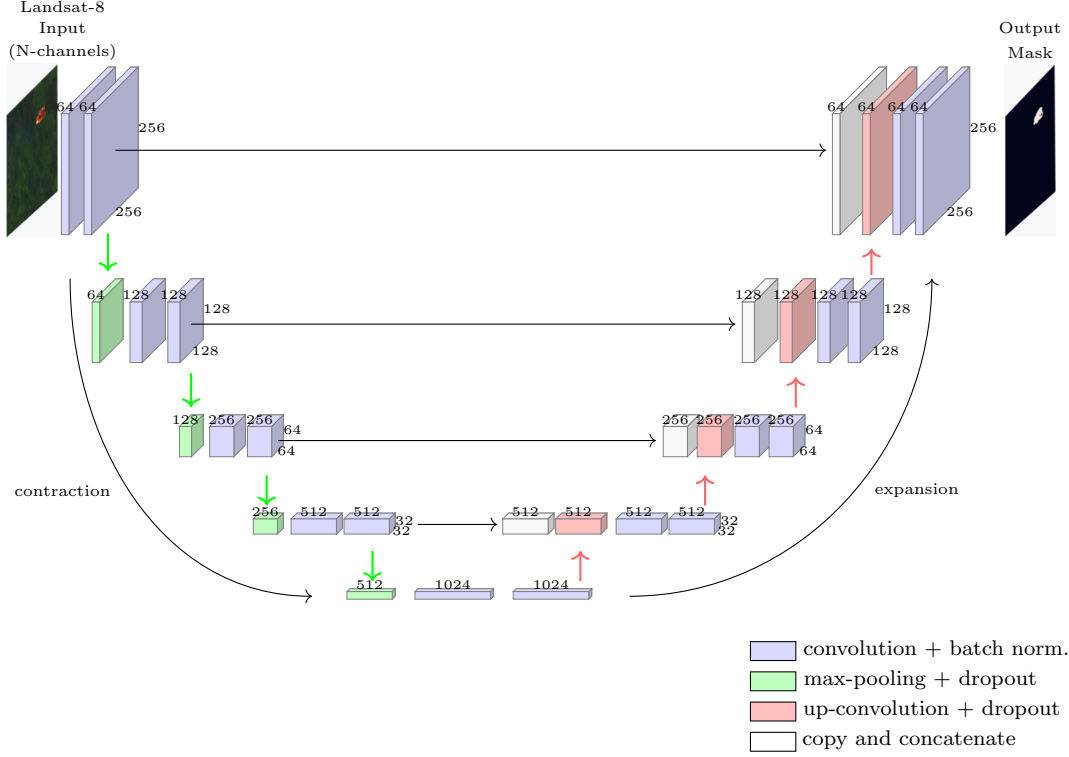


Figure 8: U-Net architecture for image segmentation.

version of the 3-channel U-Net, with the number of filters per layer being divided by 4 (i.e. 16 filters in the first layer, 32 in the second, etc.). We call this a “light” version of the original U-Net network. Relevant parameters for these architectures are shown in Table 1.

Table 1: U-Net-based architectures for active fire recognition.

|                  | # of input channels | # of trainable parameters | model size (MB) | inference time per patch (ms) |
|------------------|---------------------|---------------------------|-----------------|-------------------------------|
| U-Net (10c)      | 10                  | 34,529,153                | 132             | 36,8                          |
| U-Net (3c)       | 3                   | 34,525,121                | 132             | 36                            |
| U-Net-Light (3c) | 3                   | 2,161,649                 | 8,5             | 25,5                          |

All the variations have as output a 1-channel binary image, with  $256 \times 256$

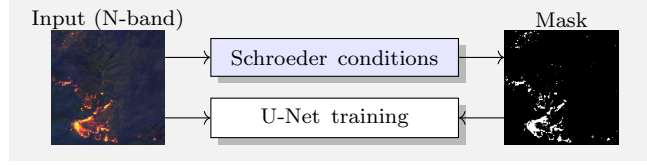
pixels, where 1 and 0 represent, respectively, fire and non-fire locations. To obtain these binary outputs, the CNN outputs are thresholded so that any pixel with value above 0.25 is set to 1 (this threshold was empirically defined after an initial test run on a small fraction of the dataset — we observed a performance gain of  $\approx 1 - 2\%$  on the F-score, computed as explained in Section 3.3, over using the standard 0.5 value).

Each of the 3 CNN architectures was trained and tested to approximate 5 different situations: each of the 3 considered sets of conditions, as well as their intersection and a “best-of-three” voting (i.e. a pixel in the mask is set as active fire if at least two sets of conditions agree that it is a fire pixel). This adds up to a total of 15 tested scenarios, as illustrated in Figure 9.

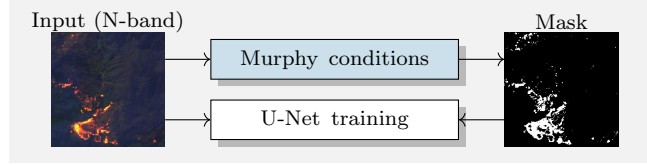
### 3.2. *Implementation and Experimental Setup*

The proposed CNN architectures were implemented and tested, to verify their ability to approximate the results from the sets of conditions from Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018). We also evaluated the performance of the CNNs and sets of conditions on a set of manually annotated images.

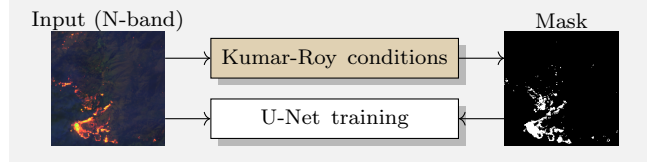
The image patches from the dataset were randomly divided in 3 sets, for training, validation and test, containing, respectively, 40%, 10% and 50% of the samples. The experiments were carried out on an Intel Core i8, 64 GB of RAM, running Linux, with a Titan Xp (12 GB) GPU. The implementation



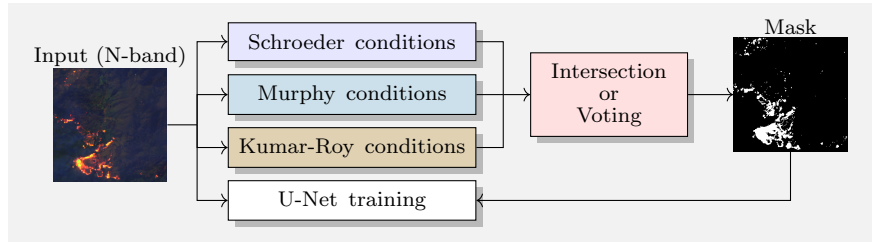
(a)



(b)



(c)



(d)

Figure 9: Different situations for training and testing the deep convolutional neural networks. Each of the 3 network architectures (U-Net (10c), U-Net (3c) and U-Net-Light (3c)) was trained and tested to approximate the outputs of 5 different sets of conditions: the 3 algorithms for active fire detection Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018), as shown in (a), (b) and (c), as well as their intersection and a “best-of-three” voting, as shown in (d).

385 was in the Python language, using the TensorFlow<sup>5</sup> library. We trained each  
 386 architecture using Adam optimization (learning rate of 0.001), a batch size of

---

<sup>5</sup>[www.tensorflow.org](http://www.tensorflow.org)

16, and binary cross entropy as the loss function, for 50 epochs, or until the loss on the validation set was not improved for at least 5 epochs. Figure 10 shows an example of one training run (loss  $\times$  epoch) — in that case, for the U-Net (3c) trained on the intersection of the 3 algorithms (Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018)). All other training runs had a similar behavior, with the number of epochs usually fluctuating between 13~50.

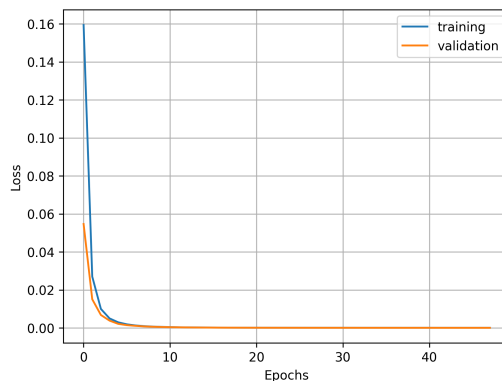


Figure 10: Training loss for the U-Net (3c), taking the intersection of 3 algorithms as target. Other training runs had a similar behavior. After around 10 epochs, there are slight variations to the loss, which are hard to see due to the scale in the  $y$  axis, which was automatically selected by the library we used.

### 3.3. Evaluation metrics

We evaluated the architectures according to the primary metrics used in the main semantic segmentation challenges such as PASCAL VOC<sup>6</sup> (Everingham et al. (2010)), KITTI<sup>7</sup> (Fritsch et al. (2013)), and COCO<sup>8</sup> (Lin et al.

<sup>6</sup>Visual Object Challenge

<sup>7</sup>Karlsruhe Institute of Technology and Toyota Technological Institute at Chicago

<sup>8</sup>Common Objects in Context — <http://cocodataset.org>

(2014)). The  $F$ -score (Chinchor and Sundheim (1993)), also known as dice coefficient, is a widely used metric for ranking purposes, and consists of the harmonic mean of precision  $P$  and recall  $R$  metrics

$$F = \frac{2}{1/P + 1/R} \quad (1)$$

such that

$$P = \frac{tp}{tp + fp} \quad R = \frac{tp}{tp + fn} \quad (2)$$

where  $tp$  (true positives) is the number of correctly classified fire pixels,  $fp$  (false positives) the number of non-fire pixels incorrectly classified as fire, and  $fn$  (false negatives) the number of fire pixels incorrectly classified as non-fire.

We also report the Intersection-Over-Union (IoU) metric (O. Pinheiro et al. (2015)), also known as Jaccard index

$$\text{IoU} = \frac{tp}{tp + fn + fp} \quad (3)$$

The IoU metric conveys the same information as the  $F$ -score —  $\text{IoU} = F/(2 - F)$  — but it is also widely in many segmentation challenges, so for clarity we report both values. All metrics were computed as per the COCO evaluation implementation:  $tp$ ,  $fp$  and  $fn$  values are accumulated for all the images and the metrics are computed only once for the entire dataset (i.e. the reported values are not the average per-patch performance, but the global

405 per-pixel performance).

406 As in the KITTI challenge (Lyu et al. (2019)), we ranked the algorithm’s  
407 performances according to the  $F$ -score metric. One common metric that we  
408 do not list in our experiments is the pixel accuracy metric, which takes into  
409 account the number of true negatives (non-fire pixels detected as such). In  
410 our tests, we observed a large class imbalance between fire and non-fire pixels  
411 — more than 99% of the total samples are non-fire pixels — and that always  
412 resulted in very high accuracy — even for an extreme case of a detector that  
413 fails to detect any fire pixels could still seem to perform reasonably well,  
414 making pixel accuracy a misleading metric.

## 415 4. Results and Discussion

### 416 4.1. Performance on the automatically segmented patches

417 For the first round of experiments, we used the CNN architectures trained  
418 and validated with 50% of the image patches against the other 50% reserved  
419 for testing, in order to show how well each trained architecture can approx-  
420 imate the original segmentations obtained by the Schroeder et al. (2016),  
421 Murphy et al. (2016) and Kumar and Roy (2018) conditions, as well as com-  
422 binations of their outputs (their intersection and a “best-of-three” voting).  
423 In total, we had 15 testing scenarios, as shown in Table 2, with 3 CNN ar-  
424 chitectures (U-Net (10c), U-Net (3c) and U-Net-Light (3c)) to approximate  
425 the 5 target segmentation masks.

426 As can be seen, all architectures were able to reproduce the behavior of

Table 2: Active fire recognition performance: we used the set of conditions from Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) to train three convolutional neural networks for active fire segmentation: U-Net (10c) uses as input a 10-channel image from Landsat-8; U-Net (3c) uses as input a 3-channel image, containing only bands  $c_7$ ,  $c_6$ , and  $c_2$ , which have a good response for active fire; and U-Net-Light (3c) is a reduced version of the 3-channel U-Net. The performances are related to the CNN ability to reproduce the masks from Schroeder et al. (2016), Murphy et al. (2016), Kumar and Roy (2018), and also the intersection and consensus of their masks, in a testing set never seen by these architectures.

| Mask                    | CNN Architecture                    | Metrics (%) |             |             |             |
|-------------------------|-------------------------------------|-------------|-------------|-------------|-------------|
|                         |                                     | $P$         | $R$         | IoU         | $F$         |
| Schroeder <i>et al.</i> | U-Net (10c)                         | 86.8        | <b>89.7</b> | 78.9        | 88.2        |
|                         | U-Net (3c)                          | 89.8        | 88.8        | <b>80.7</b> | <b>89.3</b> |
|                         | U-Net-Light (3c)                    | <b>90.8</b> | 86.1        | 79.2        | 88.4        |
| Murphy <i>et al.</i>    | U-Net (10c)                         | <b>93.6</b> | 92.5        | 87.0        | 93.0        |
|                         | U-Net (3c)                          | 89.1        | <b>97.6</b> | 87.2        | 93.2        |
|                         | U-Net-Light (3c)                    | 92.6        | 95.1        | <b>88.4</b> | <b>93.8</b> |
| Kumar-Roy               | U-Net (10c)                         | <b>84.6</b> | <b>94.1</b> | <b>80.3</b> | <b>89.1</b> |
|                         | U-Net (3c)                          | 84.2        | 90.6        | 77.4        | 87.3        |
|                         | U-Net-Light (3c)                    | 76.8        | 93.2        | 72.7        | 84.2        |
| Intersection            | Schroeder <i>et al.</i> U-Net (10c) | 84.4        | <b>99.7</b> | 84.2        | 91.4        |
|                         | Murphy <i>et al.</i> U-Net (3c)     | <b>93.4</b> | 92.4        | <b>86.7</b> | <b>92.9</b> |
|                         | Kumar-Roy U-Net-Light (3c)          | 87.4        | 97.3        | 85.4        | 92.1        |
| Voting                  | Schroeder <i>et al.</i> U-Net (10c) | <b>92.9</b> | 95.5        | <b>89.0</b> | <b>94.2</b> |
|                         | Murphy <i>et al.</i> U-Net (3c)     | 91.9        | 95.3        | 87.9        | 93.6        |
|                         | Kumar-Roy U-Net-Light (3c)          | 90.2        | <b>96.5</b> | 87.3        | 93.2        |

427 the different algorithms reasonably well. It is interesting to note that a larger  
428 network (U-Net (10c)) is not necessarily superior to a network that uses a  
429 reduced number of channels (U-Net (3c)), nor to a less complex network with  
430 a reduced number of channels (U-Net-Light (3c)). That means it is likely  
431 that this particular application does not demand very complex architectures,  
432 and more importantly, that channels  $c_7$ ,  $c_6$  and  $c_2$  indeed contain most of the  
433 information needed for detecting active fires. Moreover, there is a tendency

that models with higher precision have lower recall, and vice-versa, but not to a large margin in most cases, which shows the learning procedure is balancing false and missed detections.

Figures 11, 12 and 13 compare the outputs produced by the CNNs to the outputs of, respectively, the Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) sets of conditions, for image samples from the test set. It can be seen that the CNNs indeed succeed in approximating the target masks, especially considering the presence of fire in a certain area, with most differences at the pixel level being small. Even the 3-channel architectures are successful in most cases, once again confirming that channels  $c_7$ ,  $c_6$  and  $c_2$  contain most of the information needed for detecting active fires.

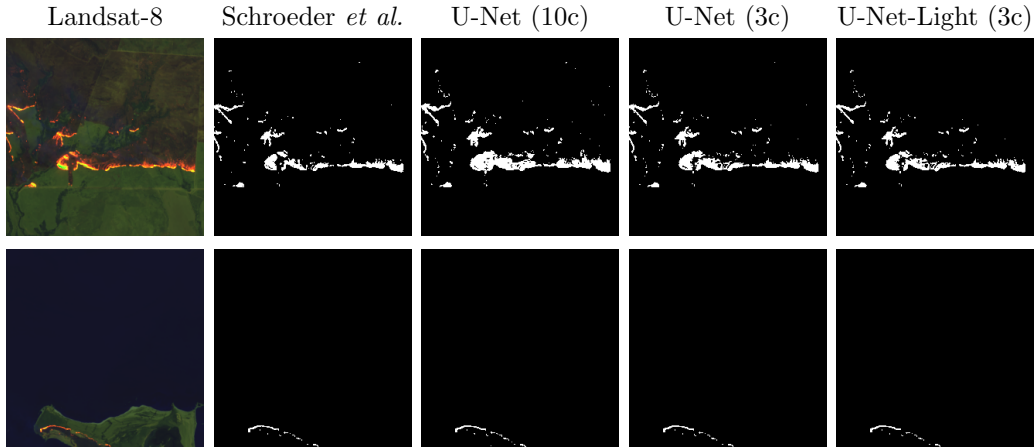


Figure 11: Active fire segmentation results obtained by the algorithm from Schroeder et al. (2016), and by the networks trained using those outputs as the target masks.

The second row from Figure 13, which has several small clusters of fire pixels detected by the original algorithm and the 10-channel CNN, but which

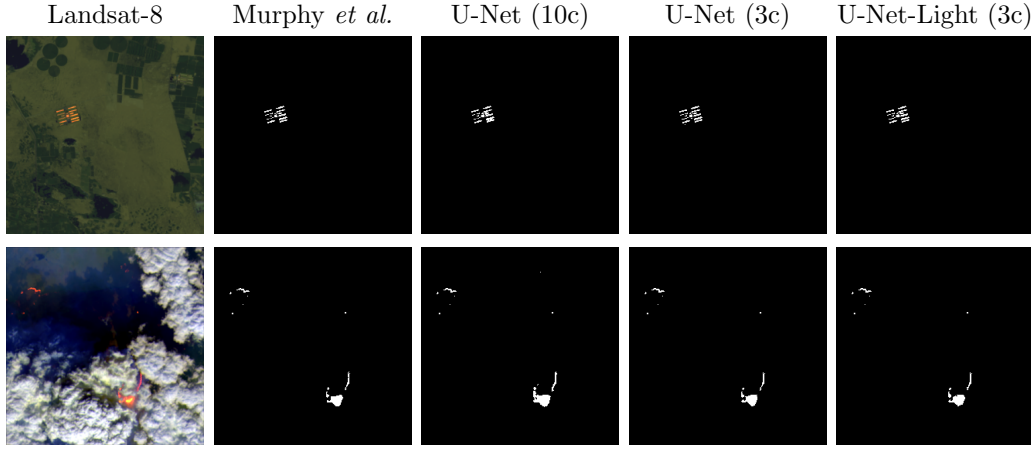


Figure 12: Active fire segmentation results obtained by the algorithm from Murphy *et al.* (2016), and by the networks trained using those outputs as the target masks.

447 were not found by the 3-channel networks (and which are, consistently, hard  
 448 to see in the 3-channel visualization), is an exception of good approximation  
 449 for the networks with 3-channels.

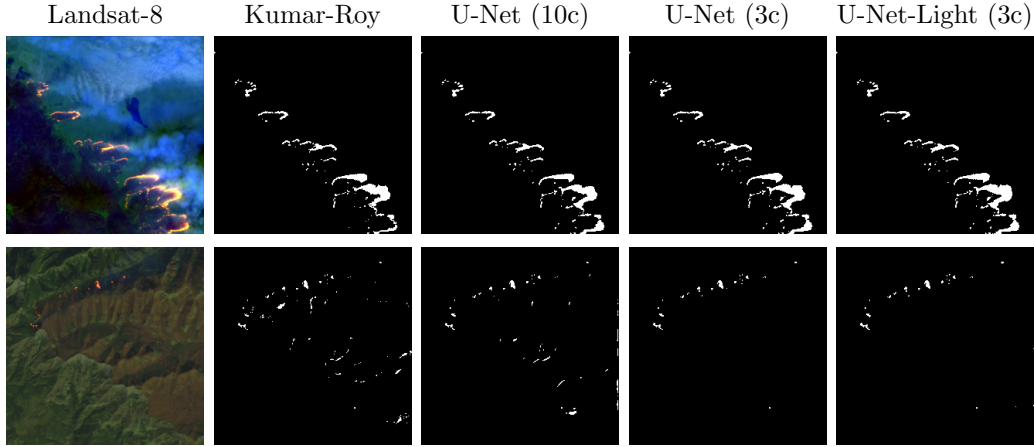


Figure 13: Active fire segmentation results obtained by the algorithm from Kumar and Roy (2018), and by the networks trained using those outputs as the target masks.

450 In Figures 14 and 15, we compare the outputs from the three sets of

451 conditions to their combination (first their intersection, then the “best-of-  
 452 three” voting scheme), as well as the output of the 3-channel U-Net trained  
 453 to approximate that combination.

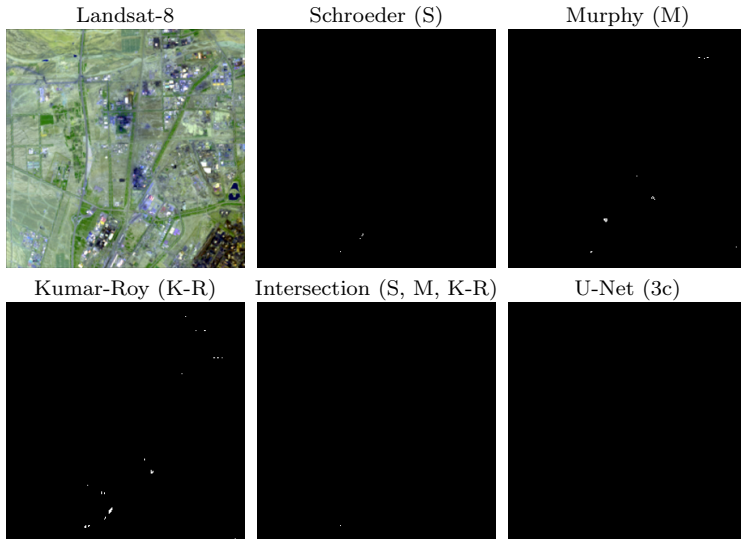


Figure 14: The Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) algorithms had false detections inside a city, but their intersection and the 3-channel U-Net trained to approximate it were able to avoid them.

454 Figure 14 shows that, while the three sets of conditions may incorrectly  
 455 detect active fires inside urban settlements, their intersection is usually more  
 456 robust (although that may also result in lower recall). As for Figure 15, it  
 457 shows the voting scheme helps reducing behaviors associated with only one  
 458 of the sets of conditions, such as the tendency of the Murphy *et al.* conditions  
 459 to produce thicker fire clusters and, in this example, the “hole” in the lower  
 460 fire region from the Kumar-Roy conditions. In both cases, the trained CNNs  
 461 were successful in approximating the combined masks (both samples come  
 462 from the test set). Another question brought by these images is: how do the

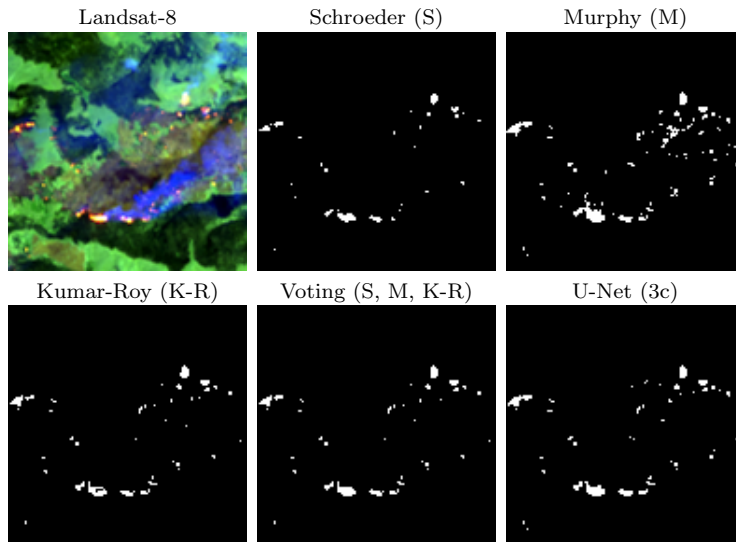


Figure 15: The 3-channel U-Net was able to approximate a “best-of-three” voting scheme that takes the outputs from the Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) algorithms. The voting scheme reduces errors caused by artifacts produced by only one of the algorithms.

463 3 sets of conditions compare to each other, and to the networks that learn to  
 464 combine them? This matter will be analyzed in more depth in Section 4.2.

#### 465 4.2. Performance on the manually annotated images

466 For the second round of experiments, we fed the CNN models trained in  
 467 the first round (Section 4.1) with Landsat-8 image patches from the manu-  
 468 ally annotated dataset, and compared the results with the hand-segmented  
 469 masks. For reference, the outputs from the Schroeder et al. (2016), Murphy  
 470 et al. (2016) and Kumar and Roy (2018) algorithms were also compared to  
 471 the hand-segmented masks. This allows us to observe how well the networks  
 472 learned from the automatically segmented patches perform when compared to  
 473 a human specialist (one must keep in mind, however, that different specialists

would produce distinct segmentations, since it is very subjective sometimes to separate the frontiers of burnt areas from areas that are still burning, all this added to the presence of smoke and many levels of fire intensity). For works on burnt areas the reader can refer to Chen et al. (2020); Malambo and Heatwole (2020).

Table 3 shows the performances obtained for each one of the 15 scenarios involving CNN architectures, as described in section 4.1, as well as, the performances for the original set of conditions of Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) without the use of any CNN architecture. As can be seen, the best F-scores were around 90% — this is expected, as there are biases from both the human specialist labeling the data and the algorithms themselves, which are tuned to perform under certain assumptions. As an example, we also show in Figure 16 a selected image patch from Landsat-8, with a large area of active fire, along with its associated manual segmentation, and also the segmentation masks produced for each scenario being considered.

As can be seen in Table 3, the original conditions from Murphy et al. (2016) performed better than those from Schroeder et al. (2016) and Kumar and Roy (2018) in the manually annotated dataset, mostly due to its tendency to detect more fire pixels than the other conditions, resulting in a very high recall, albeit at the cost of a lower precision (i.e. it is more prone to false detections). The Kumar and Roy (2018) original conditions had lower precision and recall compared to the Schroeder et al. (2016) conditions — in

Table 3: Active fire recognition performance against **manual annotations** for 13 Landsat-8 scenes. The entries with — do not use a CNN, the performances are computed directly from the Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) sets of conditions against the manual annotations. The remaining entries correspond to deep networks trained with masks automatically extracted using these sets of conditions. All entries were compared against the same manual annotations and the boldface values show the best performance for each metric.

| Mask                              | CNN Architecture        | Metrics (%)      |             |             |             |             |
|-----------------------------------|-------------------------|------------------|-------------|-------------|-------------|-------------|
|                                   |                         | $P$              | $R$         | IoU         | $F$         |             |
| Schroeder <i>et al.</i> (non-CNN) | —                       | 88.1             | 70.2        | 64.1        | 78.1        |             |
| Schroeder <i>et al.</i>           | U-Net (10c)             | 86.9             | 88.4        | 78.0        | 87.7        |             |
|                                   | U-Net (3c)              | 89.5             | 80.5        | 73.6        | 84.8        |             |
|                                   | U-Net-Light (3c)        | 89.0             | 78.8        | 71.8        | 83.6        |             |
| Murphy <i>et al.</i> (non-CNN)    | —                       | 76.6             | 96.1        | 74.3        | 85.2        |             |
| Murphy <i>et al.</i>              | U-Net (10c)             | 74.8             | 96.9        | 73.0        | 84.4        |             |
|                                   | U-Net (3c)              | 72.7             | <b>97.2</b> | 71.2        | 83.2        |             |
|                                   | U-Net-Light (3c)        | 75.5             | 96.9        | 73.7        | 84.9        |             |
| Kumar-Roy (non-CNN)               | —                       | 82.3             | 68.4        | 59.7        | 74.7        |             |
| Kumar-Roy                         | U-Net (10c)             | 82.2             | 94.2        | 78.3        | 87.8        |             |
|                                   | U-Net (3c)              | 84.5             | 93.4        | 79.8        | 88.8        |             |
|                                   | U-Net-Light (3c)        | 78.8             | 96.9        | 76.9        | 86.9        |             |
| Intersection (non-CNN)            | —                       | <b>92.6</b>      | <b>57.6</b> | <b>55.1</b> | <b>71.0</b> |             |
| Intersection                      | Schroeder <i>et al.</i> | U-Net (10c)      | 91.8        | 75.4        | 70.6        | 82.8        |
|                                   | Murphy <i>et al.</i>    | U-Net (3c)       | 90.8        | 73.5        | 68.4        | 81.2        |
|                                   | Kumar-Roy               | U-Net-Light (3c) | 90.8        | 72.8        | 67.8        | 80.8        |
| Voting (non-CNN)                  | —                       | <b>87.7</b>      | <b>79.1</b> | <b>71.2</b> | <b>83.2</b> |             |
| Voting                            | Schroeder <i>et al.</i> | U-Net (10c)      | 83.6        | 94.0        | 79.3        | 88.5        |
|                                   | Murphy <i>et al.</i>    | U-Net (3c)       | 86.4        | 93.0        | 81.1        | 89.6        |
|                                   | Kumar-Roy               | U-Net-Light (3c) | 87.2        | 92.4        | <b>81.4</b> | <b>89.7</b> |

497 fact, it fails to detect several fire pixels in the example of Figure 16, an issue  
498 that will be further discussed briefly. All the sets of conditions had problems  
499 detecting the core of the fire at the bottom of the image, which has very high  
500 intensity.

501 As for the CNNs, most differences between them and the original sets of

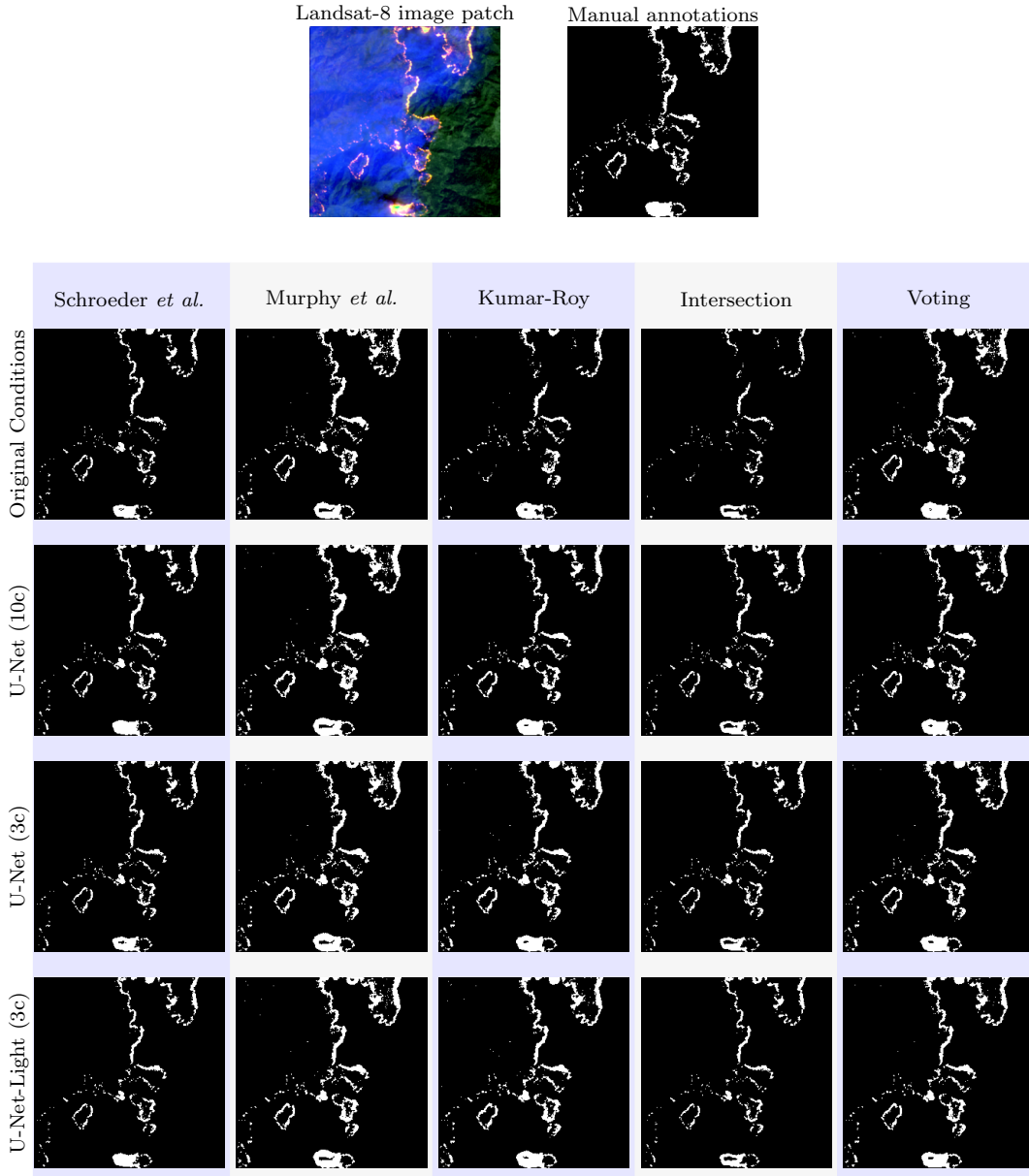


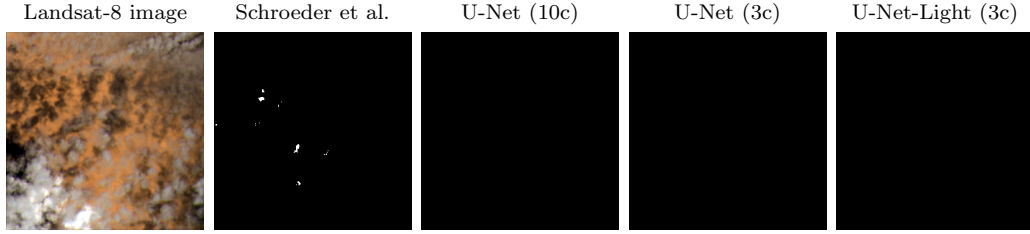
Figure 16: Active fire segmentation results produced by the Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy (2018) sets of conditions, the deep networks trained to approximate them, and their combined outputs (intersection and “best of three” voting), for a Landsat-8 image patch.

502 conditions were actually positive, with the deep networks performing better  
 503 than the Schroeder et al. (2016) and Kumar and Roy (2018) sets of conditions,  
 504 with higher recall but without a corresponding hit on precision, as shown in  
 505 Table 3. In Figure 16, the networks trained to approximate the Kumar  
 506 and Roy (2018) conditions were able to detect many fire pixels missed by  
 507 the original algorithm. The deep networks have also shown a higher recall  
 508 (but lower precision) compared to the Murphy et al. (2016) conditions. As  
 509 expected, by taking the intersection of three network outputs, we obtained  
 510 higher precision but a considerably lower recall, since this combination keeps  
 511 only those pixels that all three networks agree are fire pixels, that is, it  
 512 is very restrictive (the intersection of the three original sets of conditions  
 513 leads to the highest precision, but at the cost of a very low recall). The  
 514 best overall performance — including the original sets of conditions — was  
 515 obtained by the voting scheme applied to the outputs of three networks,  
 516 which resulted in high recall without a great impact on the precision, with  
 517 the 3-channel light U-Net being the one whose results were the most similar  
 518 to the manual segmentations, but with no significant difference for the larger  
 519 3-channel U-Net. We emphasize that these networks were trained using only  
 520 automatically segmented samples, so the improved performance cannot be  
 521 caused by some bias learned from the manually annotated patches, which  
 522 were only seen in the tests.

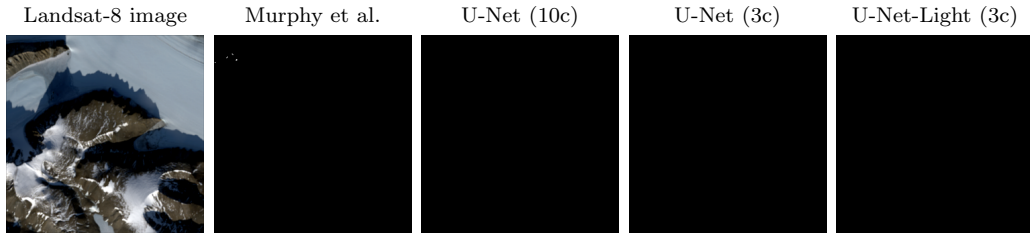
523 Besides the quantitative results discussed so far, we analyzed a number  
 524 of individual cases where the performance or recall obtained by one of the

525 approaches was particularly low. One example is the segmentation from  
 526 the Kumar and Roy (2018) original conditions shown in Figure 16, which  
 527 missed many fire pixels. We observed that these pixels were in fact detected  
 528 as fire pixels, but were later discarded after being regarded as water pixels,  
 529 due to some very similar values in bands  $\rho_2$  to  $\rho_5$ . Other cases are presented  
 530 in Figure 17, which shows failure cases for each one of the sets of condi-  
 531 tions. In the first row, Figure 17(a), the Schroeder et al. (2016) conditions  
 532 produced false detections in a Sahara desert region. In the second and third  
 533 rows, Figure 17(b, c), the Murphy et al. (2016) and Kumar and Roy (2018)  
 534 conditions produced false detections in some Greenland regions. In the map  
 535 previously shown in Figure 3 these errors can be seen as green dots in the  
 536 Sahara region and blue and magenta dots in Greenland — but note the map  
 537 refers to data from August 2020, while the examples in Figure 17 were taken  
 538 in September 2020, indicating that the algorithms are consistently producing  
 539 false detections in some situations.

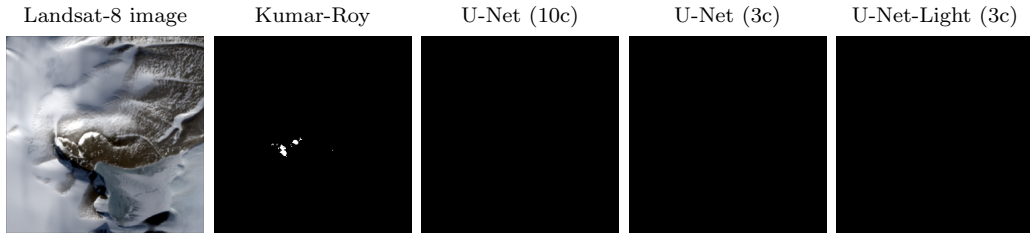
540 By analyzing these cases in further detail, by looking at each step and  
 541 component from the conditions, we noticed that most errors were not caused  
 542 by fundamental problems with the conditions themselves, but by reflectance  
 543 values that appear very close to the thresholds from one rule or another. For  
 544 example, some errors could be avoided if the 0.2 threshold from one unam-  
 545 biguous fire condition from Schroeder et al. (2016) was changed to 0.194 —  
 546 but that could potentially also lead to real fires not being detected. This  
 547 highlights one of the challenges faced when handcrafting sets of rules: some



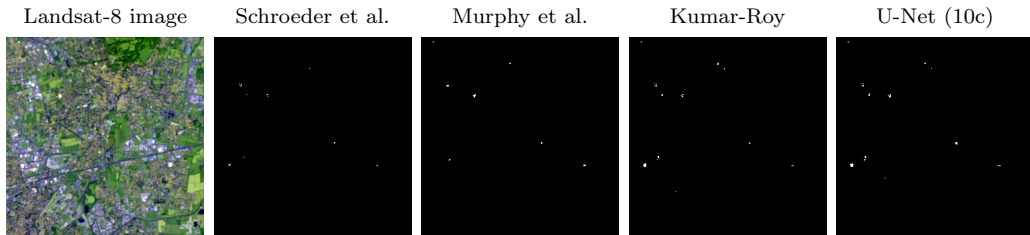
(a) Segmentations from Schoeder et al. conditions and for U-Net architectures based of his conditions.



(b) Segmentations from Murphy conditions and for U-Net architectures based of his conditions.



(c) Segmentations from Kumar-Roy conditions and for U-Net architectures based of his conditions.



(d) Segmentations from Schoeder et al., Murphy et al, Kumar-Roy and a U-Net architecture.

Figure 17: Classification errors: (a) region from Sahara desert WRS: 189/046, date 2020/09/18; (b) and (c) from Greenland, WRS: 003/006, date: 2020/09/11; and (d) city of Milan, Italy, WRS: 193/029, date: 2020/09/14, where all original conditions and CNN architectures reported misclassifications for bright and/or reflective urban elements.

548 situations must be simplified to keep the algorithm readable and understand-  
549 able, so thresholds are usually rounded to 1 or 2 decimal places, and only  
550 the most important relations between bands are considered. Machine learn-  
551 ing techniques are able to encode more complex rules, with more precise  
552 weights, coefficients and thresholds; although this may also mean that the  
553 learned rules and relations can be hard to understand and describe by a hu-  
554 man. In any case, the deep networks were able to avoid the errors in the  
555 samples shown in the first three rows of Figure 17. One case that was not  
556 solved by the CNNs is shown in the last row (d) of Figure 17 — all the  
557 approaches produced false detections around urban areas. From these ex-  
558 amples, we believe that these kinds of persistent errors must be addressed  
559 differently, by considering temporal analysis as described by Schroeder et al.  
560 (2016), since it is very unlikely that an active fire will remain active at the  
561 same location for several months.

#### 562 4.3. Discussion

563 In sections 4.1 and 4.2, we discussed particular results regarding, respec-  
564 tively, how well the deep learning models could approximate the results from  
565 handcrafted algorithms, and how the CNNs and handcrafted sets of condi-  
566 tions performed when compared to a human specialist. The main findings  
567 were that the CNNs can indeed approximate the handcrafted algorithms,  
568 even with a reduced architecture and relying on only 3 channels of sensor  
569 data. Moreover, we noted that, by combining multiple CNN outputs, it is

possible to obtain better performance compared to the original sets of conditions, and the situations where the CNNs differ from the original algorithms can actually be positive, allowing more complex rules and relations, albeit with the trade-off of being **hard to understand** or describe by a human.

It is worth noting that the CNN architectures we described in this paper were the ones that reported the most promising results, however, several other variations were tested, such as architectures without batch normalization (which failed, most likely due to the vanishing gradients problem), variations with more and less convolutional layers or filters per layer (i.e. deeper/shallower, and wider/narrower architectures). We emphasize that we did not focus on finding highly optimized CNN architectures, the feasibility of the approach being our main concern. Thus, the results we obtained could be potentially improved in a number of ways: future research could explore the use of image super-resolution ([Gargiulo et al. \(2019\)](#), [Ma et al. \(2019b\)](#)); ensembles of distinct convolutional neural network architectures ([Minetto et al. \(2019\)](#)); and spatio-temporal information ([Boschetti et al. \(2017\)](#)) to differentiate between temporally-persistent fire pixels from places like factories or rooftops from deforestation wildfires. Furthermore, machine learning-driven approaches can be quickly adapted to new sensors ([Mateo-Garca et al. \(2020\)](#)) (e.g. Sentinel-2), possibly relying on transfer learning, without requiring specifically designed sets of conditions.

## 591 5. Conclusions

592

593 In this paper, we addressed the problem of active fire detection using  
594 deep learning techniques. We introduced a new large-scale dataset, con-  
595 taining 146,214 image patches extracted from the Landsat-8 satellite, along  
596 with associated outputs produced by three well established handcrafted al-  
597 gorithms (Schroeder et al. (2016), Murphy et al. (2016) and Kumar and Roy  
598 (2018)). A second dataset was also created, with 9,044 image patches and an-  
599 notations manually produced by a human specialist. We hope these datasets  
600 can prove useful for the community, since they set a challenging target while  
601 keeping the data in a friendly format for existing machine learning tools,  
602 allowing researchers to test different architectures and approaches. Besides  
603 the datasets, we presented a study on how convolutional neural networks can  
604 be used to approximate the outputs from the considered handcrafted algo-  
605 rithms, as well as on how the outputs from multiple models can be combined  
606 to achieve better performance than the individual sets of conditions.

607 Our study showed that handcrafted sets of conditions for active fire recog-  
608 nition are hard to be defined, and small variations in fixed thresholds may  
609 cause false positives or negatives, on the other hand, deep learning techniques  
610 are able to encode more complex rules to improve over this aspect but at the  
611 cost of relations that are hard to understand and describe by a human.

612 Future work shall focus on exploring alternative methods for active fire  
613 detection over the proposed datasets, as well as better exploring the hy-

614 pothesis that models trained over samples from one satellite can perform  
615 well when applied to other satellites. It is also possible to develop means of  
616 combining detection results from multiple satellites, with different responses,  
617 spatial resolutions, and capture cycles; aiming at improving dependability  
618 and producing estimates with higher resolution in both space and time.

## 619 **Acknowledgment**

620 We gratefully acknowledge the support of NVIDIA Corporation with the  
621 donation of the Titan Xp GPU used for this research. The authors would like  
622 to thank also the research Brazilian agencies CNPq, CAPES and FAPESP.

## 623 **References**

- 624 Ba, R., Chen, C., Yuan, J., Song, W., Lo, S., 2019. SmokeNet: Satellite  
625 Smoke Scene Detection Using Convolutional Neural Network with Spatial  
626 and Channel-Wise Attention. *Remote Sensing* 11, 1–22. doi:10.3390/  
627 rs11141702.
- 628 Ban, Y., Zhang, P., Nascetti, A., Bevington, A.R., Wulder, M.A., 2020.  
629 Near Real-Time Wildfire Progression Monitoring with Sentinel-1 SAR  
630 Time Series and Deep Learning. *Scientific Reports* 10. doi:10.1038/  
631 s41598-019-56967-x.
- 632 Bermudez, J.D., Happ, P.N., Q's, R., 2019. Synthesis of Multispectral Op-  
633 tical Images From SAR/Optical Multitemporal Data Using Conditional

634 Generative Adversarial Networks 16, 1220–1224. doi:10.1109/LGRS.2019.  
635 2894734.

636 Boschetti, M., Busetto, L., Manfron, G., Laborte, A., Asilo, S., Pazhanivelan,  
637 S., Nelson, A., 2017. Phenorice: A method for automatic extraction of  
638 spatio-temporal information on rice crops using satellite data time series.  
639 Remote Sensing of Environment (Elsevier) 194, 347 – 365. doi:10.1016/  
640 j.rse.2017.03.029.

641 Cardil, A., Monedero, S., Ramirez, J., Silva, C.A., 2019. Assessing and reini-  
642 tializing wildland fire simulations through satellite active fire data. Jour-  
643 nal of Environmental Management (Elsevier) 231, 996–1003. doi:https:  
644 //doi.org/10.1016/j.jenvman.2018.10.115.

645 Chapelle, O., Scholkopf, B., Zien, A., 2010. Semi-Supervised Learning. 1st  
646 ed., The MIT Press. doi:10.5555/1841234.

647 Chen, D., Loboda, T.V., Hall, J.V., 2020. A systematic evaluation of in-  
648 fluence of image selection process on remote sensing- based burn severity  
649 indices in north american boreal forest and tundra ecosystems. ISPRS  
650 Journal of Photogrammetry and Remote Sensing (Elsevier) 159, 63 – 77.  
651 doi:10.1016/j.isprsjprs.2019.11.011.

652 Chinchor, N., Sundheim, B., 1993. Muc-5 evaluation metrics, in: Confer-  
653 ence on Message Understanding, Association for Computational Linguis-  
654 tics, USA. p. 6978. doi:10.3115/1072017.1072026.

655 Chuvieco, E., Mouillot, F., van der Werf, G.R., San Miguel, J., Tanase,  
 656 M., Koutsias, N., Garca, M., Yebra, M., Padilla, M., Gitas, I., Heil, A.,  
 657 Hawbaker, T.J., Giglio, L., 2019. Historical background and current devel-  
 658 opments for mapping burned area from satellite earth observation. *Remote*  
 659 *Sensing of Environment (Elsevier)* 225, 45 – 64. doi:10.1016/j.rse.2019.  
 660 02.013.

661 Everingham, M., Gool, L., Williams, C.K., Winn, J., Zisserman, A., 2010.  
 662 The Pascal Visual Object Classes (VOC) Challenge. *International Journal*  
 663 *of Computer Vision (IJCV)* 88, 303338. doi:10.1007/s11263-009-0275-4.

664 Ferreira, L.N., Vega-Oliveros, D.A., Zhao, L., Cardoso, M.F., Macau, E.E.,  
 665 2020. Global fire season severity analysis and forecasting. *Computers &*  
 666 *Geosciences (Elsevier)* 134, 104339. doi:10.1016/j.cageo.2019.104339.

667 Flannigan, M.D., Haar, T.H.V., 1986. Forest fire monitoring using NOAA  
 668 satellite AVHRR. *Canadian Journal of Forest Research* 16, 975–982.  
 669 doi:10.1139/x86-171.

670 Fritsch, J., Khnl, T., Geiger, A., 2013. A new performance measure and  
 671 evaluation benchmark for road detection algorithms, in: *International*  
 672 *IEEE Conference on Intelligent Transportation Systems*, pp. 1693–1700.  
 673 doi:10.1109/ITSC.2013.6728473.

674 Garcia, A., Orts-Escolano, S., Oprea, S., Villena-Martinez, V., Martinez-  
 675 Gonzalez, P., Garcia-Rodriguez, J., 2018. A survey on deep learning tech-

676     niques for image and video semantic segmentation. *Applied Soft Comput-*  
 677     ing (Elsevier) 70, 41 – 65. doi:10.1016/j.asoc.2018.05.018.

678     Gargiulo, M., Dell’Aglia, D.A.G., Iodice, A., Riccio, D., Ruello, G., 2019.  
 679     A CNN-Based Super-Resolution Technique for Active Fire Detection on  
 680     Sentinel-2 Data, in: *PhotonIcs Electromagnetics Research Symposium*  
 681     (Spring), pp. 418–426. doi:10.1109/PIERS-Spring46901.2019.9017857.

682     Giglio, L., Descloitres, J., Justice, C.O., Kaufman, Y.J., 2003. An enhanced  
 683     contextual fire detection algorithm for modis. *Remote Sensing of Environ-*  
 684     ment (Elsevier) 87, 273 – 282. doi:10.1016/S0034-4257(03)00184-6.

685     Giglio, L., Schroeder, W., Justice, C.O., 2016. The collection 6 MODIS active  
 686     fire detection algorithm and fire products. *Remote Sensing of Environment*  
 687     (Elsevier) 178, 31–41. doi:10.1016/j.rse.2016.02.054.

688     Goodfellow, I., Bengio, Y., Courville, A., 2016. *Deep Learning*. The MIT  
 689     Press.

690     Ji, Y., Stocker, E., 2002. Seasonal, intraseasonal, and interannual variability  
 691     of global land fires and their effects on atmospheric aerosol distribution.  
 692     *Journal of Geophysical Research: Atmospheres* 107, ACH 10–1–ACH 10–  
 693     11. doi:10.1029/2002JD002331.

694     Kaufman, Y.J., Justice, C.O., Flynn, L.P., Kendall, J.D., Prins, E.M., Giglio,  
 695     L., Ward, D.E., Menzel, W.P., Setzer, A.W., 1998. Potential global fire

696 monitoring from eos-modis. *Journal of Geophysical Research: Atmo-*  
697 *spheres* 103, 32215–32238. doi:10.1029/98JD01644.

698 Kondratyev, K.Y., Dyachenko, L., Binenko, V., Chernenko, A., 1972. De-  
699 tection of Small Fires and Mapping of Large Forest Fires by Infrared Im-  
700 agery, in: *International Symposium on Remote Sensing of Environment*,  
701 p. 12971303.

702 Kumar, S.S., Roy, D.P., 2018. Global operational land imager Landsat-8  
703 reflectance-based active fire detection algorithm. *International Journal of*  
704 *Digital Earth* 11, 154–178. doi:10.1080/17538947.2017.1391341.

705 Langford, Z., Kumar, J., Hoffman, F., 2018. Wildfire Mapping in Interior  
706 Alaska Using Deep Neural Networks on Imbalanced Datasets, in: *IEEE*  
707 *International Conference on Data Mining Workshops*, pp. 770–778. doi:10.  
708 1109/ICDMW.2018.00116.

709 Lecun, Y., Bengio, Y., Hinton, G., 2015. Deep learning. *Nature* 521, 436–444.  
710 doi:10.1038/nature14539.

711 Lee, T.F., Tag, P.M., 1990. Improved Detection of Hotspots using the  
712 AVHRR 3.7-um Channel. *Bulletin of the American Meteorological Society*  
713 71, 1722 – 1730. doi:10.1175/1520-0477.

714 Lin, T.Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár,  
715 P., Zitnick, C.L., 2014. Microsoft COCO: Common Objects in Context,  
716 in: *Fleet, D., Pajdla, T., Schiele, B., Tuytelaars, T. (Eds.), European*

- 717 Conference on Computer Vision (ECCV), Springer, Cham. pp. 740–755.  
718 doi:10.1007/978-3-319-10602-1\_48.
- 719 Litjens, G., Kooi, T., Bejnordi, B.E., Setio, A.A.A., Ciompi, F., Ghafoorian,  
720 M., van der Laak, J.A., van Ginneken, B., Snchez, C.I., 2017. A survey on  
721 deep learning in medical image analysis. *Medical Image Analysis (Elsevier)*  
722 42, 60 – 88. doi:10.1016/j.media.2017.07.005.
- 723 Lyu, Y., Bai, L., Huang, X., 2019. Road Segmentation using CNN and  
724 Distributed LSTM, in: *IEEE International Symposium on Circuits and*  
725 *Systems*, pp. 1–5. doi:10.1109/ISCAS.2019.8702174.
- 726 Ma, L., Liu, Y., Zhang, X., Ye, Y., Yin, G., Johnson, B.A., 2019a. Deep learn-  
727 ing in remote sensing applications: A meta-analysis and review. *ISPRS*  
728 *Journal of Photogrammetry and Remote Sensing (Elsevier)* 152, 166–177.  
729 doi:10.1016/j.isprsjprs.2019.04.015.
- 730 Ma, L., Liu, Y., Zhang, X., Ye, Y., Yin, G., Johnson, B.A., 2019b. Deep  
731 learning in remote sensing applications: A meta-analysis and review. *IS-*  
732 *PRS Journal of Photogrammetry and Remote Sensing (Elsevier)* 152, 166  
733 – 177. doi:10.1016/j.isprsjprs.2019.04.015.
- 734 Maier, S.W., Russell-Smith, J., Edwards, A.C., Yates, C., 2013. Sensitivity  
735 of the modis fire detection algorithm (mod14) in the savanna region of  
736 the northern territory, australia. *ISPRS Journal of Photogrammetry and*

737 Remote Sensing (Elsevier) 76, 11 – 16. doi:10.1016/j.isprsjprs.2012.  
738 11.005. terrestrial 3D modelling.

739 Malambo, L., Heatwole, C.D., 2020. Automated training sample definition  
740 for seasonal burned area mapping. ISPRS Journal of Photogrammetry and  
741 Remote Sensing (Elsevier) 160, 107 – 123. doi:10.1016/j.isprsjprs.  
742 2019.11.026.

743 Mateo-García, G., Laparra, V., López-Puigdollers, D., Gómez-Chova, L., 2020.  
744 Transferring deep learning models for cloud detection between landsat-  
745 8 and proba-v. ISPRS Journal of Photogrammetry and Remote Sensing  
746 (Elsevier) 160, 1 – 17. doi:10.1016/j.isprsjprs.2019.11.024.

747 Matson, M., Holben, B., 1987. Satellite detection of tropical burning in  
748 brazil. International Journal of Remote Sensing 8, 509–516. doi:10.1080/  
749 01431168708948657.

750 Minetto, R., Segundo, M.P., Sarkar, S., 2019. Hydra: An ensemble of con-  
751 volutional neural networks for geospatial land classification. IEEE Trans-  
752 actions on Geoscience and Remote Sensing 57, 6530–6541. doi:10.1109/  
753 TGRS.2019.2906883.

754 Morisette, J.T., Giglio, L., Csiszar, I., Justice, C.O., 2005. Validation  
755 of the modis active fire product over southern africa with aster data.  
756 International Journal of Remote Sensing 26, 4239–4264. doi:10.1080/  
757 01431160500113526.

758 Murphy, S.W., de Souza Filho, C.R., Wright, R., Sabatino, G., Correa  
 759 Pabon, R., 2016. Hotmap: Global hot target detection at moderate spa-  
 760 tial resolution. *Remote Sensing of Environment (Elsevier)* 177, 78 – 88.  
 761 doi:10.1016/j.rse.2016.02.027.

762 O. Pinheiro, P.O., Collobert, R., Dollar, P., 2015. Learning to Segment  
 763 Object Candidates, in: Cortes, C., Lawrence, N., Lee, D., Sugiyama, M.,  
 764 Garnett, R. (Eds.), *Advances in Neural Information Processing Systems*  
 765 (NIPS), pp. 1990–1998. doi:10.5555/2969442.2969462.

766 Paoletti, M., Haut, J., Plaza, J., Plaza, A., 2019. Deep learning classifiers  
 767 for hyperspectral imaging: A review. *ISPRS Journal of Photogrammetry*  
 768 *and Remote Sensing (Elsevier)* 158, 279 – 317. doi:10.1016/j.isprsjprs.  
 769 2019.09.006.

770 Petersson, H., Gustafsson, D., Bergstrom, D., 2016. Hyperspectral image  
 771 analysis using deep learning a review, in: *International Conference on*  
 772 *Image Processing Theory, Tools and Applications*, pp. 1–6. doi:10.1109/  
 773 IPTA.2016.7820963.

774 Pinto, M.M., Libonati, R., Trigo, R.M., Trigo, I.F., DaCamara, C.C., 2020. A  
 775 deep learning approach for mapping and dating burned areas using tempo-  
 776 ral sequences of satellite images. *ISPRS Journal of Photogrammetry and*  
 777 *Remote Sensing (Elsevier)* 160, 260 – 274. doi:10.1016/j.isprsjprs.  
 778 2019.12.014.

- 779 Portillo-Quintero, C., Sanchez-Azofeifa, A., Marcos do Espirito-Santo, M.,  
780 2013. Monitoring deforestation with modis active fires in neotropical dry  
781 forests: An analysis of local-scale assessments in mexico, brazil and bolivia.  
782 Journal of Arid Environments (Elsevier) 97, 150–159. doi:[https://doi.](https://doi.org/10.1016/j.jaridenv.2013.06.002)  
783 [org/10.1016/j.jaridenv.2013.06.002](https://doi.org/10.1016/j.jaridenv.2013.06.002).
- 784 Ronneberger, O., P.Fischer, Brox, T., 2015. U-Net: Convolutional Net-  
785 works for Biomedical Image Segmentation, in: Medical Image Computing  
786 and Computer-Assisted Intervention, Springer. pp. 234–241. doi:[10.1007/](https://doi.org/10.1007/978-3-319-24574-4_28)  
787 [978-3-319-24574-4\\_28](https://doi.org/10.1007/978-3-319-24574-4_28).
- 788 Roy, D., Wulder, M., Loveland, T., et al., 2014. Landsat-8: Science and  
789 product vision for terrestrial global change research. Remote Sensing of  
790 Environment (Elsevier) 145, 154 – 172. doi:[10.1016/j.rse.2014.02.001](https://doi.org/10.1016/j.rse.2014.02.001).
- 791 Rumelhart, D.E., Hinton, G.E., Williams, R.J., 1986. Learning Represen-  
792 tations by Back-propagating Errors. Nature 323, 533–536. doi:[10.1038/](https://doi.org/10.1038/323533a0)  
793 [323533a0](https://doi.org/10.1038/323533a0).
- 794 Schroeder, W., Oliva, P., Giglio, L., Csiszar, I.A., 2014. The new viirs  
795 375m active fire detection data product: Algorithm description and ini-  
796 tial assessment. Remote Sensing of Environment (Elsevier) 143, 85 – 96.  
797 doi:[10.1016/j.rse.2013.12.008](https://doi.org/10.1016/j.rse.2013.12.008).
- 798 Schroeder, W., Oliva, P., Giglio, L., Quayle, B., Lorenz, E., Morelli, F.,

799 2016. Active fire detection using Landsat-8/OLI data. *Remote Sensing of*  
800 *Environment (Elsevier)* 185, 210 – 220. doi:10.1016/j.rse.2015.08.032.

801 Wang, Z., Chen, J., Hoi, S.C.H., 2020. Deep learning for image super-  
802 resolution: A survey. *IEEE Transactions on Pattern Analysis and Machine*  
803 *Intelligence* 1, 1–22. doi:10.1109/TPAMI.2020.2982166.

804 Yao, G., Lei, T., Zhong, J., 2019. A review of convolutional-neural-network-  
805 based action recognition. *Pattern Recognition Letters (Elsevier)* 118, 14 –  
806 22. doi:10.1016/j.patrec.2018.05.018.

807 Yuan, Q., Shen, H., Li, T., Li, Z., Li, S., Jiang, Y., Xu, H., Tan, W., Yang, Q.,  
808 Wang, J., Gao, J., Zhang, L., 2020. Deep learning in environmental remote  
809 sensing: Achievements and challenges. *Remote Sensing of Environment*  
810 *(Elsevier)* 241, 111716. doi:10.1016/j.rse.2020.111716.

811 Zhu, X.X., Tuia, D., Mou, L., Xia, G., Zhang, L., Xu, F., Fraundorfer, F.,  
812 2017. Deep learning in remote sensing: A comprehensive review and list  
813 of resources. *IEEE Geoscience and Remote Sensing Magazine* 5, 8–36.  
814 doi:10.1109/MGRS.2017.2762307.