

Interactive comment on “A weekly, near real-time dataset of the probability of large wildfire across western US forests and woodlands” by Miranda E. Gray et al.

Miranda E. Gray et al.

miranda@csp-inc.org

Received and published: 12 April 2018

R = Reviewer Comment, A = Author Comment

R: While this manuscript presents a concept that would be valuable (near real-time large fire probability), I believe it oversells its novelty over existing products and I am suspicious that there are methodological flaws (but there is not enough information in the methods section to tell for certain). As the topical editor has already noted, the products are not available at either Figshare or Google Cloud Storage (I did not try the Google Earth Engine).

Printer-friendly version

Discussion paper



A: We've updated the access control list on all assets to public READ (anyone on the internet has read access to the objects). Products are now available at all these locations, and we apologize for the links not working initially.

R: While the authors did locate other near real-time products that deliver similar information (e.g. Preisler et al 2016), they missed other datasets currently commonly used for planning fuel treatments and other management activity (notably Short et al's "Spatial dataset of probabilistic wildfire risk components for the conterminous United States" at <https://www.fs.usda.gov/rds/archive/Product/RDS-2016-0034/>, which was calibrated to previous fires using Short's Fire Occurrence Database which also provides part of the underpinnings of this current effort). This dataset and similar efforts using the FSim model (Finney et al 2011) already deliver many of the functions that the current effort proposes it is uniquely positioned to fill, including acting as "a foundational dataset for longer-term planning and research, such as strategic targeting of fuels management" (e.g. Scott et al 2016, Thompson et al 2017), "fire-smart development at the wildland urban interface" (e.g Haas et al 2013), and "analysis of trends in wildfire potential over time" (Finney et al 2011, then Short et al's updated dataset, and also the fire potential datasets the authors reference by Dillon). Given the mature state of this previous work and its prevalence, its inclusion in the current manuscript would seem important.

A: We've inserted these important references throughout the manuscript and where appropriate, but also attempted to make it more clear that our effort is not unique in filling these positions, as we had previously implied, but still offers complementary advantages over current efforts. For example, this dataset offers an advantage of being able to more immediately draw on updated data, while maintaining the ability to look at probabilistic exposure across large extents and a high resolution. Specifically, please see pg 2, lines 10-19 and pg 10, lines 15-20.

R: In addition, the authors state that other existing models do not account for long-term fuel and climate variability, but I'm not convinced their model better accounts for these due to a number of reasons: 1) their model uses a number of variables as proxies for

[Printer-friendly version](#)[Discussion paper](#)

fuel and weather that as far as I know have not been demonstrated in the literature to relate to fuel availability and flammability or to meaningfully measure weather's effect on fire, including PDSI (which in fact has been demonstrated to not be related to large fire activity (Riley et al 2013)), EVI, and NDWI,

A: While it is true that PDSI was not found to be strongly related to large fire occurrence in Riley et al. 2013, this was independent of ecoregion, bioclimatic zones, vegetation type, and any other predictor variables. There is another study for the western US that found PDSI in the fire season to be correlated to fire activity in specific forested ecoregions (e.g. Abatzoglou and Kolden, 2013). While we did not model large fire probability specific to ecoregions, there are interactions with other predictor variables (e.g. bioclimatic and proxy fuel variables), that might have still made this an important predictor.

A: We were incorrect in stating that NDWI can be used directly for live fuel moisture without also accounting for the amount of vegetation, and also incorrect in stating that EVI was used as a proxy for fuel availability. We've added some text and references to clarify that we are using EVI and NDWI as proxies for long-term biomass production and vegetation dynamics, canopy water content, and live fuel moisture when coupled with EVI. We also clarify that we use these variables on a short-term basis as a proxy for vegetation abundance and condition, and they have been shown to help predict fire. Please see pg 5, lines 18-32 and pg 6, lines 28-32.

R: as far as I can tell, their model has a static development layer (the CSP from 2001 and 2011), thus not capturing development trends except coarsely

A: Yes, this is the most comprehensive dataset of anthropogenic development (including urban development, transportation, energy, and agriculture) covering our study extent, and is unfortunately not yet available at more frequent intervals. While it would be ideal to have an annual development layer, other proxies that are directly attributable to human development and available at an annual timescale, such as nightlights from

[Printer-friendly version](#)[Discussion paper](#)

the VIIRS satellite sensor, would not be as comprehensive and reliable of a predictor.

R: it's unclear how their model captures prior burns, which they state are important factors in fire spread (perhaps they mean the EVI to do this, and indeed it might, but citations or analysis are needed to show that the EVI captures differences in the pre- and post-fire landscapes).

A: We did intend for the EVI to capture this, and have added a citation and text on Pg 5, lines 19-21.

R: Other functions that they state their model can accomplish are already delivered during wildfire incidents by a suite of models in the Wildland Fire Decision Support System, including FlamMap, FARSITE, and FSPRO – which do in fact function in real time, with runs taking on the order of 15 minutes to 1.5 hours and being delivered the same day to suppression forces.

A: We agree that these functions are already provided in these other tools. We have clarified that the functions and advantages of our current effort can be both different and complementary to the functions of the WFDSS tools. For example, this data is not meant to be used on individual fires or at a local scale, but can still provide near real time data at a high resolution and broad spatial extents. Please see pg 2, lines 11-19.

R: However, the approach taken here is novel and could be a valuable addition to the suite of existing models – if methodological questions can be addressed.

Moving on to model structure, I can't discern from the manuscript what the response variable in random forests is. Is it probability of large fire? Probability of small fire? Wouldn't these two be related and increase together? Is probability that the pixel doesn't burn calculated? Or is it simply predicting binary large versus small fire? Or to predict binary burning by large fire versus not burned?

A: We've clarified the model structure, including the response variable. The response variable was binary large ('1') vs. small ('0') fire, resulting in a classification of the

probability that an area would burn in a large fire. See pg 4, starting line 1.

R: Also, it's unclear how predictor variables were chosen. Did the authors choose a set of variables based on their hunches about what's important? Or was variable importance in random forests used to guide selection?

A: Variables were chosen based on previous literature and hypotheses of what drives large fire. Random forest is a predictive framework that does well with many, often correlated predictors, and can uncover some complex interactions between predictors that need not be specified a priori (see Cutler et al., 2007). Therefore, variables were also chosen independently and based on the fact that their interactions may help predict large fire (eg, EVI and NDWI).

R: I am concerned about a potential flaw in model design: it appears that the authors model probability of each pixel independently – however, probability of burning is not independent and is affected by contagion from neighboring pixels. Each pixel's burning can't be considered an independent event, since burning is spatially related to its neighbors. To that end, each large fire might properly be regarded as the unit of prediction, not individual pixels. If the authors have accounted for this, methods should be shared in the paper.

A: We have accounted for the non-independence of individual burning pixels by taking the mean values of predictor variables in a circular kernel with a radius of 1135 meters. This results in a window size that is approximately the size of a large fire (i.e., 405 ha), and assumes that this entire area has an influence on whether a pixel can burn (e.g., fire may spread from a distant source). See pg 5, lines 13-16.

R: The manuscript is largely lacking in assessment of goodness-of-fit of the model.

A: Please see the revised section 'Dataset Evaluation', in which we have added more detail about variable importance, and accuracy of models when evaluated against training and independent testing data.

[Printer-friendly version](#)[Discussion paper](#)

R: Also, is it feasible to predict burn probability without first predicting ignition? Burning is predicated on ignition taking place first, and then on spatial contagion: thus, the burn probability of one pixel wouldn't be independent of the nearby ones. Some discussion of this, demonstrated by out-of-bag error rates and additional goodness-of-fit metrics is needed.

A: We have clarified that we are modeling the probability of an area burning in a large fire, conditional on either an ignition event or fire spreading to that area. This is different from modelling overall burn probability that accounts first for the occurrence of an ignition. We've added a caveat in the text that overall burn probability would be influenced by non-random ignitions, which we have not accounted for in our analysis. See pg 3, lines 17-23.

R: Also, better demarcation of prescribed fires versus wildfires is needed.

A: Large prescribed fires were excluded from the data used in our analyses by only selecting large fire samples from the MODIS BA dataset from within MTBS wildfires. Prescribed fires are not included in the FOD dataset, and thus were not included in the small fires. See pg 4, lines 25-27.

R: Lastly, how were the outputs validated? By comparison with other existing burn probability or fire regime datasets?

A: Outputs were evaluated with the training data and independent testing data (i.e., small and large fires that occurred from 2015-2016.). We did not compare our outputs with existing datasets, but we do believe that future comparisons of this sort will be important.

Specific comments follow:

R: Page 2 Line 4: the sentence that ends "characteristics of fire regimes" would seem to require citations. Please add.

A: We've added citations to pg 2, lines 4-5.

[Printer-friendly version](#)[Discussion paper](#)

R: Page 2 Line 7: also see Finney et al 2011, Preisler et al 2016, and Preisler et al 2004, for example, who have already produced this type of work.

A: We have revised the first paragraph and removed this original line, but included these references elsewhere and where appropriate. See pg.2 lines 21 and 29.

R: Page 2 Line 14: it seems that Finney et al 2011 (FSim paper below) should also be referenced here. The comment that follows is not relevant to FSim (“it requires detailed specification of many model inputs and is highly sensitive to misspecification of these parameters”), which is calibrated to fire occurrence in Short’s FOD.

A: Finney et al. 2011 is now referenced here. We believe that the following comment is relevant to FSim, since it is input-intensive (requiring a fuel model assignment itself with many parameters, fuel moisture estimation, and weather parameters). The underlying simulation models (i.e., FlamMap) are sensitive to all of these parameters.

R: Page 2 Line 16-18: I am confused by this comment, since as noted above, FlamMap runs in seconds, FARSITE in minutes, and FSPro in 15-90 minutes, meaning that they can and are updated subdaily during active fires.

A: This comment was meant more in reference to models such as FSim, but we have also clarified what we see as the limitations of this type of model and more real-time models like FARSITE. See pg 2, lines 11-19.

R: Page 2 Line 19-20: these lines state that models like Preisler et al 2016 are constrained by availability of accurate high-resolution fire, weather and fuels data – however, later in the manuscript the authors correctly acknowledge that the Preisler et al 2016 model was run daily last year with updated weather data, with outputs available on the WFAS website.

A: This was meant to imply that the spatial or temporal scale is constrained in such models (e.g., sacrificing spatial scale for temporal scale, or vice versa), but not necessarily the temporal scale alone. We have clarified on pg 2 lines 30-31.

[Printer-friendly version](#)[Discussion paper](#)

R: Page 3 Line 24-26: I'm curious how fuel type (grass, brush, timber litter, etc) is accounted for in the model, as fuel type is directly related to fire spread and probability.

A: Fuel type was not directly accounted for in the models. Several of the predictor variables are meant to provide meaningful approximations of water and energy balances that determine vegetation types (i.e., bioclimatic predictors, topographic aspect, elevation, and slope). Long-term remotely sensed indices of EVI and NDWI are meant to provide inter-annual summaries of the amount and dryness of vegetation, which also helps differentiate vegetation and fuel types. See pg 5, lines 18-32.

R: Page 4 Line 6: What are "large fire event days"? I'm thinking you mean an individual burned pixel. Please elaborate in this paragraph on how you decided whether burned pixels were part of the same fire – from Figure 1, it looks like you assigned pixels to MTBS fires.

A: We removed this term and have also clarified how we used MTBS to constrain large fire sampling. See pg 4, lines 25-27.

R: Page 4 Line 23: what is meant by "there are methods that may be adapted to associate active fire information with small fire events"?

A: We determined that this sentence was unnecessary and removed it, since it was only meant to suggest a potential different approach in the future.

R: Page 4 Line 24: Why was it important to have the same number of small and large fires? There are many many more small fires than large fires in Short's FOD.

A: This was done to maintain a balance in the sampling ratio so that the models were not biased towards predicting small fires (see Breiman, 2001). Because training RF models takes bootstrapped resamples of the data to build trees, when there is extremely imbalanced data, there is a significant probability that a bootstrap sample contains few or even none of the minority class. This would result in a tree with poor performance for predicting the minority class. There is then the risk of loss of infor-

[Printer-friendly version](#)[Discussion paper](#)

mation from the many small fires, but we have partially attempted to overcome this by taking 10 random sample seeds.

R: Page 4 Line 25-26: There are other (perhaps more effective) ways to remove prescribed burns from your dataset. For example, fire type is an attribute in MTBS. By using April-October fires only, you'd include most Rx burns in northern states like Montana and Idaho, and exclude large southern California fires like the recent Thomas Fire that often take place in December.

A: See response to comment above and pg 4, lines 25-27 and pg 5, line 1.

R: Page 4 Line 30: I must confess I'm not familiar with the NDWI, but when I Google it, USGS calls it the "Normalized Difference Water Index" rather than wetness index, and describes it as being used to discern water from non-water. Are you talking about the same index? It seems improbable that there are two MODIS NDWI's with different calculations: : but perhaps that is the case. Please clarify. In either case, I've not seen the NDWI used in any studies relating it to fire occurrence. So it is a good choice here? The relationship of canopy moisture and flammability is quite complex and not well understood (see for example McAllister et al 2012). Despite the lack of study of the NDWI, it could be a good predictor in your model, but not enough information on variable importance is presented in the current version for me to assess.

A: You are correct - Gao (1996) first referred to NDWI as the Normalized Difference Water Index, and we have revised this in the text. Although Gao proposed it for remote sensing of vegetation liquid content, we have included some more recent references that find it to correlate with pixel water content, and live fuel moisture when coupled with EVI. See Pg 5, lines 12-25 and mention of NDWI variable importance on pg 8 line 3.

R: Page 5 Line 4-5: If I am understanding correctly, you only used MODIS values from inception up to the date of the fire to assign percentile values. Why not use the whole record? It seems in your current method, the percentile assignments would be

sensitive to the date of the fire (so if a fire occurs at an index value of 100 in 2005 and another fire at a value of 100 in 2010 in the same pixel, these could be calculated to be different percentile values since the underlying distribution of values would be different).

A: We did not use the whole record because fires affect spectral reflectance of the vegetation and surface. Therefore, post-fire MODIS estimates may not give a reliable and consistent estimate of what is driving the fire of interest.

R: Page 5 Line 6-7: What is the index of human modification supposed to signify with regards to burn probability? Why is it included? What was the variable importance score?

A: We hypothesized that more developed landscapes, because they are less natural and generally more fragmented, were less likely to burn in large fires. We also assumed that suppression resources and mandates were more readily called upon nearer to urban development, and so included an intuitive measure of Euclidean distance to urban development. Human modification or distance to urban variables were consistently in the top 5 important variables across the 10 models.

R: Page 5 Line 20: Can you explain more about why the CV of temperature and precipitation is “seasonality”? I don’t follow. Similarly, why are temperature of the wettest and driest months and precipitation of the coldest and wettest months included as predictors? Have these been demonstrated to correlate with fire probability? Do they have high variable importance scores?

A: Temperature seasonality, which we’ve revised to be the standard deviation of mean monthly temperature, is defined by the amount of temperature variation over a year (see O’Donnell and Ignizio, 2012). Amongst the other bioclimatic predictors you mention, we’ve now only included the temperature of the wettest month, which is meant to describe the coincident interactions of energy and water balances, in the absence of more direct, long-term water balance metrics. This was demonstrated to correlate with

[Printer-friendly version](#)[Discussion paper](#)

fire probability - see updated references in the text on pg 6, line 24.

R: Page 5 Line 23: I think you are saying that EVI is related to fuel availability. I think you are working only in forested ecosystems, so most times EVI will be correlated with the canopy rather than the understory. However, surface fire propagates in the understory and crown fires are relatively rare. So is EVI really related to fuel availability? Also, see McAllister et al regarding live fuel moisture and flammability.

A: We have revised this to clarify that we are using long-term EVI to characterize biomass production, but not fuel availability explicitly. Our hypothesis was that a multi-year time-series of EVI may differentiate between levels of biomass production across the western US, but also interannual, pixel-wise dynamics of vegetation that may help predict large fire. Please see pg 5. lines 18-24.

R: Also in this paragraph, if you have only five years of data in some cases and you are calculating anomalies, you would have only 5 observations, right? Again, why not use the full MODIS record (or did you)? Also, are the average LSTs for both night and daytime temperatures?

A: Yes, for the closest day-of-year anomalies, there was only one observation from each year. Please see the above response for why we did not use the full MODIS record. The LSTs were for daytime temperatures only, and we have clarified this in the text.

R: Page 6 Line 5: Why was PDSI included as a variable when it has been demonstrated not to be strongly correlated with large fire activity (e.g. Riley et al 2013)?

A: Please see response to the PDSI comment above.

R: Page 6 Line 8: Please state which NFDRS fuel model the ERC was calculated for. I believe Abatzoglou's product is for fuel model G.

A: Yes, this was fuel model G.

R: Page 6 Line 11: fm1000 represents the previous 42 days ($1000 \text{ hr}/24 = 41.666$ days).

A: It seems that fm1000 was only calculated by using the previous 7 days (see Schlobohm and Brain, 2002, pg 22), but please correct us if you know this to be different for the GRIDMET dataset.

R: Page 6 Lines 18-29: Were small fires assigned to a single pixel? Please explain why only one year of fires was used in evaluation (do you expect these relationships to be stationary from year to year when there is so much annual variability in area burned?). For the rest of this paragraph and the following paragraph I'm quite confused. I don't understand what the response variable in the model is (as stated above). Also, can you briefly define sensitivity and specificity? If the response variable is probability, how do you define a false negative and false positive?

A: Yes, small fires were assigned to a single pixel. Please see response to comment above for clarification of the response variable. Sensitivity and specificity have been defined in the text. Although predicted probability was extracted at the testing data points, these were assigned a binary predicted response of '0' or '1' based on the optimal cutoff, thereby allowing us to assign false negatives and false positives.

R: Page 7 Line 10: I would like to see each of the bands illustrated by a figure, otherwise it's quite difficult for a reviewer to visualize and assess the product. How prevalent were pixels with a rating of 1 (at least one MODIS pixel was not processed or had bad quality)?

A: Unfortunately, we did not get to this suggestion but really like the idea, and can create this figure for a revised submission, if deemed appropriate.

R: Page 8 Line 7: again, see other literature including but not limited to Thompson et al 2017 and Scott et al 2016.

A: See pg 9, lines 25-26.

[Printer-friendly version](#)[Discussion paper](#)

R: Page 8 Line 11: There are other products updated daily that account for changes in weather and fuel moisture, including Preisler et al 2016. Some of the inputs to your model appear to be static, including the CSP (human development layer) and it's not clear how past disturbances (burns) are included in your model. Perhaps the EVI captures previously burned areas, but I know of no study that documents that. Have you assessed how your model works in recently burned vs. burned areas?

A: Please see responses above related to the human development layer, and EVI to capture prior burns. In future analysis with this dataset, we would like to evaluate how it compares in recently burned areas across the west, but we have not done this evaluation yet.

R: Page 8 Line 20: It's not clear how this model would provide better information to managers during active fires than the suite of models in WFDSS (FlamMap, FARSITE, and FSPro), which output information on predicted fire intensity, fire spread, and burn probability in near real-time. Please clarify.

A: Please see response to one of your first comments above, related to this, and pg 10, lines 15-20.

R: Figure 1: This figure nicely illustrates how incomplete MODIS data is!! I've noticed this while daily following nearby fires in my area. MODIS often misses surface fires where the canopy is dense or even crown fires where the smoke plume is dense. Is MODIS then a good basis for predicting burned pixels (especially when it can be difficult to eliminate Rx fire)?

A: We chose to use MODIS because it provides the day-of-burn, which can then be related more precisely to predictor variables. While it is rather incomplete, we've noticed that it still captures pixels within most of the MTBS perimeters (as illustrated in this figure), and provided enough data for our modeling purposes.

R: Figure 3: Please present actual values rather than "high" or "low". I don't feel I can

[Printer-friendly version](#)[Discussion paper](#)

validate the product without them.

A: We have added the actual values to this figure, and per a recommendation from another reviewer, have only showed one prediction date. We've also included MTBS fires that occurred in the month immediately after this prediction, as another coarse way to evaluate the product.

R: Figure 4: I'm confused here. Why not present at-pixel values? Are these the sum or average of false positives for an ecoregion? I'm also confused as to what these mean: there is always a probability of fire, so what does it mean to have a false negative or false positive if you are predicting probability? Of course, as I said earlier, I'm confused about what the response variable is, so when I understand that perhaps I won't be confused here.

A: We believe that providing at-pixel values would make this figure too hard to read, if we understand you correctly. In an attempt to make it readable and still informative, what we've presented is the rate (e.g., # of false positives/total number of testing data points) in each ecoregion. Although predicted probability was extracted at the testing data points, they were assigned a binary predicted response of '0' or '1' based on the optimal cutoff, thereby allowing us to assign false negatives and false positives. That said, we are very open to revising this figure if deemed appropriate, perhaps by using at-pixel values.

R: Figure 5: I'm confused here too. Is the white squiggly line the probability of small fire and the black squiggly line the probability of large fire? If so, the y-axis is incorrect. Is the vertical white line the date of a small fire, and the black vertical line the date of a large fire? What does it mean to randomly pair a large and small fire? Should they be related? Why in some cases are the black and white trends similar and in some cases different?

A: This figure caused confusion to multiple reviewers, so we have decided to eliminate it from the manuscript.

References:

Abatzoglou, J. T. and Kolden, C. A.: Relationships between climate and macroscale area burned in the western United States, *Int. J. Wildl. Fire*, 22(7), 1003–1020, 2013.

Breiman, L.: Random forests, *Mach. Learn.*, 45(1), 5–32, doi:10.1023/A:1010933404324, 2001.

Cutler, D. R., Edwards, T. C., Beard, K. H., Cutler, A., Hess, K. T., Gibson, J. and Lawler, J. J.: Random Forests for Classification in Ecology, *Ecology*, 88(11), 2783–2792, doi:10.1890/07-0539.1, 2007.

Gao, B. C.: NDWI - A normalized difference water index for remote sensing of vegetation liquid water from space, *Remote Sens. Environ.*, 58(3), 257–266, doi:10.1016/S0034-4257(96)00067-3, 1996.

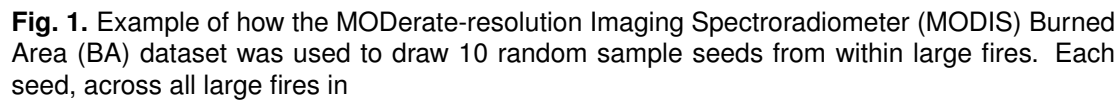
O'Donnell, M. S. and Ignizio, D. A.: Bioclimatic Predictors for Supporting Ecological Applications in the Conterminous United States, *U.S Geol. Surv. Data Ser.* 691, 10, 2012.

Schlobohm, P. and Brain, J.: Gaining an understanding of the National Fire Danger Rating System, 2002.

Please also note the supplement to this comment:

<https://www.earth-syst-sci-data-discuss.net/essd-2017-136/essd-2017-136-AC2-supplement.pdf>

Interactive comment on *Earth Syst. Sci. Data Discuss.*, <https://doi.org/10.5194/essd-2017-136>, 2018.



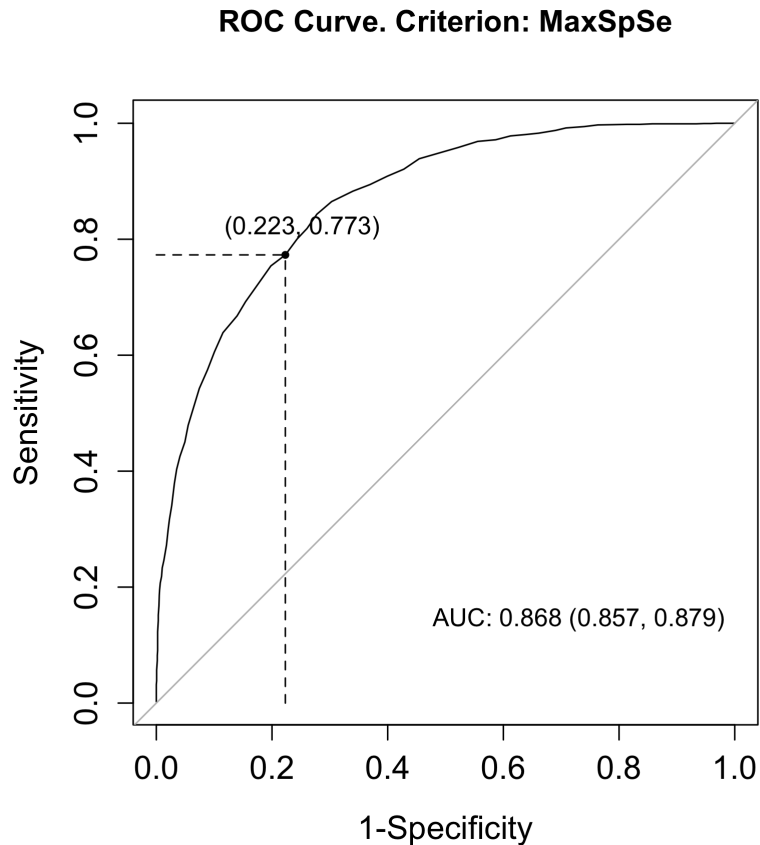


Fig. 2. Receiver Operating Curve (ROC) for an independent testing dataset of small and large fires that occurred from 2015-2016. Sensitivity and (1-Specificity) values are shown for the point where large fire

[Printer-friendly version](#)[Discussion paper](#)

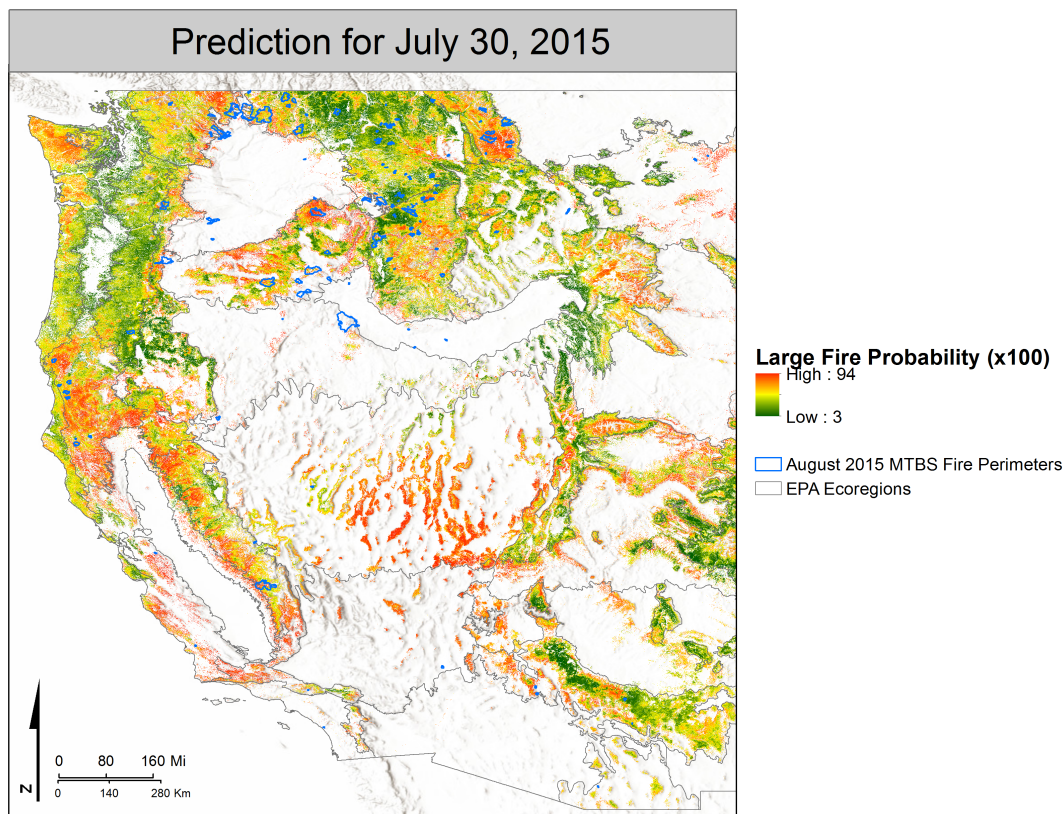


Fig. 3. Large fire probability for the week of July 30, 2015. MTBS fires greater than 405 hectares, and that started in August 2015, are overlaid on the map.

[Printer-friendly version](#)[Discussion paper](#)

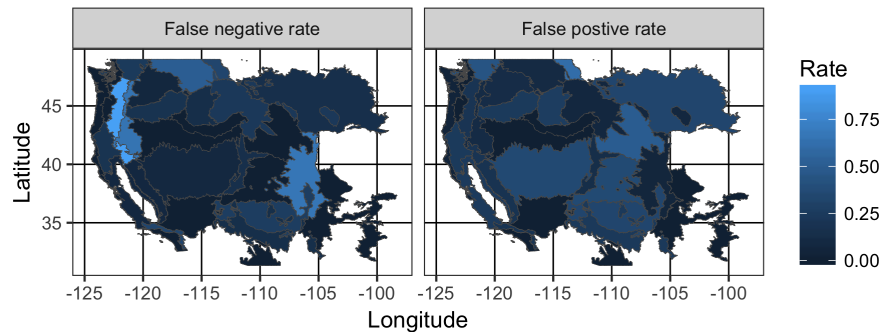


Fig. 4. False positive (FP) and false negative (FN) rates of an independent testing dataset of small and large fires from 2015-2016, mapped across EPA level three ecoregions. No testing data was available for

[Printer-friendly version](#)[Discussion paper](#)