

Predicting the Onset of Australian Winter Rainfall by Nonlinear Classification

LAURA FIRTH

Data Analysis Australia, Perth, Australia

MARTIN L. HAZELTON

University of Western Australia, Perth, Australia

EDWARD P. CAMPBELL

CSIRO, Perth, Australia

(Manuscript received 17 November 2003, in final form 23 July 2004)

ABSTRACT

A method for predicting the timing of winter rains is presented, making no assumptions about the functional form of any relationships that may exist. Ideas built on classification and regression trees and machine learning are used to develop robust predictive rules. These methods are applied in a case study to predict the timing of winter rain in five farming towns in the southwest of Western Australia. The variables used to construct the model are mean monthly sea surface temperatures (SSTs) over a 72-cell grid in the Indian Ocean, Perth monthly mean sea level pressure (MSLP), and monthly values of the Southern Oscillation index (SOI). A predictive model is constructed from data over the period 1949–99. This model correctly classifies the onset of the winter rains approximately 80% of the time with SST variables proving to be the most important in deriving the predictions. Further analysis indicates a change point in the mid-1970s, a well-known phenomenon in the region. The prediction rates are significantly worse after 1975. Furthermore, the important region of the Indian Ocean, in terms of SSTs for prediction, moves from the Tropics down toward the Southern Ocean after this date.

1. Introduction

Australian rainfall has been the subject of a considerable body of research. The existence of a relationship between Australian rainfall and the El Niño–Southern Oscillation (ENSO) has been well established. Nicholls (1991) defines El Niño as a marked temperature increase that occurs every few years in the Pacific Ocean. El Niño is linked with the Southern Oscillation, which is a global pattern of fluctuations that occur both in the ocean and in the atmosphere. The Southern Oscillation index (SOI) (the standardized difference in pressure between Tahiti and Darwin) is used as a measure of the behavior of ENSO. Nicholls (1985) finds that ENSO is related to many aspects of Australian rainfall, including the onset date of the wet season in north Australia. However, the relationship between the SOI and rainfall is not as strong in the southwest as it is in the north and east of Australia (McBride and Nicholls 1983). Nicholls

(1989) finds that variations in rainfall in the southwest occurred at the same time as changes in Indian Ocean sea surface temperatures (SSTs) and thus suggests that SST may be a factor influencing rainfall.

In the last five years there has been much research into the relationship between Indian Ocean SSTs and rainfall in the southwest. This research has shown that winter rainfall in the southwest is highly correlated with Perth mean sea level pressure (MSLP; Hunt et al. 2000; Nicholls et al. 2000). Nicholls et al. (2000) find that, although there is a weak relationship between variations in Indian Ocean SSTs and rainfall in the southwest, this relationship is most likely due to the relationship between SOI and both SST and rainfall. Their climate models involving SSTs also failed to pick up the decline in mean rainfall that has occurred since the mid-1970s. Thus, they conclude that the variations in SST and rainfall are unlikely to be causally related. However, Campbell et al. (2000) raise a concern about using methods that assume a linear relationship to assess nonlinear climate processes. The dangers of assuming linearity have been noted previously by a number of authors, including Palmer (1999) and Graf and Castanheira (2001). Both of these papers argue

Corresponding author address: Dr. Martin L. Hazelton, School of Mathematics and Statistics, University of Western Australia, M019, 35 Stirling Highway, Crawley, WA 6009, Australia.
E-mail: martin@maths.uwa.edu.au

that the climate system consists of a number of quasi-stationary regimes or modes of variation. This implies that conventional linear methods in effect collapse separate, coherent modes into one mode of variation from which a signal is, at best, difficult to extract. We therefore regard Indian Ocean SST, SOI, and MSLP as all being potential explanatory variables for our non-linear model to predict the onset of the southwest wet season.

In our case study the onset of the wet season is regarded as a categorical variable with three categories: early, mid, or late season onset. We therefore seek to develop a model to classify each wet season as either early, mid, or late according to the values of the covariates (SST, SOI, and MSLP) relating to that year. Our aim here is not to model the underlying mechanisms of the rainfall system but to use statistical techniques to identify important variables that can be used for predictions. These statistical techniques include linear regression, logistic regression (for a two-class problem), discriminant analysis, and classification and regression tree (CART) analysis (Gower 1998). Nicholls (1984) used logistic regression when seeking to predict whether or not the wet season in northern Australia will be late. However, we are faced with far more explanatory variables. The SST, SOI, and MSLP sets all contain monthly averages over 12 months and the SST data are measured over 72 grid cells in the Indian Ocean. The data covers the years 1949–99 and we used five farming towns to represent the southwest region. This means that there are nearly 900 explanatory variables for 50 years of data for five towns. Application of logistic regression would require a huge reduction in the number of predictors, which can result in a loss of the amount of information available for the prediction. Furthermore, it can also lead to a situation where there are a multitude of models that are essentially equally good, but each gives a different idea of the relationship between the response and explanatory variables (Breiman 2001b).

Our preferred classification method is random forests, a technique developed by Breiman (2001a) based on CART analysis (Breiman et al. 1984). This algorithmic approach has a number of advantages over classical model-based classification techniques. A random forest classifier does not assume linearity nor that the data are drawn from a particular distribution. It has the ability to handle a large number of explanatory variables without requiring optimal subsets to be selected. Rather, a random forest averages the results of many classification/regression trees, a process that can improve overall prediction accuracy (Breiman 2001b).

The remainder of this paper is structured as follows. The next section contains an overview of classification trees and the extension to random forests. The prediction error for random forests is calculated using a technique called out-of-bag estimation, which is described.

In section 3, we provide a case study and the particulars of our implementation of random forests. A discussion of the results obtained and some concluding remarks are given in section 4.

2. Methodology

a. Classification trees

Classification trees seek to assign individuals to predetermined classes by means of an hierarchical partition of the data. The aim is to produce a tree-based classifier that will assign any newly observed individual to its correct class with high probability. The fundamental principles are perhaps most easily understood through an illustrative example. Suppose that data on two (explanatory) variables are observed on five individuals, as given in Table 1. These are the training data used to build the classifier. A tree can be constructed by defining a sequence of binary splits with respect to the variables. The result is an increasingly fine partition of the data. Figure 1 provides an example. Here t_1 represents the root node, a group containing all the individuals in the dataset. This node is then branched into two daughter nodes, labeled t_2 and t_3 , by a binary split using the second variable. The splitting process continues until only terminal nodes (i.e., those which are not split further, e.g., t_5) remain.

Three fundamental questions in the construction of the tree are (i) how are the splits selected?; (ii) when should a node become a terminal node?; and (iii) which class should be assigned to each terminal node?

Node splitting: At each node the choice of variable to split by, and the value at which the split occurs, is determined so as to maximize the “purity” of the resulting daughter nodes. Intuitively speaking, impurity is the extent to which there is a mixture of classes among the individuals at a node. There are a number of ways of quantifying node purity, including entropy and the Gini index (e.g., see Hastie et al. (2001, chapter 9). All the results reported in this article were obtained using the Gini index. Note that any given variable may be used to define multiple binary splits on any given branch of the tree, allowing for a very flexible representation of the relationship between class membership and the predictor variables.

TABLE 1. Data for illustrative example.

Individual	Class	Variable 1 (var 1)	Variable 2 (var 2)
1	A	26	374
2	A	22	303
3	B	26	396
4	B	29	377
5	B	20	341

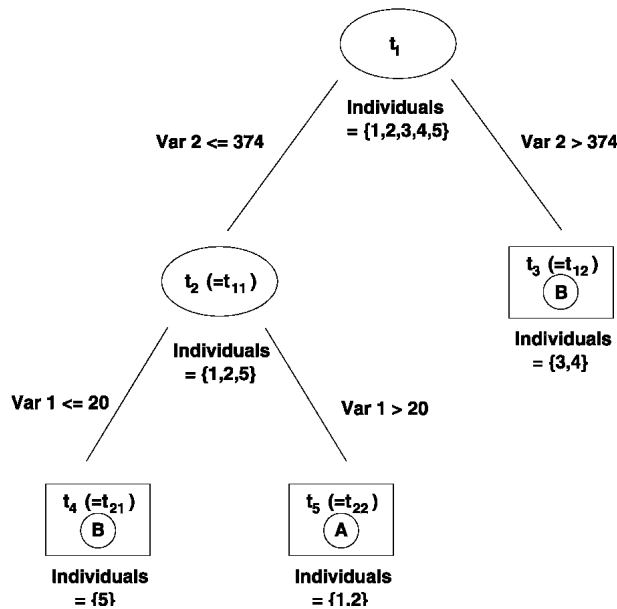


FIG. 1. Classification tree for data in Table 1.

Terminal nodes: A node can be declared terminal for a number of reasons. First, if all the individuals at a node are from the same class (so that the node is pure), then the node is terminal. Second, it is usual to set a minimum node size so that a node cannot be split further if its daughters would be too small. Third, a node will be terminal if there are no possible splits that would decrease the impurity by a “significant” amount.

Terminal node assignment: Each terminal node is assigned a class. Any new observation can then be classified by running it down the tree and giving it the class of the appropriate terminal node. A straightforward approach to this problem is to allocate a class according to majority rule. The example tree in Fig. 1 is particularly simple in that each terminal node is pure and, hence, is assigned the class of a 100% majority of its members.

See Breiman et al. (1984) for further details regarding classification (and regression) tree analysis.

b. Random forests

Classification trees are easy to interpret and implement. However, they tend to be unstable in the sense that, if the original dataset is changed slightly, the resulting tree can be different but with a similar error rate (Breiman 1996a). For example, if the value of variable 2 for the first individual in Table 1 was changed to 378, then a tree very different to that in Fig. 1 would result.

Random forests is a method that seeks to improve the instability of the classification trees. Breiman (1996a) shows that an unstable classifier can be stabilized by aggregating the results over a large number of

these unstable classifiers. The aggregated classifier also has a lower prediction error. In terms of random forests, aggregating means that a large number of trees are grown and each is used to classify the individuals. The class that is most often assigned to each individual is chosen as its final class label. There are many ways in which the multiple trees can be constructed. Some of these methods involve generating new training sets (datasets that are used to fit the model) while others involve altering the internal construction of the tree. We employ *bagging* and *random feature selection* to grow the forest for our analyses.

Bagging (bootstrap aggregating): In principle, we would like to use a large number of training sets, all drawn from the population of interest, to create the random forest. However, this is not possible in practice. Nevertheless, we may approximate this idea by using bootstrap sampling (Efron and Tibshirani 1993). That is, we take samples (with replacement) from the original training set and grow a tree from each of the new samples. We grew forests of 1000 trees using this approach.

Random feature selection: Breiman (1998) shows empirically that, although bagging reduces the variance of the aggregated classifier, it can increase the bias. Random feature selection involves randomly choosing a subset of covariates from which to derive the optimal explanatory variable at each split for each tree. This technique does not reduce the variance as much as bagging, but the bias does not change (Dietterich and Kong 1995). Therefore, the two methods can be combined to produce better results. Breiman (2001a) empirically showed that bagging improves the accuracy of the prediction error when random feature selection is used.

The performance of a classifier is determined by its misclassification rate. If this rate is calculated directly from all the data to which the forest is fitted, then this will give an optimistic estimate of the true prediction error because the forest will be fine-tuned to the particular training data at hand. A much better approach is out-of-bag estimation, which utilizes for each tree those individuals that were not selected for the bootstrap sample.

Each tree is used to classify members of its out-of-bag sample, with final classification for an individual being determined by majority vote over all trees in the forest. The forest classification is compared with the true group for each individual, allowing a misclassification rate to be computed. Empirical testing shows that this out-of-bag estimate is generally very close to the value calculated with truly new data (Breiman 1996b).

c. Measuring variable importance

Although random forests produce a better prediction error than classification trees, they are not as easy to

interpret. It can therefore be difficult to see the relationship between the response and explanatory variables. One way that these relationships can be examined in the random forest is by determining the variables that were important in terms of predicting the response variable (Breiman and Cutler 2001). The importance of a variable indicates the magnitude of its effect in producing accurate predictions. There are a number of ways that this importance can be measured. Some popular ones measure how much the predictions change when the values of a given variable for each individual in the out-of-bag sample are randomly permuted. The more important the variable is for prediction, then the greater the effect of such spoiling on the misclassification rate. An alternative approach is to consider the effect that each variable has on the node impurity measure when each tree is constructed.

d. Practical implementation

Our implementation of random forests was done using the statistical package R (Ihaka and Gentleman 1996, with library: randomForest) running under Linux on a 1.9GHz Pentium 4 personal computer. We used the default settings for most tuning parameters in our work, but set the forest size to 1000 trees (rather than the default of 100) in order to stabilize estimates of misclassification rates. We investigated the effect of changing other tuning parameters, and in particular the terminal node size, as recommended by (Breiman and Cutler 2001). We found our results to be very robust to changes around the default values. In assessing variable importance we examined the four measures implemented in the randomForest library, the first three of which are based on permutation methods and the last of which employs a node impurity measure.

3. Case study

Our case study is focused on the wheat growing region of southwest Western Australia. The region experiences a “mediterranean” climate with mild, wet winters and hot, dry summers. On average 80% of rainfall occurs in the period from May to October with the wettest months being June and July. The winter season is a very important time of year for farmers in the region planting crops. The onset of the winter rains is difficult to predict, although there are many rules of thumb in the farming community. This makes it difficult and financially risky to determine appropriate sowing times for crop plants. The aim of this study is to develop a model to predict the timing of the winter rains. From a practical viewpoint this work is important not only for farmers but for the whole country, as Western Australian agriculture makes a significant contribution to the nation’s economy. In 2000/01, the agricultural production value was \$4.664 billion, with \$3.802 billion of this

exported. Cereal crops made up 54% of these exports (Robertson 2001). From a theoretical perspective this study has the potential to improve our general understanding of southwestern Australian climate.

a. Data and study area

As mentioned in the introduction, our aim in this case study is to predict the onset of the wet season for five Western Australian farming towns, namely Corrigin, Merredin, Geraldton, Wongan Hills, and Lake Grace (see Fig. 2). For data from 1949 to 1999 we used soil moisture levels to determine the onset of winter rainfall, rather than the rainfall level itself, because this gave a better indication of when there was enough rainfall to penetrate the soil and aid crop growth. In the original data the onset date was set as the date of the beginning of (approximately) continuous positive soil moisture. We chose to categorize the onset date (in line with the recommendations of Nicholls 1984), with the result that April was defined as early season onset, May as midseason, and June as late season.

The explanatory variables in our analysis are the average monthly SSTs, SOI, and MSLP from 1948 to 1999. The SSTs came from version 2 of the Global Sea Ice and Sea Surface (GISST2) dataset and covered a region of the Indian Ocean ranging between the Australian and African coasts and extending from the Southern Ocean to the Tropics (9.5° – 39.5° S, 50.5° – 110.5° E). This area was divided in a 6×12 grid of cells, each $5^{\circ} \times 5^{\circ}$, for which SST was recorded. MSLP data for Perth were supplied by the Australian Bureau of Meteorology.

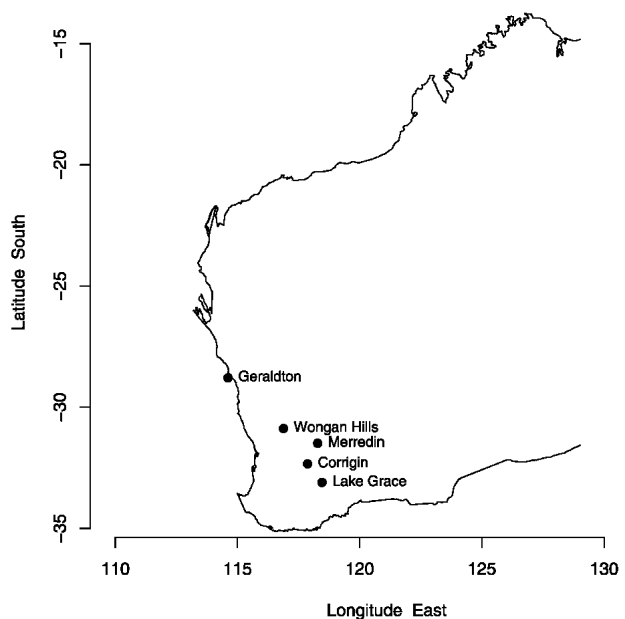


FIG. 2. Map of Western Australia showing the case study locations.

b. Results

For each year from 1949 to 1999, SST, SOI, and MSLP for the 12 months prior to the onset date were included as potential predictors. As the 12 months prior to the onset differed according to the onset class, we set up two slightly different prediction problems. Problem 1 is to predict in March whether or not the onset will be early using explanatory variables dating back to April of the previous year. For those years for which the onset is not early, problem 2 is to predict in April whether the onset will be mid or late season using explanatory variables dating back to May of the previous year. The results from both problems were combined to compute the misclassification rate, with a year being classified correctly only if the classification for problems 1 and 2 (if required) were both correct.

The fact that we have multiple response variables (i.e., five, rather than one, towns for which we seek predictions) needs to be addressed. Originally we considered two approaches. First, we considered the prediction problems for each town separately. In the second approach we constructed a single random forest for prediction in all towns simultaneously, but included an indicator variable for each town as an explanatory variable. The latter approach provided better predictions, as one might expect, because it allowed the data from one town to “borrow strength” from data for the other towns. Only Geraldton, which is coastal and located at a much higher latitude than the other towns, showed no difference in misclassification rate between the two approaches.

We shall refer to our initial random forest using all predictor variables (SST, SOI, and MSLP, as described above) as Classifier A. This classifier gave an overall misclassification rate, and town-specific rates, as given in Table 2. The random forest predictions have an accuracy rate of around 80% overall, with Geraldton suffering in comparison to the other towns. This compares very favorably to predictions using a single classification tree, which had an accuracy rate (estimated by leave-one-out cross-validation) of 57.3%, and classification using principal components logistic regression, which had an accuracy rate of 61.6%. The measures of variable importance for Classifier A provided some evidence that SST was more influential in determining predictions than SOI or MSLP, but the picture was not

TABLE 2. Misclassification rates (MRs) for Classifier A as a percentage.

Town	MR (%)
All	20.39
Corrigin	15.69
Geraldton	27.45
Merredin	19.61
Wongan Hills	19.61
Lake Grace	19.61

TABLE 3. Overall MRs for Classifiers B and C.

Classifier	MR (%)
B	20.39
C	20.78

entirely clear. To investigate this issue further, we created two new random forest classifiers.

- **Classifier B:** Random forest using only SST predictor variables.
- **Classifier C:** Random forest using only SOI and MSLP predictor variables.

The results for these new classifiers are displayed in Table 3. We see that the misclassification rate for the random forest constructed using only SST predictor variables is identical to that from the original random forest, while the random forest grown with just MSLP and SOI variables suffers only slightly in comparison. We discuss this matter further in the final section of this article.

c. Before and after 1975

On close inspection, the results seemed to indicate that the misclassification rates got worse after the mid-1970s. This is illustrated by Fig. 3, which gives the number of misclassified towns by year and according to the classification problem, as described in section 3b. A formal statistical comparison of the overall error rates before and after 1975 does not indicate a significant drop in performance for the classifier

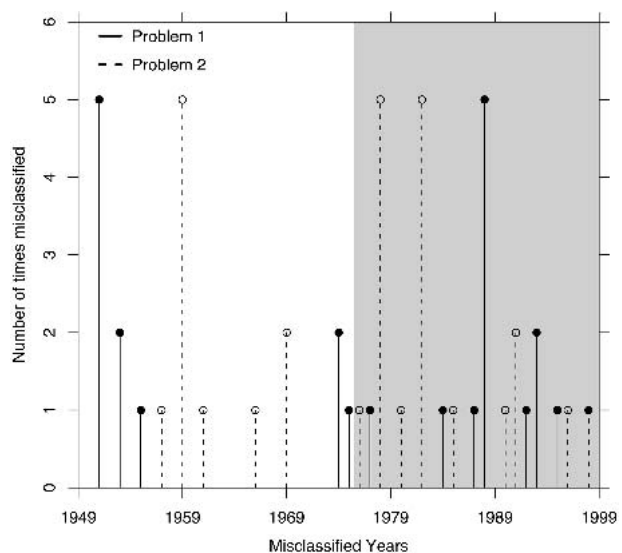


FIG. 3. The misclassification rates for Classifier A for problems 1 and 2. The white background is for years up to 1975 (inclusive), the gray background is for years after 1975. The dotted line with solid black dot on top at 1998 indicates that this year was misclassified once for both problems.

TABLE 4. The misclassification rates for the 1949–75 and 1976–99 forests using all predictor variables.

Year range	MR (%)
1949–75	15.56
1976–99	25.83

(p value = 0.119). However, there is a significant increase in the misclassification rate for problem 2 (p = 0.022), indicating that it becomes more difficult to predict mid versus late season onset after 1975. The timing of this change in the predictability of the onset of the wet season coincides with a decline in the southwestern rainfall in the mid-1970s, as documented by a number of researchers including Nicholls et al. (2000).

We further explored the changes that occurred in the mid-1970s by dividing the data into two subsets, one containing observations for years 1949–75 and the other containing observations for 1976–99. Random forest classifiers were grown from both subsets. The misclassification rates are given in Table 4, and show the expected worsening of performance of the classifier in the latter period. For both periods the SST explanatory

variables were most influential in determining the predictions, but a detailed analysis unearthed marked changes in the important SST grid cells between the two periods. These results are displayed pictorially in Figs. 4–7, which show important SST cells (by area and month) for prediction under various scenarios. Important cells are further differentiated using a scale of 1 (moderate) to 3 (high) corresponding to the number of importance measures that recognized the cell's influence.

Figures 4 and 5 indicate that the area of the Indian Ocean where the SST is important for problem 1 (predicting early/not early onset) has changed. Prior to 1975 the important region was in the Tropics. This now seems to have moved down toward the Southern Ocean. Figures 6 and 7 also show this change, although it is not as strong. Before 1975 the important region stretched over all latitudes from where SSTs were measured (10° – 40° S), but after 1975 the Tropics no longer appears to be important. Furthermore, these problem 2 (mid/late onset) forests also show a change in the months in which the SST is important. After 1975, March was the important month for the purposes of prediction, but before 1975 the important SSTs oc-

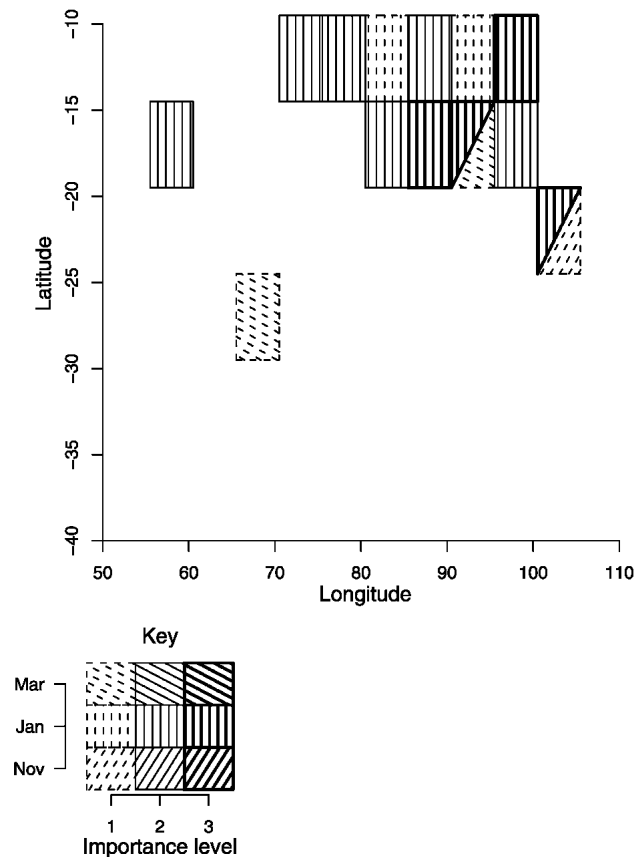


FIG. 4. Important SST cells for problem 1 using 1949–75 data. The heavier the shading for a cell, the greater its importance for rainfall prediction. The month for which SST at a cell was important is indicated by the key.

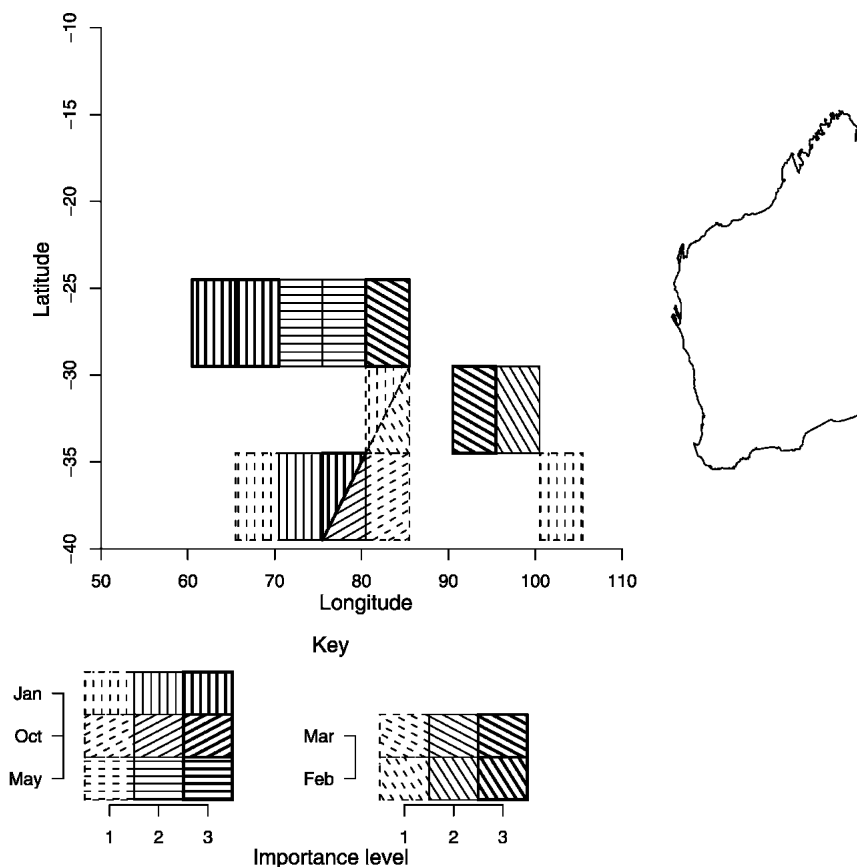


FIG. 5. As in Fig. 4 but using 1976–99 data.

curred three months earlier in December. This suggests that not only has the important region in the Indian Ocean been altered, but that the lead time for SST influences has also changed.

4. Discussion and conclusions

The aim of our case study was to develop a model to predict when the wet season in the southwest of Western Australia would begin using SST, SOI, and MSLP as potential explanatory variables. The model that we developed correctly predicted the onset approximately four out of five years. The SST variables proved the most influential in deriving the predictions. Indeed, the misclassification rate did not change when a random forest was constructed with SOI and MSLP data omitted. However, this finding should be interpreted with care for two reasons: First, there were many more SST variables than SOI or MSLP variables since SST data is recorded over 72-cell grid. Hence the amount of information per SST cell may be quite modest. Second, our statistical classification algorithm aims to detect associations rather than causal links. One might well expect the association between SST and the onset date for the

winter rains to be explained by other factors. See Nicholls et al. (2000) for related comments.

An interesting finding in our analysis was the change in the misclassification rate, and the important explanatory variables, that occurred in the mid-1970s. This change occurred at the same time as the rainfall decline observed in much of southwest Western Australia (Nicholls et al. 2000). These results are suggestive of some significant change in the climatic processes for this region at this time, reflected in a changed region of Indian Ocean SST influence on wet season onset. The region of influence appears now to be in the south Indian Ocean rather than the Tropics. It is interesting to note that the decline in rainfall over southwest Western Australia is predominantly in early winter rainfall. One hypothesis to explain this is that bands of moisture from the Tropics now appear to be less common so that there is less scope for interaction between these so-called northwest cloud bands and frontal systems from the south [the Indian Ocean Climate Initiative (IOCI) 2002]. Such interactions produce substantial rainfall when they occur. Our case study suggests that there may be a link with Indian Ocean SSTs. Dynamic modeling of the climatic system might begin to provide some insight into the underlying physical causes.

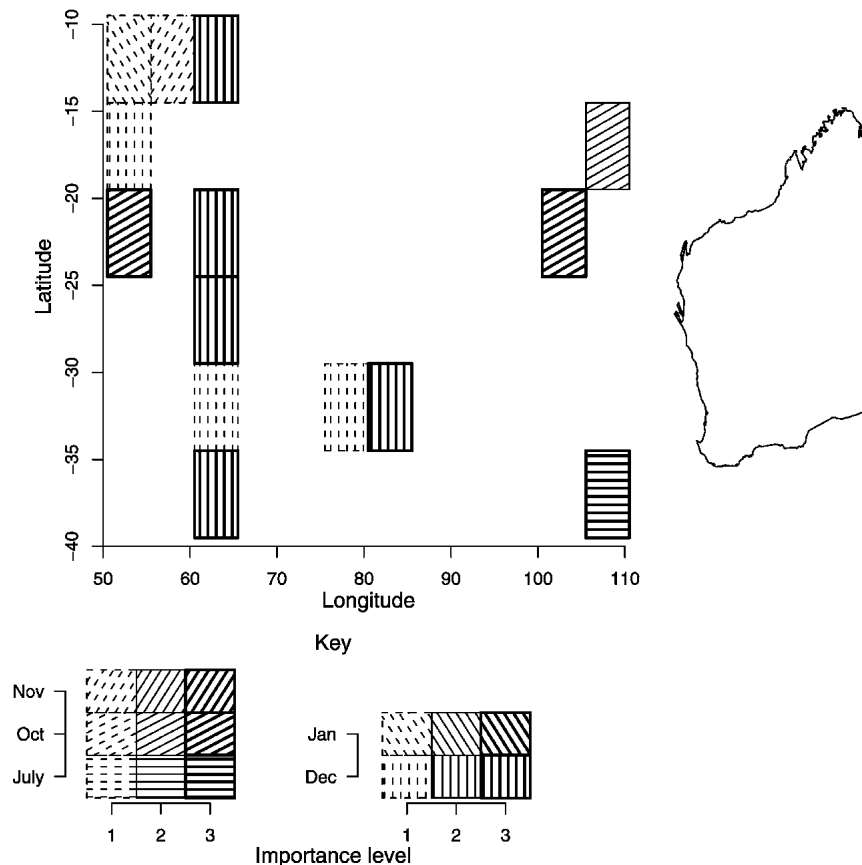


FIG. 6. Important SST cells for problem 2 using 1949–75 data. The heavier the shading for a cell, the greater its importance for rainfall prediction. The month for which SST at a cell was important is indicated by the key.

We note that this analysis could be extended to include the effects of the Southern Ocean. Recent research has found the existence of the Antarctic “Circumpolar Wave,” a current that moves continuously around Antarctica, completing the circuit approximately every eight to nine years (White and Peterson 1996). This current advects alternating regions of warm and cold water, which seem to interact with atmospheric pressure and wind patterns. There is some evidence linking the warm regions with warm, wet winters and the cold regions with cool, dry winters in the south of Australia. The effects of the Indian Ocean, Pacific Ocean, and Southern Ocean anomalies are most likely dependent. Thus, the Antarctic circumpolar wave could be an important factor in the timing of the wet season onset.

We observe that classification trees do not assume normally distributed observations or that the responses are related to the predictors in a linear fashion. Given the complex nature of the climate system it is a clear advantage to avoid simplistic assumptions such as these. For example, the Australian Bureau of Meteorology’s seasonal forecast system uses principle compo-

nents of sea surface temperature anomalies over the Indian and Pacific Oceans (Drosowsky and Chambers 2001). These are then used in a linear discriminant analysis to produce probabilities that seasonal rainfall will be above or below the long-run median; in some applications terciles are forecast (“below normal,” “near normal,” and “above normal”). The discriminant analysis model assumes that each observation belongs to one of a set of mutually exclusive groups, which is a form of nonlinear relationship. In our implementation of the CART methodology this is also true, but the contribution of predictors to the modeling is more general as no particular functional form for the relationship between explanatory variables and prediction is assumed. The random forests approach is also able to handle a large number of predictors, so there is no need to first reduce the number of predictors via principle component analysis, for example.

Our case study suggests that random forests is a potentially useful tool in climate research. It is capable of identifying important predictors, and physical insights, and of generating forecasts that appear to have significant skill.

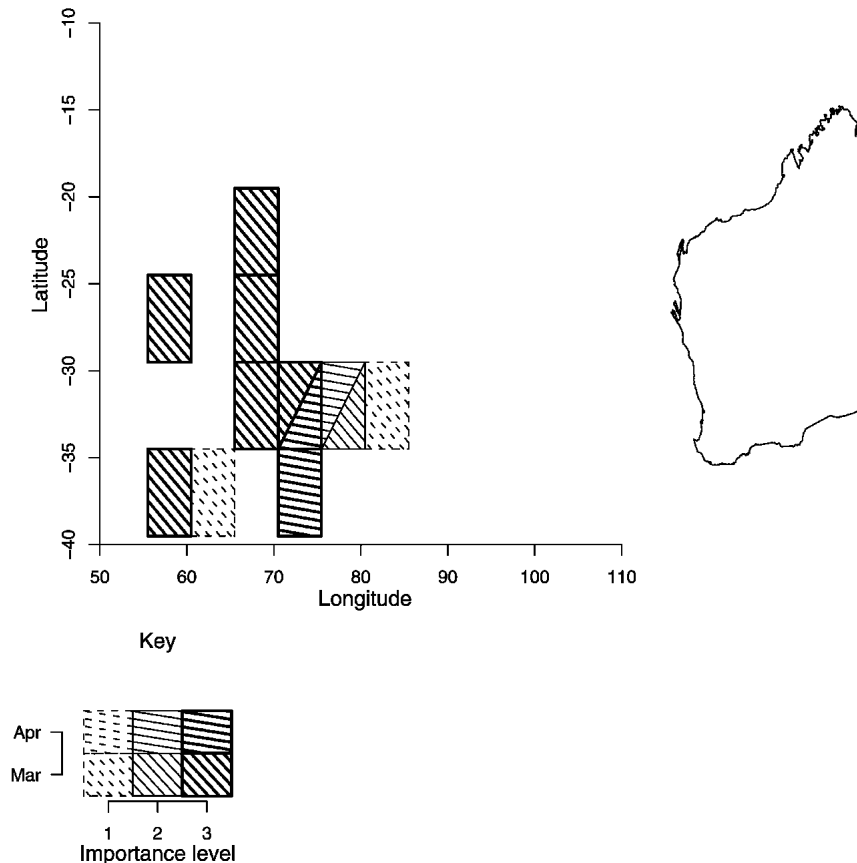


FIG. 7. As in Fig. 6 but using 1976–99 data.

Acknowledgments. The first author was funded by a top-up scholarship from CSIRO Mathematical and Information Sciences. The authors are grateful to the Department of Agriculture Western Australia for access to soil moisture information, and particularly to Ian Foster for his time and advice. This work was sponsored by the Indian Ocean Climate Initiative of the Western Australian State Government, which provided access to air pressure (courtesy of the Commonwealth Bureau of Meteorology) and sea surface temperature data. We are grateful to Neville Nicholls and an anonymous reviewer for their thoughtful comments that greatly improved the final paper.

REFERENCES

- Breiman, L., 1996a: Heuristics of instability and stabilization in model selection. *Ann. Stat.*, **24**, 2350–2383.
- , cited 1996b: Out-of-bag estimation. [Available online at [ftp://ftp.stat.berkeley.edu/pub/users/breiman/OOBestimation.ps.Z](http://ftp.stat.berkeley.edu/pub/users/breiman/OOBestimation.ps.Z).]
- , 1998: Arcing classifiers. *Ann. Stat.*, **26**, 801–849.
- , 2001a: Random forests. *Machine Learn.*, **45**, 5–32.
- , 2001b: Statistical modeling: The two cultures. *Stat. Sci.*, **16**, 199–231.
- , and A. Cutler, cited 2001: Manual on setting up, using, and understanding random forests V3.1. [Available online at http://oz.berkeley.edu/users/breiman/Using_random_forests_V3.1.pdf.]
- , J. H. Friedman, R. A. Olshen, and C. J. Stone, 1984: *Classification and Regression Trees*. Wadsworth Advanced Books and Software, 368 pp.
- Campbell, E., B. Bates, and S. Charles, 2000: Nonlinear statistical methods for climate forecasting. *Towards Understanding Climate Variability in South Western Australia*, Indian Ocean Climate Initiative, 179–237. [Available online at http://www.ioci.org.au/publications/pdf/IOCI_SPR_5.pdf.]
- Dietterich, T. G., and E. B. Kong, 1995: Machine learning bias, statistical bias, and statistical variance of decision tree algorithms. Department of Computer Science Tech. Report, Oregon State University, 14 pp. [Available online at <http://web.engr.oregonstate.edu/~tgd/publications/tr-bias.ps.gz>.]
- Drosowsky, W., and L. E. Chambers, 2001: Near-global sea surface temperature anomalies as predictors of Australian seasonal rainfall. *J. Climate*, **14**, 1677–1687.
- Efron, B., and R. J. Tibshirani, 1993: *An Introduction to the Bootstrap*. Chapman and Hall, 436 pp.
- Gower, J., 1998: Classification overview. *Encyclopedia of Biostatistics*, P. Armitage and T. Colton, Eds., Vol. 1, John Wiley & Sons, 656–667.
- Graf, H.-F., and J. M. Castanheira, 2001: Structural changes of climate variability. Max-Planck-Institut für Meteorologie Tech. Rep. 330, Hamburg, Germany, 12 pp.
- Hastie, T., R. Tibshirani, and J. Friedman, 2001: *The Elements of Statistical Learning: Data Mining, Inference and Prediction*. Springer, 522 pp.
- Hunt, B., I. Smith, R. Allan, and T. Elliot, 2000: Multi-seasonal

- predictability and climatic trends for south-west Western Australia. *Towards Understanding Climate Variability in South Western Australia*, Indian Ocean Climate Initiative, 53–124. [Available online at http://www.ioci.org.au/publications/pdf/IOCI_FPR_3.pdf.]
- Ihaka, R., and R. Gentleman, 1996: R: A language for data analysis and graphics. *J. Comput. Graph. Stat.*, **5**, 299–314.
- IOCI, cited 2002: Climate variability and change in south west Western Australia. [Available online at http://www.ioci.org.au/publications/pdf/IOCI_TechnicalReport02.pdf.]
- McBride, J., and N. Nicholls, 1983: Seasonal relationships between Australian rainfall and the Southern Oscillation. *Mon. Wea. Rev.*, **111**, 1998–2004.
- Nicholls, N., 1984: A system for predicting the onset of the North Australian wet-season. *J. Climatol.*, **4**, 425–435.
- , 1985: Impact of the Southern Oscillation on Australian crops. *J. Climatol.*, **5**, 553–560.
- , 1989: Sea surface temperature and Australian winter rainfall. *J. Climate*, **2**, 965–973.
- , 1991: The El-Niño/Southern Oscillation and Australian vegetation. *Vegetatio: Vegetation and Climate Interactions in Semi-Arid Regions*, A. Henderson-Sellers and A. Pitman, Eds., Vol. 91, Kluwer Academic, 23–36.
- , L. Chambers, and W. Drosowsky, 2000: Climate variability and predictability for south-west Western Australia. *Towards Understanding Climate Variability in South Western Australia*, Indian Ocean Climate Initiative, 1–52. [Available online at http://www.ioci.org.au/publications/pdf/IOCI_SPF_2.pdf.]
- Palmer, T. N., 1999: A nonlinear dynamical perspective on climate prediction. *J. Climate*, **12**, 575–591.
- Robertson, G., cited 2001: Director General's overview: Department of Agriculture annual report 2000–2001. [Available online at http://agspsrv34.agric.wa.gov.au/agency/pubns/annualreport/2000_2001/overview.htm.]
- White, W. B., and R. G. Peterson, 1996: An Antarctic circumpolar wave in surface pressure, wind, temperature and sea-ice extent. *Nature*, **380**, 699–702.