

Background Spectrum Classification for Cognitive Radio

Shridatt Sugrim*, Melike Baykal-Gursoy[†], and Predrag Spasojevic[1][‡] *Member, IEEE*

* WINLAB, Rutgers University, 671 Rt. 1 South, North Brunswick, NJ 08902, USA
email: ssugrim@winlab.rutgers.edu

[†] Department of Industrial and Systems Engineering, CAIT, Rutgers University
96 Frelinghuysen Rd, Piscataway, NJ 08854-8018, USA
email: gursoy@rci.rutgers.edu

[‡]Department of Electrical and Computer Engineering, Rutgers University, Piscataway, NJ, USA
email: pspasojevic@rci.rutgers.edu

I. ABSTRACT

Recent changes in policy regarding the opportunistic use of licensed radio spectrum have paved the way for new innovative technologies like cognitive radio (CR). This technology puts tight demands on systems built to sense spectrum occupancy. Any strategy employed for opportunistic spectrum usage has to consider the tradeoffs between time spent searching for empty channels and time spent using those empty channels. In most cases the spectrum sensing that is employed by a CR system starts with no prior information about the occupancy of the channels it intends to use. A classifier can be run before the CRs attempt transmission to provide the CRs' spectrum sensing sub-systems with a set of occupancy probability categories for some of the channels. By providing priors it may be possible for the CR to reach a transmission strategy in a shorter time frame.

We propose a novel method of addressing this lack of prior knowledge by employing an efficient strategy that classifies some of the channels the CR intends to use within a fixed time limit. Our classification algorithm is based on multiple sequential probability ratio tests (multi-SPRT) and a heuristic allocation strategy for measurements that considers the completion time of each multi-SPRT. We will show that this strategy will achieve a bounded error by prioritizing channels that give consistent measurement results. We also compare the performance of the proposed system to simpler systems that do not require as many computations.

II. INTRODUCTION

Radio spectrum is a limited resource that is currently under heavy contention. Regulatory bodies like the FCC take on the daunting task of fairly distributing this resource among the many spectrum hungry users. Recent changes in policy regarding the opportunistic use of licensed spectrum has paved the way for new innovative technologies like cognitive radio [1].

This adaptive radio design allows secondary users (SU) to use spectrum slotted for a different purpose if the primary user (PU) is not currently using it. A key element of this technology is sensing whether the primary user is present or not. This sensing process typically requires a very extensive sensing period because of the requirement that the interference that the PU sees be minimized (although there have been recent proposals to mitigate this requirement, see [2]). If the sensing takes too long, the utility of secondary usage of the spectrum goes down rapidly. When a sensing system takes too long to converge on a strategy, it may miss spectrum opportunities or fail to meet transmission service quality requirements. Since the radio is adaptive, an on-line sensing plan with short convergence times is required to optimally utilize spectrum. This area of research is very active, and there are several strategies proposed for sensing the environment. Techniques such as C-SPRT, and finite horizon dynamic programming are used to identify optimal stopping times and balance points between exploration and exploitation. Many of these strategies however assume a prior distribution that is either unknown or uninformed. All of these strategies could benefit from any amount of a priori information about the occupancy of the channels they intend to use.

Spectrum classification can help reduce convergence times and facilitate efficient spectrum usage. Since data usage patterns are coupled with the daily routine of users, classification based on geography and time will be of great utility. However practical approaches to spectrum classification are not very well explored. Classification can possibly be done by two different approaches, an infrastructure approach and a mobile approach.

In the infrastructure case measurement opportunities are infinite, however mobility is not possible. In the mobile case, we are not tied to fixed geographic locations, however because we are mobile we have a limit on the amount of measurement resources per geographic location. In this discussion, we focus on the mobile case. The mobile measurement case can arise from regulatory bodies doing characterization sweeps across geographic areas, to assess usage and possibly to coordinate SUs via a control channel.

¹The authors would like to thank CACI Technologies for funding in part this research.

The proposed technique can be used as part of a distributed or hybrid interference temperature measurements like those described in [3]. It could also be done by mobile SUs, in an attempt to build a personal geographic map of the occupancy for the SU. This map can be consulted before the initial spectrum scan that precedes a transmission. This will speed up the initial convergence and may identify recurrent spectral artifacts that need to be avoided. If the latter is done in some form of a handset, the measurement resources might be very limited because of battery conservation policies, or other resource constraints. Such maps will naturally be geographically limited, however techniques such as kriging spatial interpolation (as described in [4]) can be used to extend the map's utility.

In the following work, we propose a novel strategy for distributing a limited number of measurements amongst a set of channels with unknown occupancy. Our goal is to classify as many channels as we can with the limited amount of resources. We will achieve this goal by minimizing a cost function that depends on the expected number of measurements needed to satisfy an error bound. The classification will partition a set of channels into subsets with similar probability of occupancy. This classification can be used to adjust future sensing policies, perhaps by starting with channels that are in the set of low occupancy channels first.

The rest of the paper is laid out as follows. In Section III we survey some related approaches and describe the model we employ in Section IV. Simulation results and performance analysis is done in V. Finally we examine future work and conclusions in VI.

III. RELATED WORKS

The broad category of spectrum sensing for CRs is actually quite well explored. In [5], Yucek charts various strategies for sensing spectrum ranging from simple energy detection to multi-dimensional techniques. The key distinction between our work and many of the methods suggested is that their primary goal is to maximize throughput by efficiently finding spectral holes. In almost every case their primary metric is maximum throughput under different constraints, e.g. fair access to all users. In our proposed method, on the other hand, we will explore different performance metrics for an entirely different class of problems, that of classification of channels.

In [6], Lai et al. propose an adaptive allocation model for sensing that employs solutions to multi-armed bandit problems. In the single user case, they use a network similar to the proposed model of our work and consider a time slotted collection of channels. At each time slot, the CR must first listen to the channel. If the channel is empty, the CR will transmit. If not the CR will update the channel's occupancy model to avoid picking this channel again. At each step the authors are computing the expected future gains based on the learned history. Each choice is then made to maximize the expected throughput.

Our approach has a different utility function that aims to limit the amount of time spent sensing any individual channel. In the bandit formulation if a channel is unoccupied throughout

the duration of the sensing interval, this channel is always the best choice. However in our problem, such a channel is only of limited utility. For an unoccupied channel, we require only a small number of samples to confirm that it is indeed unoccupied. Once we have determined that it is empty with a sufficiently small probability of error, this channel is decided and thus sampling it more is of no utility.

Many techniques fall back on the classical Neyman-Pearson methods of error analysis, however some (e.g. Xin and Lai in [7]), employ sequential testing techniques. Here the authors focus on a similar class of problem that minimizes a cost function that depends on the expected sequence length and the error probability. Their model comes from noisy power observations where the signal is required to be present during the entire observation period. We do not assume a signal model that requires a constant signal, instead we expect that an underlying detector can hand us an instantaneous binary decision as to whether the channel is occupied or not. Our channels may be partially occupied with occupancy probability p_n .

IV. PROPOSED MODEL

Our model assumes a set of N independent channels, indexed by $n \in \{1, \dots, N\}$. Channel n is occupied with probability p_n . We assume that each p_n is drawn from a uniform distribution on $[0, 1]$. If we consider the set $P_N = \{p_n\}$, the proposed system will classify a subset of N channels by distributing a limited number of measurements L among the N channels. Each measurement of a channel $x_n(\cdot)$ is considered to be a Bernoulli random variable $B(p_n)$ with 1 indicating occupied, and 0 indicating free. The argument of $x_n(\cdot)$ is the current sample index for the respective channel.

Time will be normalized by samples and indexed by the global time index $j \in \{1, 2, \dots, L\}$. Each channel n has a sample index, $m_n(j) \in \{1, \dots\}$, the number of samples received at time j . Our approach will selectively distribute samples among the channels by building a selection function $\hat{n}(j)$, which will be defined in Section IV-B. $m_n(j)$ is then given as:

$$m_n(j+1) = \begin{cases} m_n(j) + 1 & : \text{if } n = \hat{n}(j), \\ m_n(j) & : \text{otherwise.} \end{cases} \quad (1)$$

Due to the constrained nature of the problem, not all channels will complete a classification. We index the stopping time for channels which complete the classification as M_n and the largest sample index for channels which do not complete as M'_n . Thus we have $L = \sum_{\text{classified}} M_n + \sum_{\text{unclassified}} M'_n$.

A driving observation to this method is that as the parameter of the Bernoulli random variable p_n gets closer to the edges of the $[0, 1]$ parameter space, the variance of the random variable decreases, see Figure 1. This drop in variance should be coupled with an increase in the consistency of samples. Channels that have more consistent measurements will take fewer samples to characterize while maintaining a bound on error. The proposed system will prioritize these channels over channels that are not consistent. This preference for consistency will be reflected in our performance metrics.

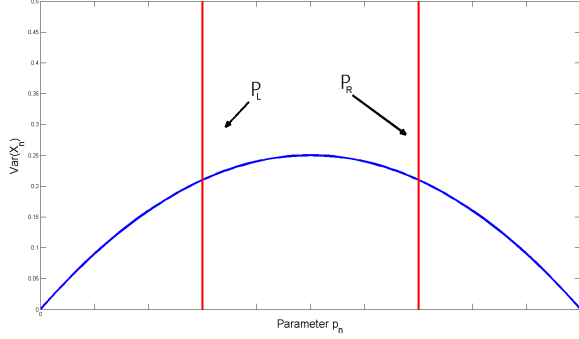


Fig. 1. Variance of $x_n(\cdot)$ as a function of p_n . P_L and P_R are extrinsic design choices. The position of the boundaries trades the efficacy of the search and the precision of the answer.

A. Single Channel Decision

For each n , we formulate our hypothesis to reflect our preference for consistent measurements. In Figure 1 we have partitioned the parameter space into 3 sets, $[0, P_L]$, (P_L, P_R) , $[P_R, 1]$. Our hypothesis can be stated as:

$$\begin{aligned} H_{C,n} : p_n &\in [0, P_L] \text{ or } p_n \in [P_R, 1] \\ H_{I,n} : p_n &\in (P_L, P_R) \end{aligned} \quad (2)$$

$H_{C,n}$ is a claim that channel n is consistent. In choosing P_L, P_R we must balance the performance of the system, as discussed in Section V-B, against precision of the answer returned. The width of the closed subsets directly effects the efficiency of the search. P_L, P_R have no n dependence as they are global parameters.

For this set of hypotheses we need a test that will satisfy a bound on $\beta = P\{\text{Declare } H_{C,n} | H_{I,n} \text{ is true}\}$ while making a decision with the fewest samples. In [8], it is shown that the SPRT satisfies these criteria.

Because the hypotheses given in (2) partition the parameter space into three distinct regions we will use multiple SPRTs and a combining rule to evaluate these hypotheses. If we consider only the “left” side of (2) our hypothesis simplifies to:

$$\begin{aligned} H_{L,n} : p_n &\leq P_L, \\ H_{L',n} : p_n &> P_L. \end{aligned} \quad (3)$$

The SPRT was derived to address exactly this form of composite hypotheses. To build the test, we identify an indifference region about P_L , a region in which we are indifferent to errors in the declaration of the test. For simplicity of the formulation we can choose a small δ about P_L and then make the assignments:

$$\begin{aligned} \theta_1 &= P_L + \delta, \\ \theta_0 &= P_L - \delta. \end{aligned} \quad (4)$$

In [9] it is shown that the test that decides simple hypothesis

of the form:

$$\begin{aligned} H_0 : p_n &= \theta_0, \\ H_1 : p_n &= \theta_1, \end{aligned} \quad (5)$$

can be used to decide (3). The likelihood ratio can therefore be written as:

$$\Lambda(x_n(m_n(j))) = \frac{f(x_n(1), \theta_1) \dots f(x_n(m_n(j)), \theta_1)}{f(x_n(1), \theta_0) \dots f(x_n(m_n(j)), \theta_0)}, \quad (6)$$

where $\Lambda(x_n(m_n(j)))$ is a function of the sample index, and $f(x_n(m_n(j)), \theta_1)$ is the probability of the $m_n(j)^{th}$ sample takes on the value $x_n(m_n(j))$ when the parameter equals θ_1 . It is compared to two thresholds A and B , which are chosen to enforce a bound on β . The test ends when one of the thresholds is crossed. The decision rule for this SPRT is then given as:

$$\varphi_L(x_n(M_n)) = \begin{cases} 1 & : \text{Declare } H_{L,n} \Leftrightarrow \Lambda(x_n(M_n)) < A, \\ 0 & : \text{Declare } H_{L',n} \Leftrightarrow B < \Lambda(x_n(M_n)), \end{cases} \quad (7)$$

Similarly we can derive a test for the “right” side of (2) which will decide between the hypotheses:

$$\begin{aligned} H_{R,n} : p_n &\geq P_R, \\ H_{R',n} : p_n &< P_R, \end{aligned} \quad (8)$$

and produce an analogous decision rule $\varphi_R(x_n(M_n))$. A full derivation is given in [10].

Utilizing these two tests we can build multi-SPRT system that decides between $H_{C,n}$ and $H_{I,n}$ by using a combining rule of the form:

$$\phi_C(x_n(M_n)) = \varphi_R(x_n(M_n)) + \varphi_L(x_n(M_n)). \quad (9)$$

$\phi_C(x_n(M_n))$ serves as an indicator for the declaration of $H_{C,n}$. The function does not require a **mod**₂ operation as the declaration of $H_{L,n}$ and $H_{R,n}$ are mutually exclusive. Since the sufficient statistic is given as

$$d_n(j) = \sum_{m_n(j)=1}^{M_n} x_n(m_n(j)), \quad (10)$$

we can use $\phi_C(d_n(j))$, a similar indicator based on the sufficient statistic, to compute β for this test. Because the test is sequential, we can define an operator characteristic curve (OC curve) as

$$L(p_n) = P(\text{declare } H_{C,n})_{p_n} = \int \phi_C(d_n(M_n)) dF_{p_n}(d_n(M_n)), \quad (11)$$

where $dF_{p_n}(d_n(M_n))$ is a measure due to the distribution on $d_n(M_n)$ with parameter p_n . $L(p_n)$ is the probability that sequential test will terminate with a decision of $H_{C,n}$ if the samples from channel n have p_n as their parameter. We can then compute β , the probability that this test will miss-classify a channel n with parameter $p_n \in (P_L, P_R)$ as:

$$\beta = \int_{P_L}^{P_R} L(p_n) f(p_n) dp_n = \frac{1}{P_R - P_L} \int_{P_L}^{P_R} L(p) dp, \quad (12)$$

since p_n was drawing from a uniform distribution.

To implement this test one can appeal to (6), however in [11] it is shown that we can evaluate the test by considering the sufficient statistic $d_n(j)$ as a random walk between two pairs of threshold lines which are derived from (6), (7), and their “right” side analogues. In Figure 2, we can see a sample run of the graphical interpretation. This graph also shows the regions where $\phi_C(x_n(M_n))$ will make it’s final decision.

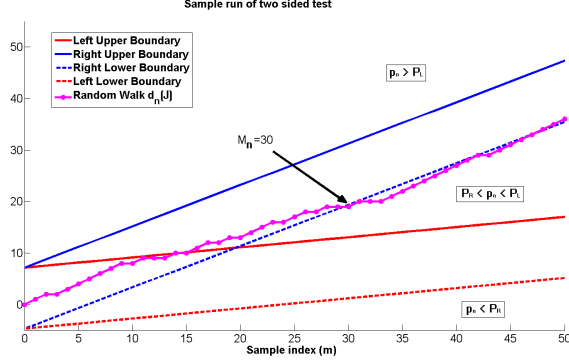


Fig. 2. Sample run of the multi-SPRT. The Left boundary lines decide if the parameter is above or below P_L . The right boundary lines do the same for P_R . When the random walk $d_n(j)$ crosses into one of the labeled regions, the test makes that declaration.

B. Multi-Channel Measurement Allocation Strategy

The multi-SPRT of Figure 2 gives us a metric for assessing the proximity to completion, the distance to the nearest decision making threshold line. We use this metric and a greedy allocation strategy to decide which channel to allocate future measurements to. We formulate a cost function for each channel as:

$$C_n(m_n(j)) = \begin{cases} e_{n,m_n(j)} + (\lambda * m_n(j)) & : \text{unclassified,} \\ \infty & : \text{classified,} \end{cases} \quad (13)$$

where $e_{n,m_n(j)}$ is a distance metric which is measured as minimum samples required to reach a decision from the current position, see Figure 3. The $\lambda * m_n(j)$ is a term that penalizes oscillations in the $e_{n,m_n(j)}$ term. Once a channel has been classified there is no longer any utility in allocating additional measurements to it, thus we set the cost to infinity to ensure no further allocations will be made.

To minimize the time spent on any one channel we take a greedy approach of choosing channels with minimum cost at each sample allocation. These channels are the closest to a decision while having taken the fewest measurements to reach this decision. Our selection function is then given by:

$$\hat{n}(j) = \arg \min_n (C_n(m_n(j))). \quad (14)$$

In many cases this minimum will not be unique, when this occurs we choose an index randomly according to a uniform distribution on the set of minimum cost channels.

At $j = 0$, all channels have the same initial cost which is the number of samples required to cross the nearest threshold

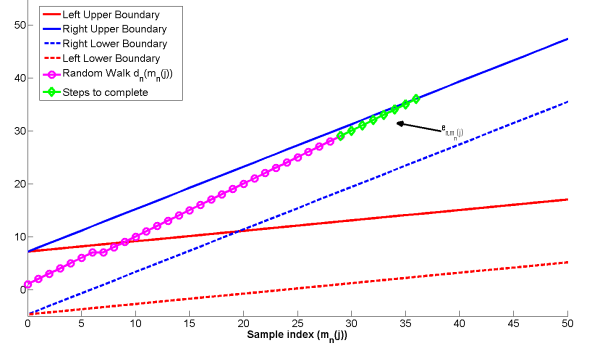


Fig. 3. $e_{n,m_n(j)}$ seen graphically. At the current step $m_n(j) = 30$. If the random walk takes 7 more steps upward, it will make a decision. Since this is the smallest number of steps required to complete the test $e_{n,m_n(j)} = 7$.

line without any changes in direction, i.e., the shortest path to a decision. Once a channel has received a few samples its random walk may take a straight path towards a threshold line, in which case it will reduce the cost by lowering $e_{n,m}$. The greedy choice will continue to pick this channel until it reaches a decision or starts raising the cost by changing direction. The λ is chosen so that when a channel’s random walk changes direction, it forces the cost to go up, and thus makes the greedy choice pick a different channel in the next allocation. Under this strategy channels that reach decisions in the fewest samples are selected over channels that require long sequences to characterize. Channels for which the samples are consistent will have fewer direction changes in their random walks, and thus be prioritized.

V. SIMULATION RESULTS

The determining factor in the utility of our approach is the resource regime we have to work with. If the number of samples is large compared to the number of channels, then there is no need to be frugal with measurements. The benefits of running this allocation scheme are grossly out weighed by it’s computational overhead. However, in a scenario when the number of channels is greater than the number of samples we have to allocate, this system will return meaningful answers while satisfying a bound on the error.

A. Alternative Approaches

For comparison we examine two alternative allocation schemes. The first scheme is the most simple, which is to divide the measurements we have equally among all the channels we have to search through, and to compute the classical estimate of the parameter p_n . Each channel will get $\Gamma = \frac{L}{N}$ samples. The estimate becomes $\hat{p}_n = \sum_{m_n(j)=1}^{\Gamma} x_n(m_n(j)) / \Gamma$. We place said channel in their respective category based on the comparison of $\hat{p}_n$ to P_L, P_R . While this approach is very easy to implement and has minimal computational overhead, it does not adapt to different resource regimes.

The second approach we consider is more adaptive. We begin by allocating some fraction of our samples for an initial pass. In this initial pass we will sample each channel a fixed number of times, Ψ , up to our budget of samples. For each channel that has received at least Ψ samples, we identify channel that have sufficient agreement between samples, e.g. if $\Psi = 4$, we would allow for at most 1 sample that is different from the others. Channels that satisfy this policy will be marked as candidates for the second pass, and in this second pass will again receive Ψ samples. This process repeats until we run out of samples. When we have run out of samples (or satisfied some sampling depth limit) the channels that received samples will have their parameter estimates computed as $\hat{p}_n = \sum_{m_n(j)=1}^{M'_n} x_n(m_n(j))/M'_n$, and these estimates will be used to categorize the channels as before, see algorithm 1.

This approach iteratively builds a tree of allocated samples. Channels with a high degree of consistency between samples have a low variance, and thus a shorter expected sequence length. By allocating more samples to consistent channels the error in the estimator $\hat{p}_n$ is lowered. This approach can be seen as an approximation to the proposed method, with a lower computation overhead.

```

#First Pass;
reserve fixed fraction of samples;
while reserve not empty do
    | allocate  $\Psi$  samples to channel;
end

#All Subsequent Passes;
while still have samples do
    #identify consistent subset;
    forall the channels that received samples do
        if consistent then
            | keep
        else
            | discard
        end
    end
    foreach consistent channel do
        | allocate  $\Psi$  samples to channel;
    end
end

#Classify;
forall the channels that received samples do
    compute  $\hat{p}_n$ ;
    switch  $\hat{p}_n$  do
        case  $< P_L$  declare  $H_{C,n}$ ;
        case  $> P_R$  declare  $H_{C,n}$ ;
        otherwise
            | declare  $H_{I,n}$ 
        end
    endsw
end

```

Algorithm 1: Tree Scheme for measurement allocation

B. Performance

To compare performance we choose a fixed $L = 1000$ and then consider an increasing N . As the resource regime gets progressively worse, the utility of our approach will become apparent. We will consider two performance metrics that demonstrate the trade off we make when employing this approach. The first is the number of consistent channels discovered as an increasing function N , that is channels declared $H_{C,n}$. This metric is simple to quantify, but somewhat misleading when used to measure performance. To compute the metric, we first define an indicator function on the results of a system that identifies which hypothesis was declared for each channel:

$$I_C(n) = \begin{cases} 1 & : \text{if } H_{C,n} \text{ was declared,} \\ 0 & : \text{otherwise.} \end{cases} \quad (15)$$

For the case of multi-SPRT (15) is the same as (9). We can then compute the metric as $\sum_n I_C(n)$. As a first look consider Figure 4. In Figure 4 we can see for $\Gamma < 5$ the simple scheme starts to break down. The simple scheme does not have enough samples to make meaningful estimates for each channel, and is thus is declaring more channels to be consistent than there actually are. Because the results of the simple scheme skew the scale of the graph, consider Figure 5 where we only compare the greedy-SPRT and tree building approaches. Figure 5 shows that the tree and greedy-SPRT approaches are very close in terms of discovery.

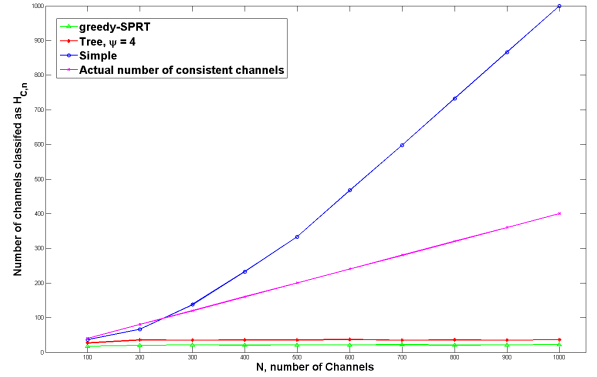


Fig. 4. Number of channels found for all 3 schemes as functions of N . As N increases, Γ goes from $10 \rightarrow 1$. The number of consistent channels is defined by P_L, P_R . In this simulation it's $\approx .40N$.

For the second metric we like to consider the probability that the system misclassifies a channel as consistent. Since we claim that channels with low probability of occupancy (consistently empty channels) are fit for opportunistic use, it is very important that we do not mislabel partially occupied channels as empty. We also deem channels that are consistently full as unusable, however if they are mislabeled, we have missed a spectral opportunity.

From Section IV-A we noted that test for each channel n misclassifies with probability β . The proposed system will

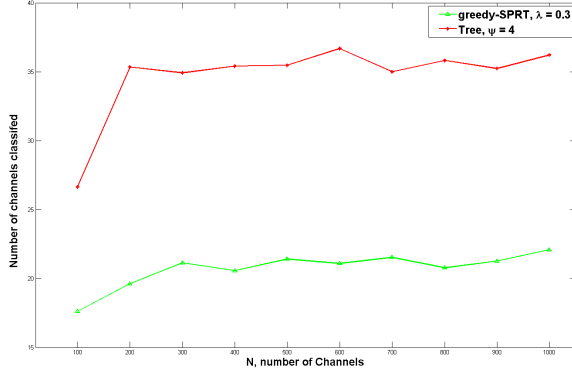


Fig. 5. Number of channels found for greedy-SPRT vs tree schemes as functions of N . The tree scheme uses a $\Psi = 4$. The greedy-SPRT scheme uses a lambda of 0.3.

combine several of these decisions into its final result. Since each channel is evaluated independently of every other channel, the probability of a missed classification is geometric with parameter β . For comparison we can directly compute the resulting probability of misclassification from the system outputs. By comparing the actual channel parameters p_n to their classifications we get the function:

$$I_M(n) = \begin{cases} 1 & : \text{if } H_{C,n} \text{ was declared but } p_n \in (P_L, P_R), \\ 0 & : \text{otherwise.} \end{cases} \quad (16)$$

Then the system misclassification rate is computed as

$$\beta_{\text{system}} = \frac{\sum_n I_M(n)}{\sum_n I_C(n)}. \quad (17)$$

We examine the β_{system} as a function of N in Figure 6.

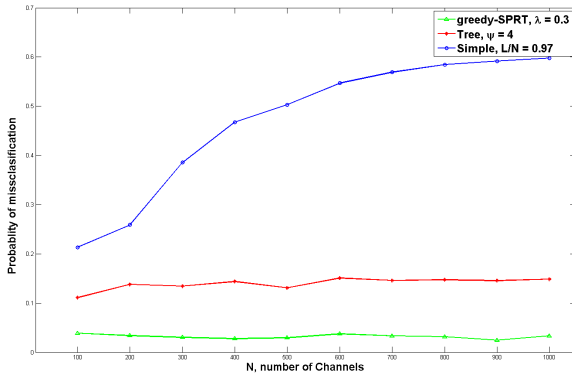


Fig. 6. β_{system} for all 3 schemes as functions of N . The error of the simple scheme approaches 0.6 as N grows large. When the system has so few samples per channel, it labels everything as consistent.

In this graph we see what goes wrong with the simple scheme, it makes too many errors. As the ratio $\frac{L}{N} \rightarrow 1$, the

greedy-SPRT approach maintains a bound on the classification error. Both the simple and tree approaches see an increase in misclassifications as N increases. In the simple case this error grows very large.

VI. CONCLUSION

Any adaptive system can benefit from having a reliable set of prior probabilities to start from. The proposed adaptive allocation approach was built to maintain a bound on error in the low resource regime where the distribution of samples matters. The approach was designed with the admission that it may run out of samples before it classifies every channel, however for any channel that is classified, a bound on error will be satisfied. In the cases where the available resources are very high, the simple scheme still made a significant number of classification errors. In contrast the greedy-SPRT kept the probability of misclassification well below 0.1 for all N .

In our model we assumed that all channels were independent as this represents a worst case scenario for the channel model. If we had additional information about correlations between channels, the system could be modified to incorporate such information (e.g. it could be used to make a more informed choice when we have to pick a random index).

The channel occupancy at any given time is a function of the usage patterns of the PUs (and SUs). In the uninformed case each user is equally likely to be streaming vs bursting packets. If at least some of the traffic is streamed, then some of the discovered p_n should be long lived relative to the duration of measurement interval.

While the system seems to perform well on paper and in simulation, its key drawbacks are its reliance on a uniformly distributed set of p_n and the requirement that the greedy choice be computed for every sample. In addressing the former we have started examining the worst case performance under non-uniform distributions. In hostile environments (environments where every channel is nearly occupied) the system identifies the “worst of the worst”, channels which should be avoided because they have no open time slots.

REFERENCES

- [1] A. J. Goldsmith and L. J. Greenstein, *Principles of Cognitive Radio*. Cambridge University Press, 2012.
- [2] S. W. Boyd, J. M. Frye, M. B. Pursley, and T. C. Royster, “Spectrum monitoring during reception in dynamic spectrum access cognitive radio networks,” *Communications, IEEE Transactions on*, vol. 60, no. 2, pp. 547–558, 2012.
- [3] P. J. Kolodzy, “Interference temperature: a metric for dynamic spectrum utilization,” *International Journal of Network Management*, vol. 16, no. 2, pp. 103–113, 2006.
- [4] R. C. Dwarakanath, J. D. Naranjo, and A. Ravanshid, “Modeling of interference maps for licensed shared access in lte-advanced networks supporting carrier aggregation,” in *Wireless Days (WD), 2013 IFIP*. IEEE, 2013, pp. 1–6.
- [5] T. Yucek and H. Arslan, “A survey of spectrum sensing algorithms for cognitive radio applications,” *Communications Surveys & Tutorials, IEEE*, vol. 11, no. 1, pp. 116–130, 2009.
- [6] L. Lai, H. El Gamal, H. Jiang, and H. V. Poor, “Cognitive medium access: Exploration, exploitation, and competition,” *Mobile Computing, IEEE Transactions on*, vol. 10, no. 2, pp. 239–253, 2011.

- [7] Y. Xin and L. Lai, "Fast wideband spectrum scanning for multi-channel cognitive radio systems," in Information Sciences and Systems (CISS), 2010 44th Annual Conference on. IEEE, 2010, pp. 1–6.
- [8] A. Wald and J. Wolfowitz, "Optimum character of the sequential probability ratio test," The Annals of Mathematical Statistics, pp. 326–339, 1948.
- [9] A. Wald, Sequential analysis. Courier Corporation, 1973.
- [10] S. Sugrim, "A strategy for classifying a set of dissimilar channels by their a priori channel occupancy probability," Ph.D. dissertation, Rutgers University-Graduate School-New Brunswick, 2014.
- [11] M. Sobel and A. Wald, "A sequential decision procedure for choosing one of three hypotheses concerning the unknown mean of a normal distribution," The annals of mathematical statistics, pp. 502–522, 1949.