

Validation of the ROVER-F Code for ROP Trip Probability Calculations

John Pitre and Frank Laratta

Reactor Core Physics Branch
Atomic Energy of Canada Limited
Sheridan Science and Technology Park
Mississauga, Ontario L5K 1B2

Introduction

An important task in the operation of CANDU[®] reactors is the prevention of fuel damage as a result of fuel dryout that can occur when the fuel sheath temperature exceeds the temperature at which the coolant can efficiently remove heat. The power at which fuel dryout is expected to occur is called the critical channel power and is a function of flux shape and the fuel channel thermalhydraulics. In CANDU reactors, protection against overpowers large enough to cause dryout is provided by two regional overpower protection (ROP) systems of in-core flux detectors, arrayed through the core, each organized into three safety (or logic) channels. Each of the two independent ROP systems is associated with one of the two independent shutdown systems (SDS-1 and SDS-2). The detectors in one ROP system (associated with SDS-1) are placed in vertical penetrations, whereas the other system (associated with SDS-2) uses detectors in horizontal penetrations in the core. Each ROP system is capable of initiating the shutdown of the reactor by actuating the corresponding shutdown system. Each ROP system must be so designed that in each safety channel at least one detector will reach its setpoint before there is damaging overpower in any fuel channel. The trip of a single detector in a safety channel will trip that channel, and the trip of two of the three safety channels in an ROP system will trip that ROP system.

The “trip probability” is the probability that each of the ROP systems will generate a signal to actuate a shutdown system before any fuel channel reaches dryout. The trip setpoints for ROP systems are set by a licensing requirement of a 98% probability of tripping for each of a “design basis” set of flux shapes, before any fuel channel reaches dryout. This is a function of the layout of the ROP system, the flux shapes and the uncertainties and biases associated with the channel powers, the critical channel powers, and the detector responses. The ratio of the critical channel power to the actual channel power is called the critical power ratio (CPR). The CPR is related to the total reactor power, decreasing as the total power increases. A ripple conservatism factor quantifies the effect of the local ripple relative to the allowance made to the detector calibration based on the maximum ripple in the high-power region of the entire core.

The uncertainties are divided into three groups. Detector-random errors are random errors that vary from detector to detector (e.g. recalibration errors). Channel-random errors are random errors that vary from fuel channel to fuel channel (e.g. uncertainty in channel power). Common-random errors are random in expected value but are common to all fuel channels or detectors (e.g. uncertainty in the total reactor power). The channel-random and common-random errors are combined to form a common-mode error.

The flux shapes used for the analysis of the ROP system are based on two components:

- flux shapes consisting of the nominal time-average flux shape and perturbations thereon (reactivity device positions, xenon transients, etc.); and
- instantaneous flux distributions of the reactor, in the form of channel power refuelling ripples.

ROVER-F is a FORTRAN program which calculates the trip probability and the setpoint adjustment required to attain the target trip probability, for a given set of flux shapes. This calculation is performed with the assumption that the most effective safety channel is unavailable and therefore the remaining two safety channels must both trip. The calculation of trip probability itself is non-iterative, but once the trip probability of the specified system has been calculated, a convergence iteration using a binomial search is

used to determine the adjustment to the trip setpoints that is required to attain the required probability target. ROVER-F is the translation of modules of the previously existing ROVER/REFORM code from APL to FORTRAN. The impetus for this translation is software quality assurance, which is becoming more and more important to apply and demonstrate. Validated, properly documented and maintained codes are crucial for the licensing of all reactors. Ensuring that the main design and safety codes exist in portable version in standard programming languages (such as FORTRAN, as opposed to APL) will ensure continuity of our capability to apply the codes with confidence

Concurrent with the translation of ROVER/REFORM to FORTRAN, a number of additional capabilities were built into the code. One advantage is that the code user has the ability to define the integration increments used for calculations of probability distribution and convolution integrals. These integrals are calculated over ranges of common-mode error and rippled CPR (and thus, by inference, total reactor power). The use of smaller step sizes for these calculations has the potential to increase the accuracy of the results as compared to the theoretical values. The potential of this increased accuracy will be examined in the individual benchmark results.

To simplify the computation of trip probability, the rippled critical power ratios are binned, or grouped, based on their relation to the limiting critical power ratio. Thus the uncertainty distributions may be calculated for each bin, which may be representative of a large number of channels, greatly decreasing the computation effort. The position of the limiting channel is the first bin. Although ROVER-F permits the specification of the bin size, for these tests a rippled CPR bin interval of 0.5% has been used in all cases. A smaller bin size should improve results when the channel random error is small.

ROVER-F also supports fully-variable array dimensioning, permitting it to be used with any detector channelization scheme. ROVER-F is a stand-alone code. A useful feature is that all intermediate variables calculated internally may be accessed through optional printouts. Setpoint adjustments are calculated automatically, and the code can perform tasks directed by input, such as trip-probability calculations assuming single-detector failure, and trip probability calculations for individual channel power ripple maps (instantaneous trip probability). On an HP 715 computer, ROVER-F performs trip-probability calculations for over 900 flux shapes, including the calculation of required setpoint, in only minutes of CPU time.

To ensure that ROVER-F produces the theoretically correct values (when these are known) for trip-probability calculations, it was applied to the ROSE-ROVER benchmarks. These benchmarks are a series of tests designed to rigorously test ROP trip calculation codes. Each benchmark test is designed to test an aspect of the trip-probability calculation. For each of these tests, a comparison has been made between ROVER-F, ROVER/REFORM and the theoretical solution to the benchmark problem.

The benchmark tests were run for two calculation resolutions. The first series of tests used an integration increment of 1%, the fixed value as that used by ROVER/REFORM. It was found that for the same calculation resolution ROVER-F gives results similar to those of ROVER/REFORM, with some improvement achieved in cases with low trip probability. In the second series of tests, the integration interval in ROVER-F was decreased to 0.04%. In these tests, there is an improvement in results and the trip probabilities calculated by ROVER-F match the theoretical results very closely for all trip probability values.

The benchmarks may be sorted into six suites of tests. The individual suites will be described and examples from each of these series will be presented, comparing the results produced by ROVER-F both with the results produced by ROVER/REFORM and with the theoretical results. In all cases the integration increment used in ROVER/REFORM trip probabilities has a fixed value of 1%.

Test Descriptions and Results

1. Channel-Random and Common-Random Errors

The first suite of tests determines the effect of varying the channel- and common-random errors with a zero detector random error. Because of calculational limitations, the detector-random error cannot be set to exactly zero and therefore is set to the low uncertainty limit of 0.005%. The channel- and common-random errors are then varied to determine the effect on the accuracy of the trip-probability calculation.

These tests work on a simplified ROP model. A single fuel channel is put at risk, by applying a special rippled channel power map that results in all channels but the single channel in question having a critical power ratio much greater than the limit. Similarly, a single detector in only one safety channel (in each shutdown system) is used. The remaining detectors within the safety channel are set in such a manner that they never trip, and the detector readings in the other two safety channels are set to very high values, ensuring that they trip. Thus the trip-probability will be dependent on a single detector.

The tests themselves consist of a series of cases. Each successive case increases the detector readings, thereby increasing the trip probability. In the benchmark definitions, the increase in the detector reading is 1% in cases 2 to 21, and 0.5% in cases 22 to 43.

In the first test, the channel-random error and common-random uncertainty both are set to 5%, relatively large errors as compared to those typically experienced in CANDU 6 ROP analysis. As can be seen from Figure 1, with 1% increments ROVER-F produces results comparable to those of ROVER/REFORM.

ROVER-F has slightly better accuracy at low trip-probabilities, but at high trip-probabilities they are comparable. Figure 2 demonstrates the increase in accuracy when smaller step sizes are used. The 'stepping' of the ROVER-F calculated trip-probabilities in the later cases of Figure 1 are a result of the step size of the integration increment being close in value to the step size of the change in detector reading from case to case. These are not in evidence in the calculations with the fine integration increment.

In the second test, the channel-random error and common-random uncertainty both are set to 0.5%, small errors as compared to typical CANDU 6 values. As can be seen from Figure 3, with 1% increments ROVER-F produces somewhat more accurate results (as compared to the theoretical solution) than ROVER/REFORM. The negative values calculated for low trip probability values are a result of the comparable size of the uncertainties and the integration step size and are not physical. Figure 4 demonstrates a great increase in accuracy when smaller step sizes are used, as the difference between ROVER-F and the theoretical result is negligible for all values.

In the final test presented in this suite, the channel-random error is set to the theoretical minimum and the common random uncertainty is set to 5%. As can be seen from Figure 5, with 1% increments ROVER-F produces somewhat more accurate results (as compared to the theoretical solution) than ROVER/REFORM, particularly at low trip-probabilities. Figure 6 demonstrates a further increase in accuracy when smaller step sizes are used, as the difference between ROVER-F and the theoretical result becomes negligible for all values.

2. Bins and Ripple Conservatism

The second suite of tests examines the effect of various rippled power maps. Ripple maps are applied to the channel power maps to cause various effects on the binning procedure at a near-zero detector-random error and common-random error with nominal channel-random errors. There are two separate tests, which examine different binning effects.

The probability of an individual channel exceeding its critical channel power is a function of the channel random uncertainty and the critical power ratio of the limiting fuel channel (the minimum critical power ratio); and to diminishing degrees, the critical power ratios of the remaining fuel channels. As the critical power ratio of the channels increases, it becomes less probable that a channel will exceed the critical channel power. This suite of tests examines the effect of different groupings of non-limiting critical channel powers. As previously mentioned, these calculations are performed for 'bins' or groups of fuel channels all regarded as having the same critical channel power. The limiting channel is defined as the first bin. These tests apply different numbers of channels to the remaining bins (representing the non-limiting channels) to determine their effect on calculation accuracy.

The first test of this suite examines the effect of cumulative loading of bins. In this case, the number of channels in all the bins is increased cumulatively. Each successive case (44 to 65) includes more channels. The results of these calculations are presented in Figures 7 and 8. Again ROVER-F shows good agreement with ROVER/REFORM.

In the second test in this suite, the previous test is reversed, with successive channels decreasing in rippled CPR ratio. The results of these calculations are presented in Figures 9 and 10. Again ROVER-F shows better agreement with the theoretical trip probabilities than does ROVER/REFORM.

In both of these tests, the residual error between the ROVER-F results and the theoretical results is due to the common-random uncertainty (0.5%). As this value is decreased, the error should also decrease. However, this will require a corresponding decrease in the integration increment. Thus the results with the 0.04% step size could be improved through the use of a smaller common-random uncertainty.

3. Scaling Tests

If the critical power ratio is scaled by some factor, the relative error approach requires that the setpoints should be scaled by the same factor to maintain the same trip probability. This was verified in two ways. In the first, the critical channel power correction factors and setpoints were scaled by factors of 10 and 0.1. In the second test, the systematic error was set to +10% and -10%, with the setpoints again adjusted equivalently. The results are compared between the four tests and the trip probability calculation with no adjustments.

As can be seen from Figure 11, the results of these tests are identical even at low integration resolution.

4. Detector Uncertainty

The fourth suite of tests investigates the effect of varying the detector-random errors, similar to the first suite, except that here the channel-random and common-random uncertainties are held to very small values, to amplify any mathematical errors in the calculation of the error function. The channel-random error is set to the random uncertainty limit of 0.005%, and the common random error is set to a low value of 0.5%. The trip probabilities are then calculated for detector uncertainties ranging from 1% through 10%. The theoretical results for both the common-random and channel-random uncertainty are set to zero.

Figures 12 and 13 present the results of the trip-probability calculations for a detector uncertainty of 1%, although, as in earlier tests, the similar size of the uncertainties and the integration step size results in integration errors (negative values) at low trip probabilities. The results improve to exactly the theoretical results when the integration increment is decreased. Figures 14 and 15 present the results for a detector uncertainty of 10% and demonstrate similar results to those for the smaller detector error.

5. Detector Redundancy

The fifth suite of tests examines the effects of having two detectors active.

Detector redundancy is examined in two ways. First, a previous case is recalculated with two detectors active in the chosen safety channel. Second, one detector in each of two safety channels is made active. In both cases, all other detectors in the targeted safety channel(s) are set in such a manner that they can never be tripped, and the other safety channel(s) are set as always tripped.

The results of both tests are presented in Figures 16 to 19. In all cases, the trip probabilities calculated by ROVER-F correspond closely with the theoretical results, with the error associated with the integration step size at low trip-probabilities tending to zero as the integration step size is decreased.

6. Ripple Averaging

The final suite of tests determines the effect of ripple averaging. The trip probability calculation averages the results of the ripples applied to each case. In theory, the trip-probability for any case should be the same whether the trip-probability is averaged over all ripples or whether the trip-probabilities for each ripple, calculated separately, are averaged.

Figure 20 demonstrates that this holds true for CANDU 6 data. This test is performed for the 232 basis flux shapes and 50 ripples. The calculated trip probability, as an average of the trip probabilities for the individual ripples, is slightly smaller than the trip probability for the average of the ripples.

Conclusions

In conclusion, ROVER-F, the FORTRAN version of the trip probability calculation module of ROVER/REFORM has been validated using standardized benchmark problems. The results of trip probability calculations made with ROVER-F are at least as accurate as those made with ROVER/REFORM over the entire calculation range of the benchmarks. Some cases show improvement in matching theoretical results over ROVER/REFORM results, particularly at low trip probability. Calculations made with small integration increment give very good agreement with the theoretical calculations, at the cost of increased calculation time.

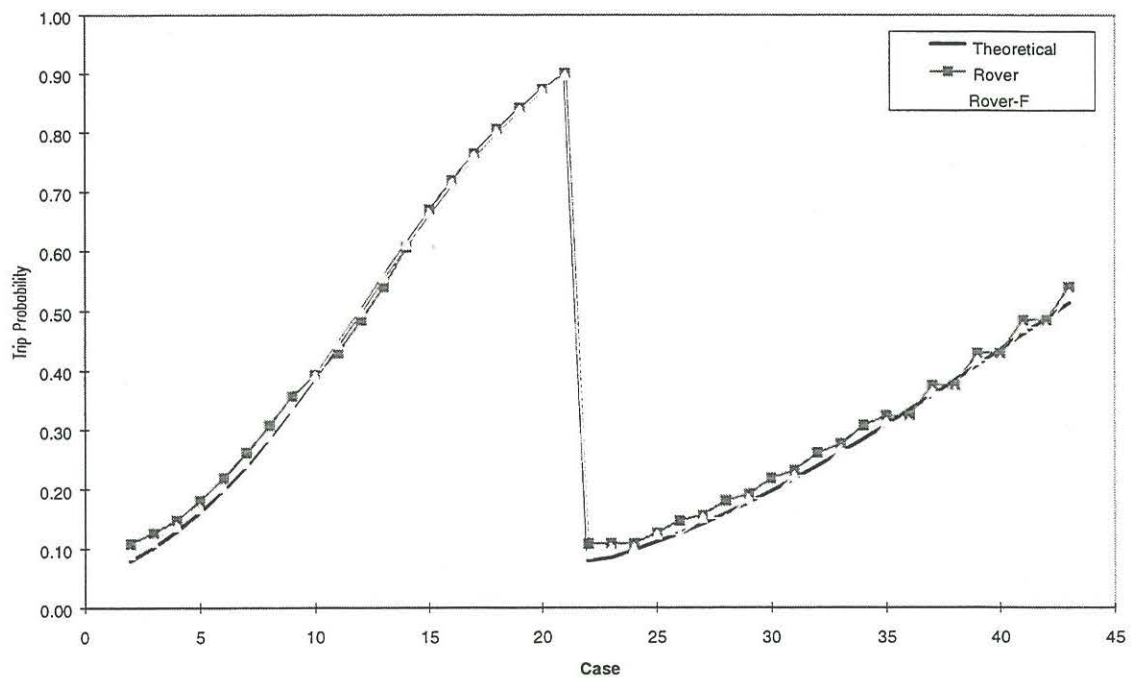


Figure 1: Channel Random $\sigma_{ch} = 5\%$, and Common Random Errors $\sigma_{cm} = 5\%$. 1% Integration Step.

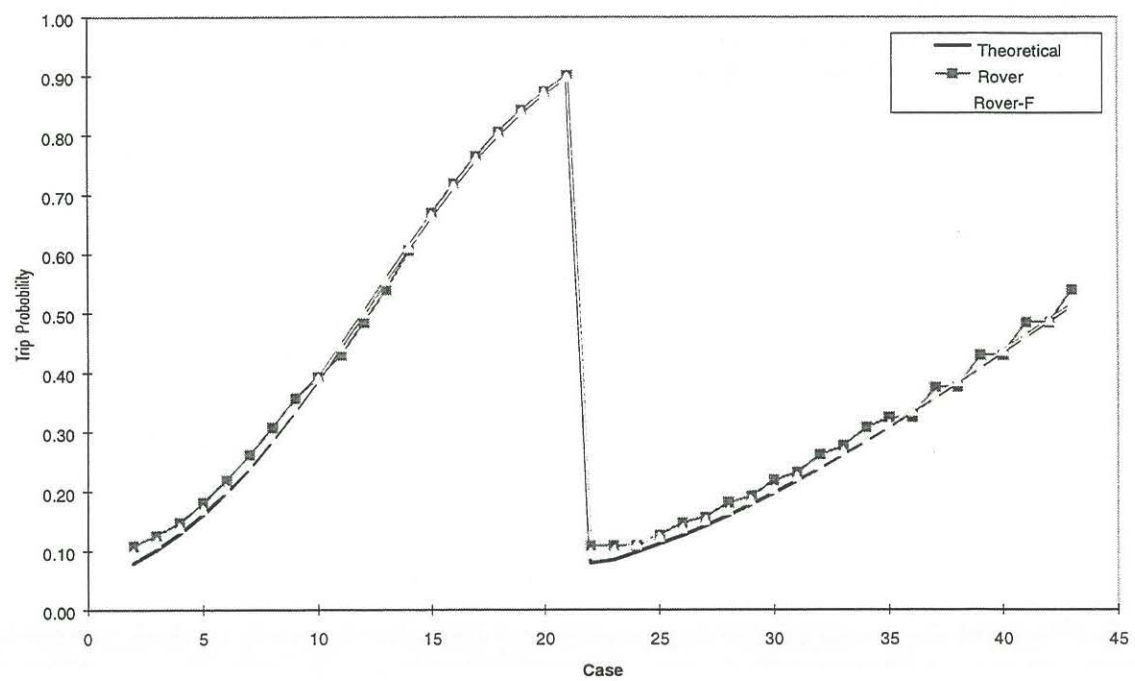


Figure 2: Channel Random $\sigma_{ch} = 5\%$, and Common Random Errors $\sigma_{cm} = 5\%$. 0.04% Integration Step.

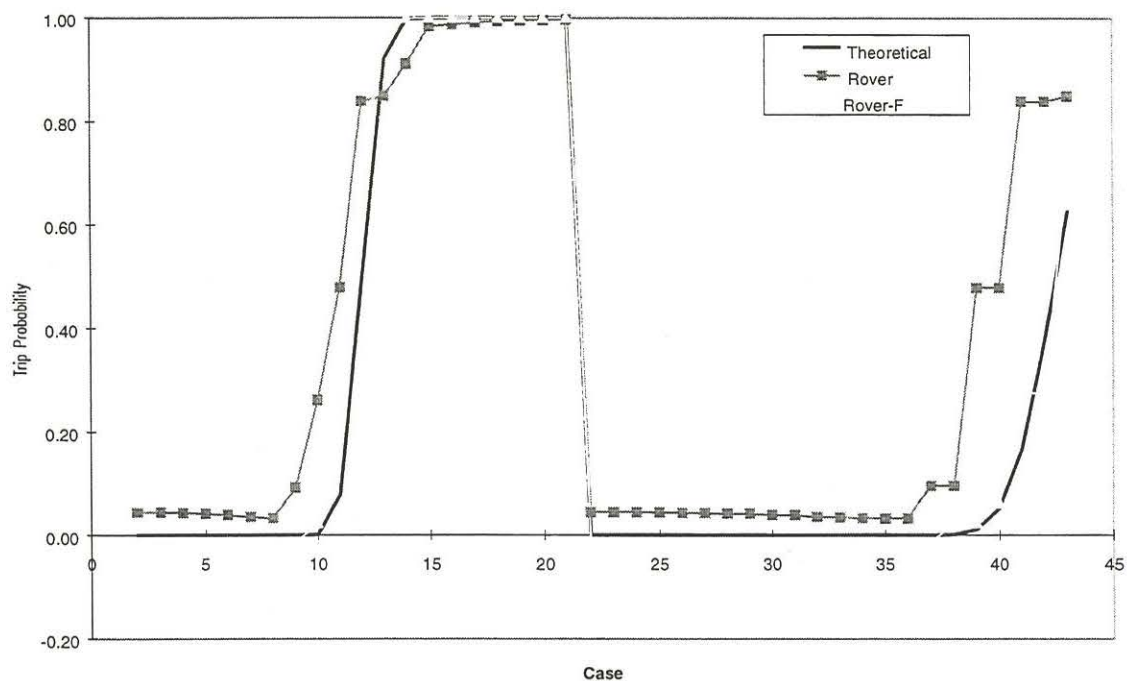


Figure 3: Channel Random $\sigma_{ch} = 0.5\%$, and Common Random Errors $\sigma_{cm} = 0.5\%$. 1% Integration Step.

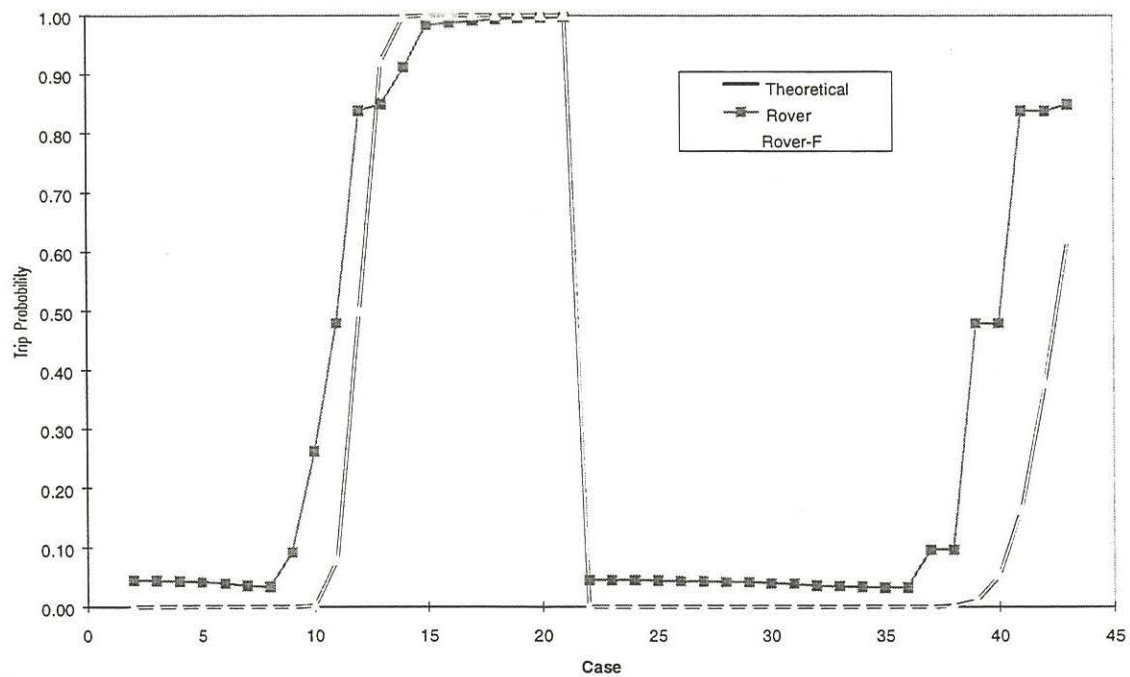


Figure 4: Channel Random $\sigma_{ch} = 0.5\%$, and Common Random Errors $\sigma_{cm} = 0.5\%$. 0.04% Integration Step.

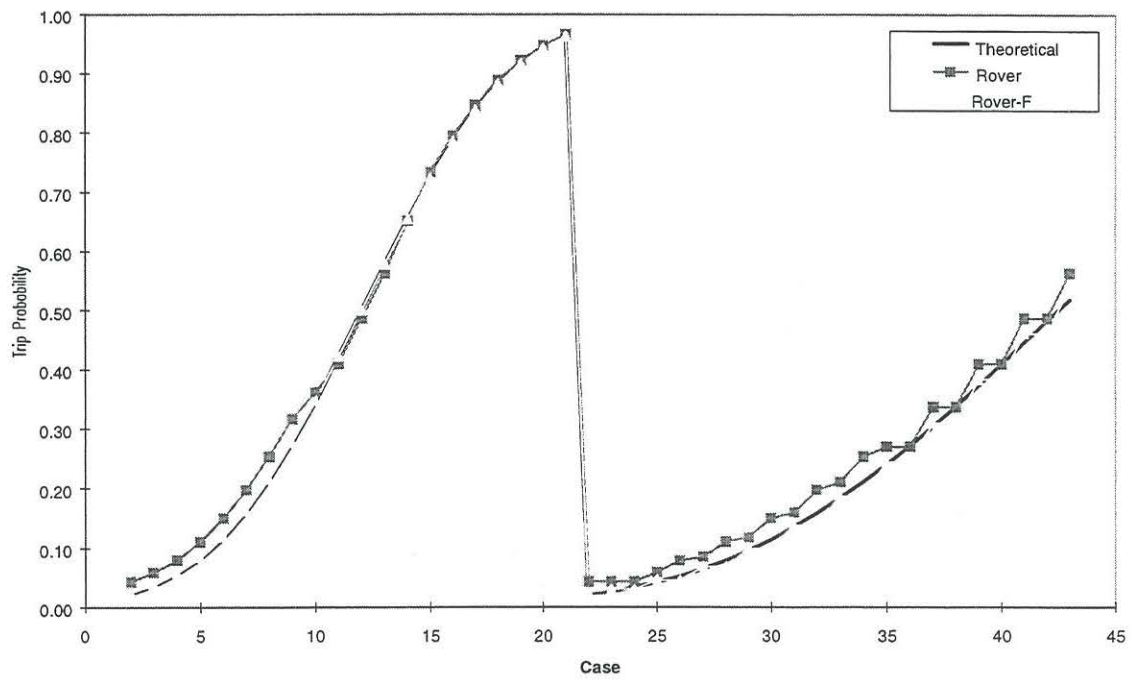


Figure 5: Channel Random $\sigma_{ch} = 0.005\%$, and Common Random Errors $\sigma_{cm} = 5\%$. 1% Integration Step.

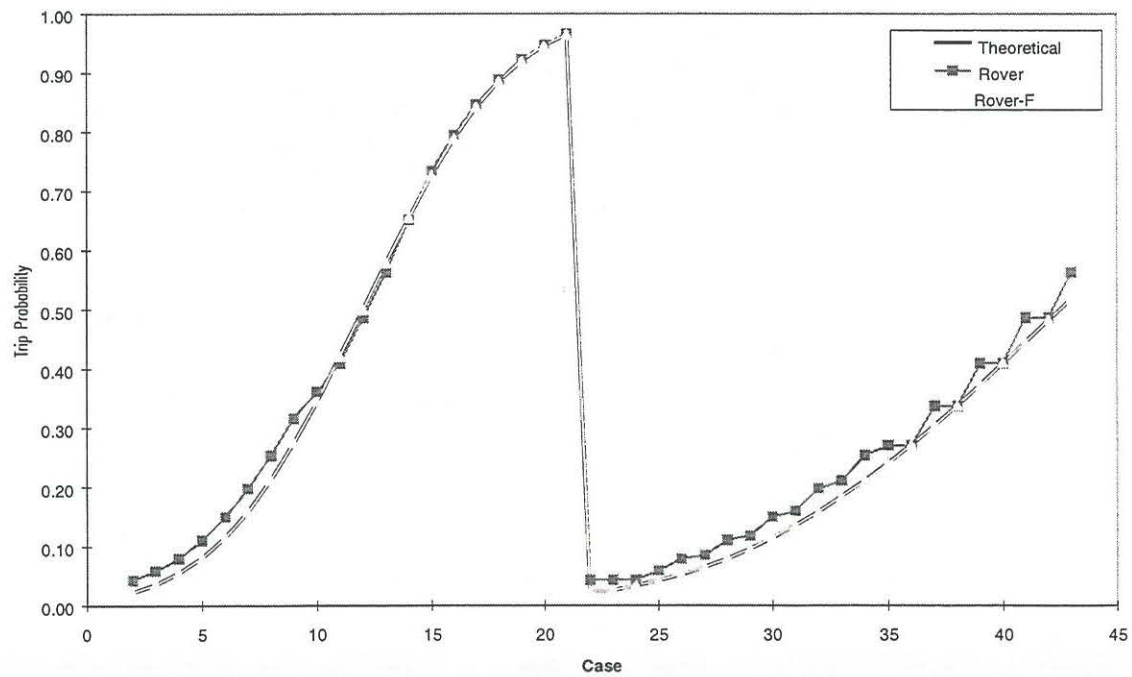


Figure 6: Channel Random $\sigma_{ch} = 0.005\%$, and Common Random Errors $\sigma_{cm} = 5\%$. 0.04% Integration Step.

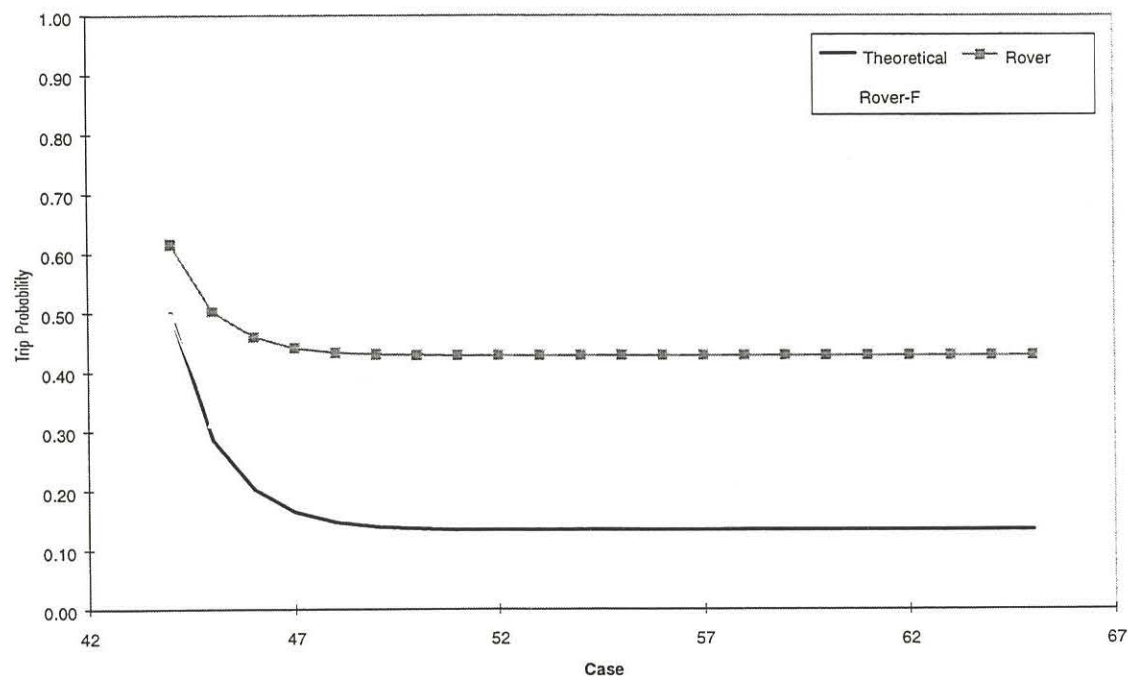


Figure 7: Bins and Ripple Conservatism, 1% Integration Steps, Cumulative Binning

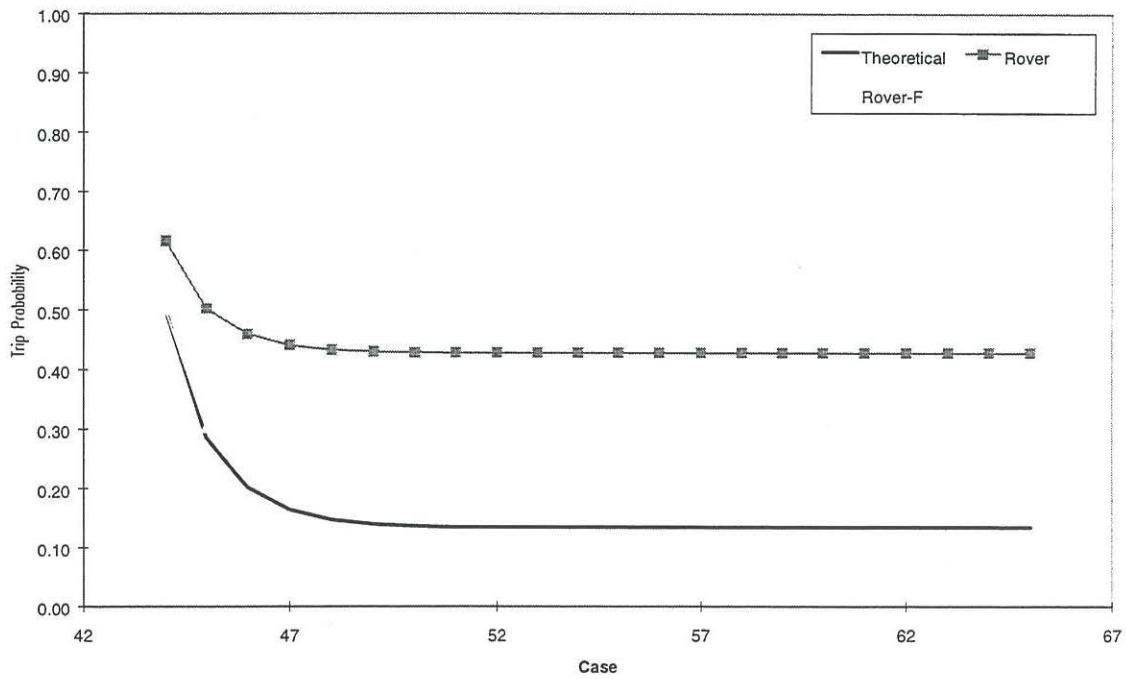


Figure 8: Bins and Ripple Conservatism, 0.04% Integration Steps, Cumulative Binning

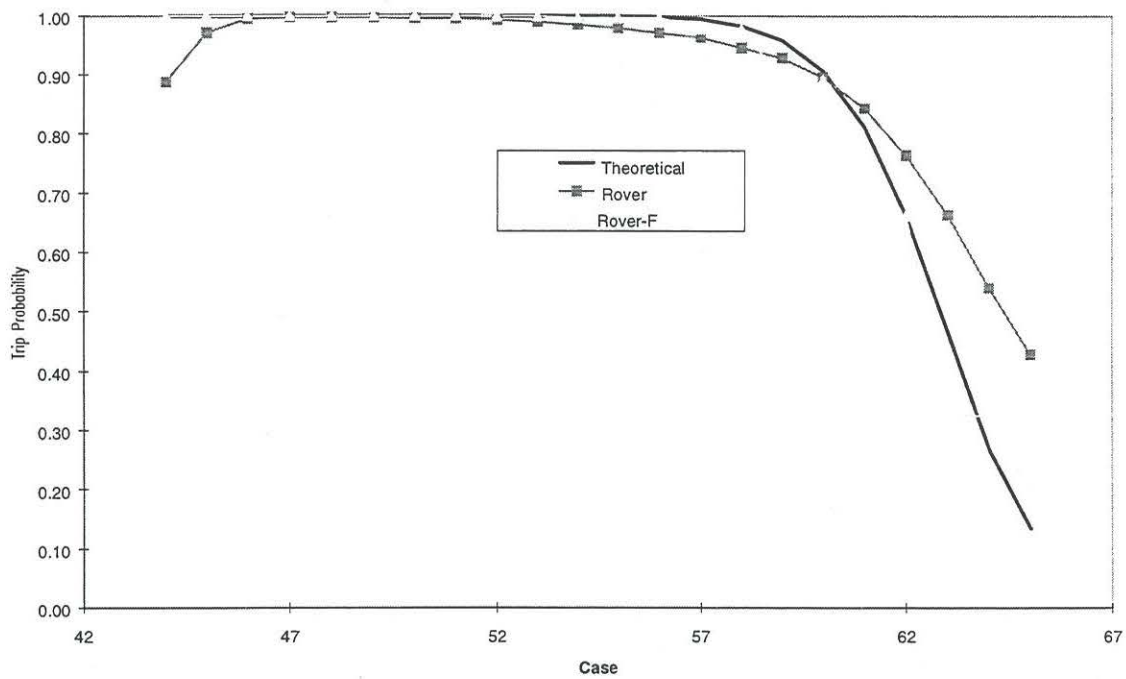


Figure 9: Bins and Ripple Conservatism, 1% Integration Steps, Decremental Binning

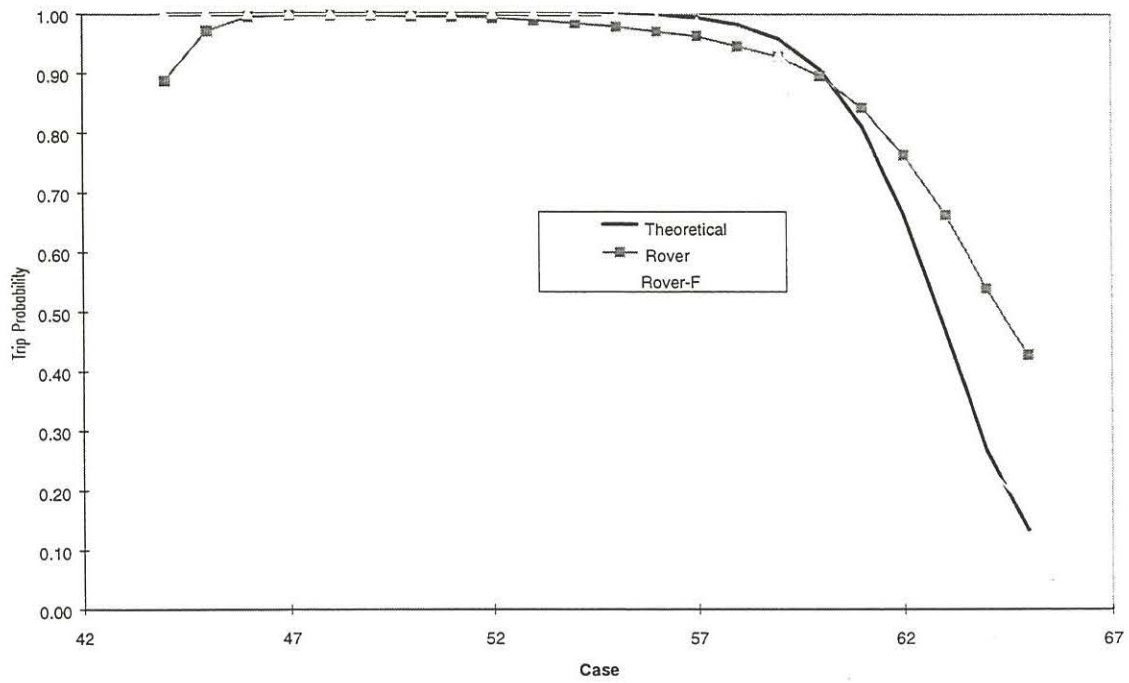


Figure 10: Bins and Ripple Conservatism, 0.04% Integration Steps, Decremental Binning

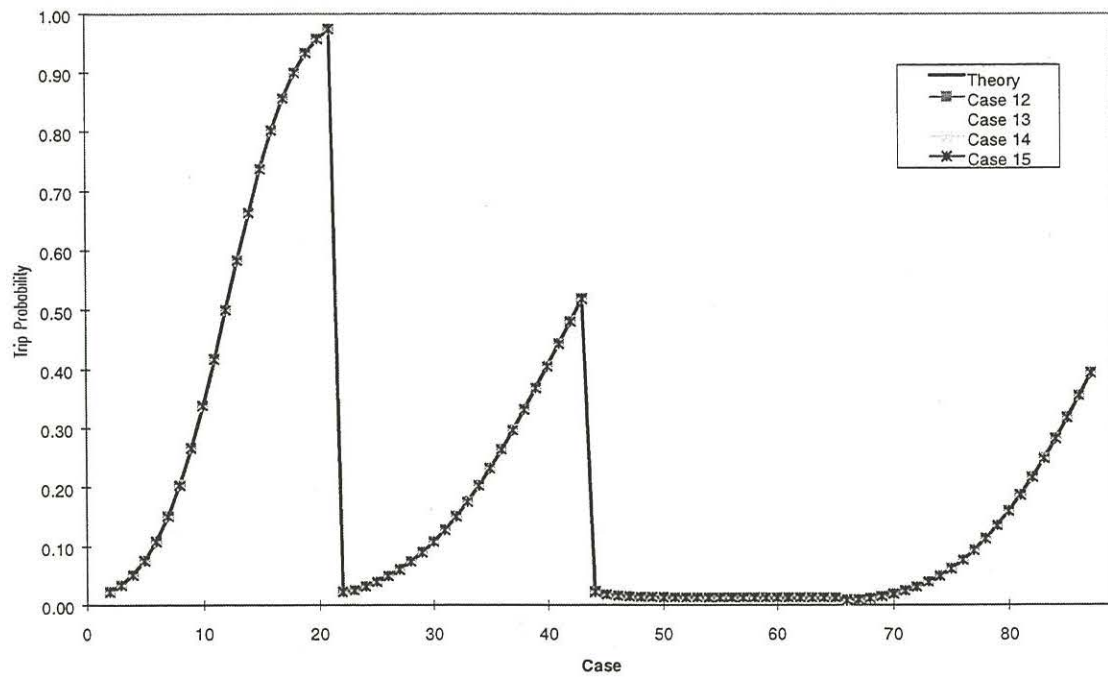


Figure 11: Scaling Tests, 1% Integration Steps

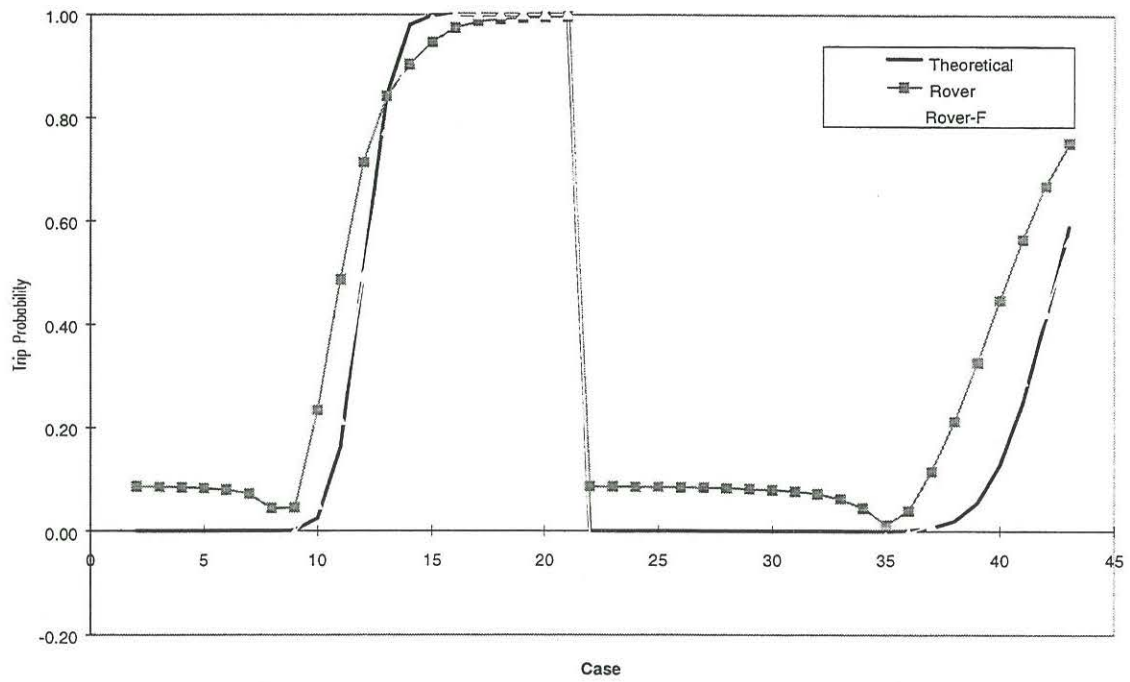


Figure 12: Detector Uncertainty $\sigma_{\text{det}} = 1\%$, 1% Integration Steps

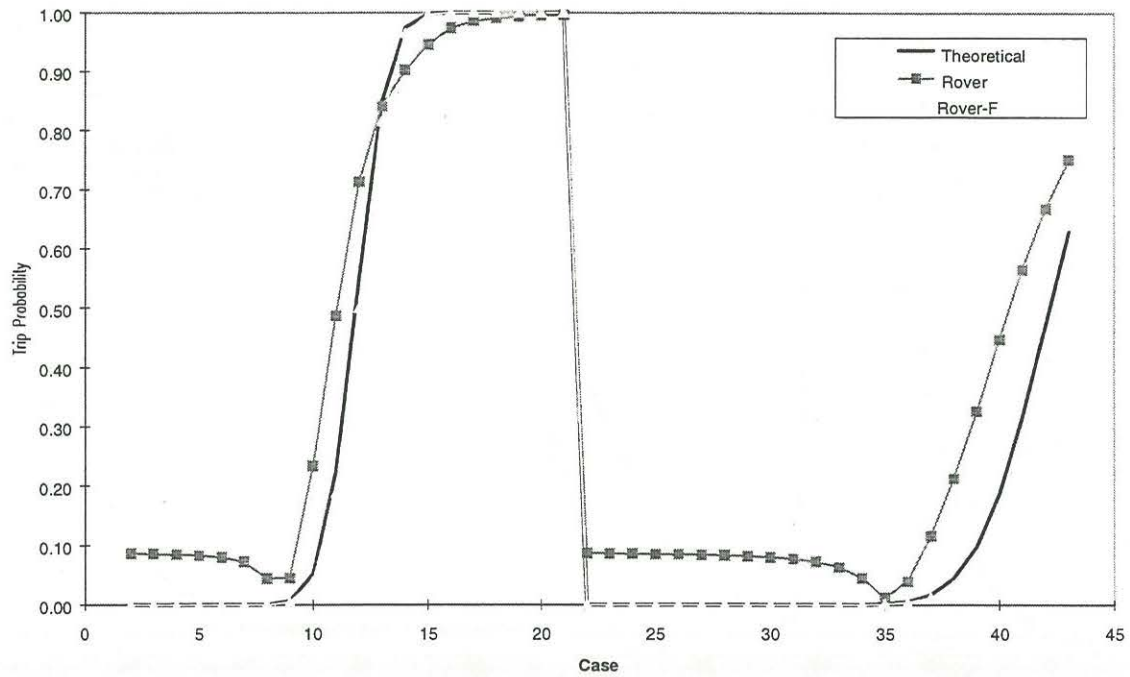


Figure 13: Detector Uncertainty $\sigma_{\text{det}} = 1\%$, 0.04% Integration Steps

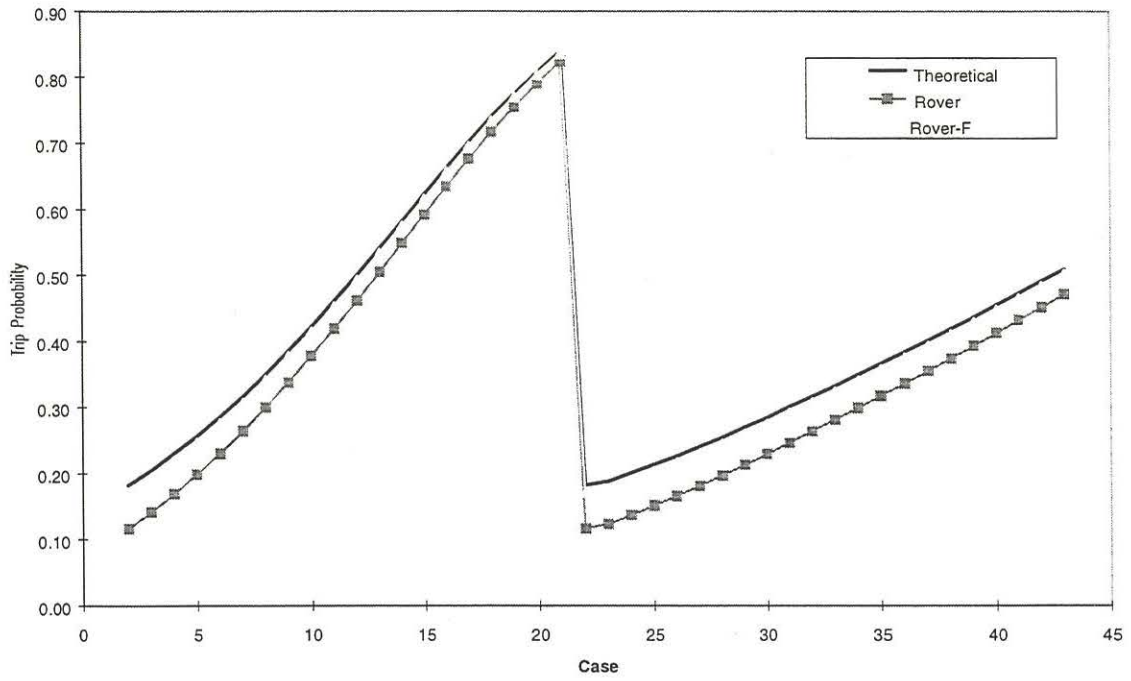


Figure 14: Detector Uncertainty $\sigma_{\text{det}} = 10\%$, 1% Integration Steps

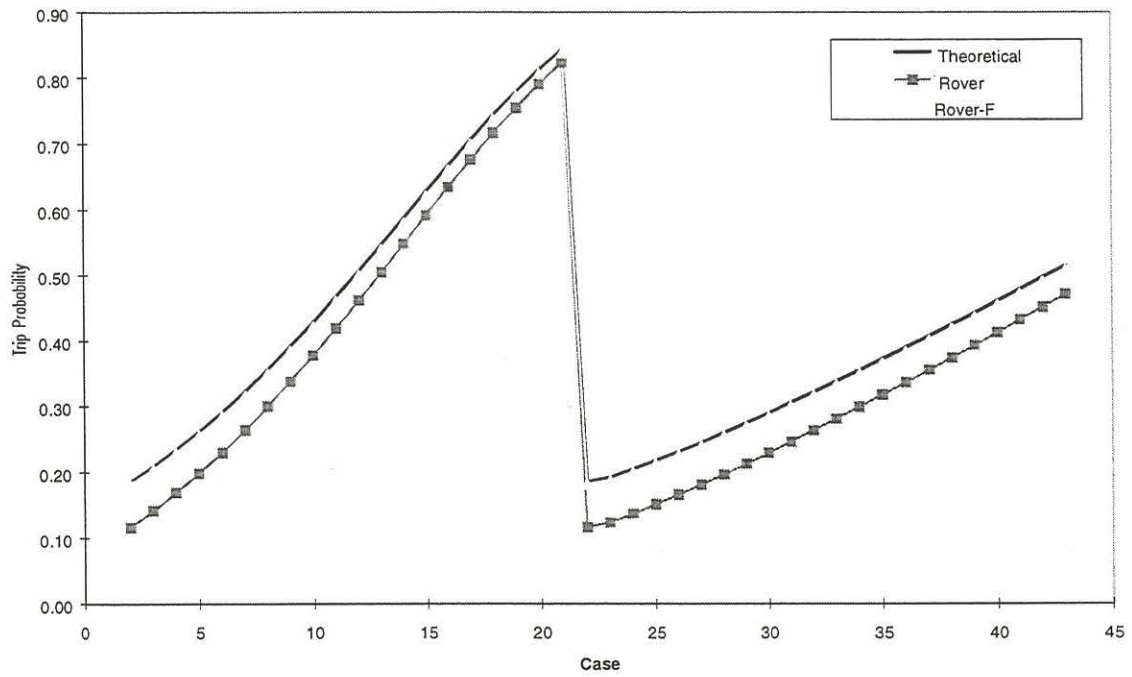


Figure 15: Detector Uncertainty $\sigma_{\text{det}} = 10\%$, 0.04% Integration Steps

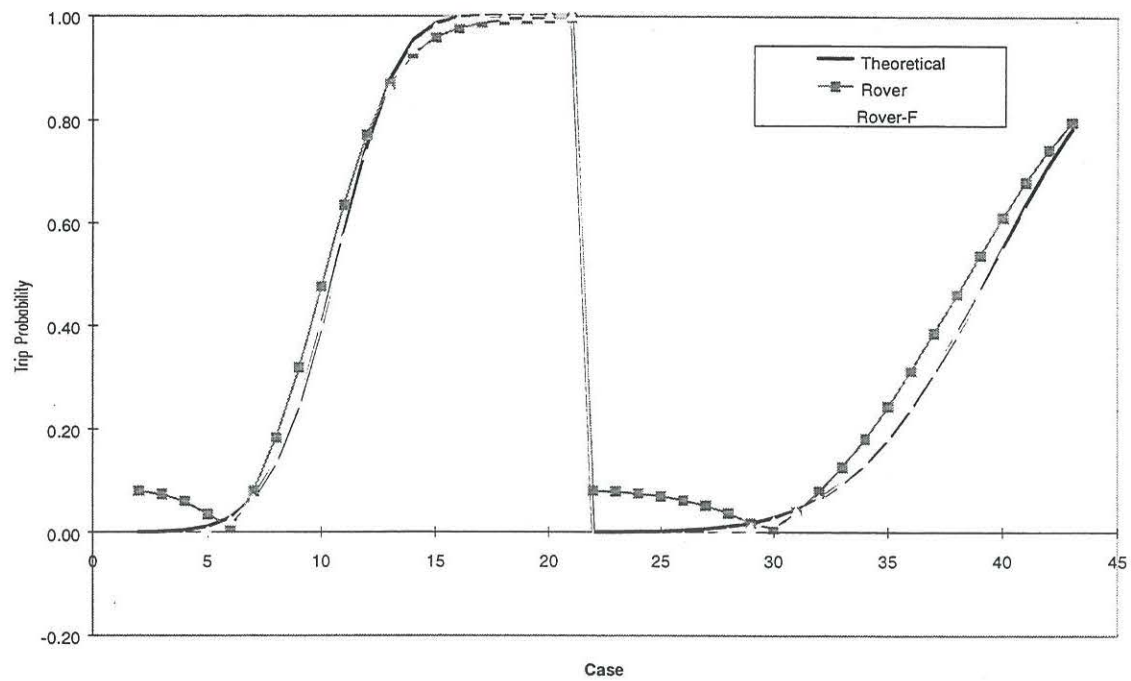


Figure 16: Detector Redundancy, 1% Integration Steps, Two Detectors in Channel D

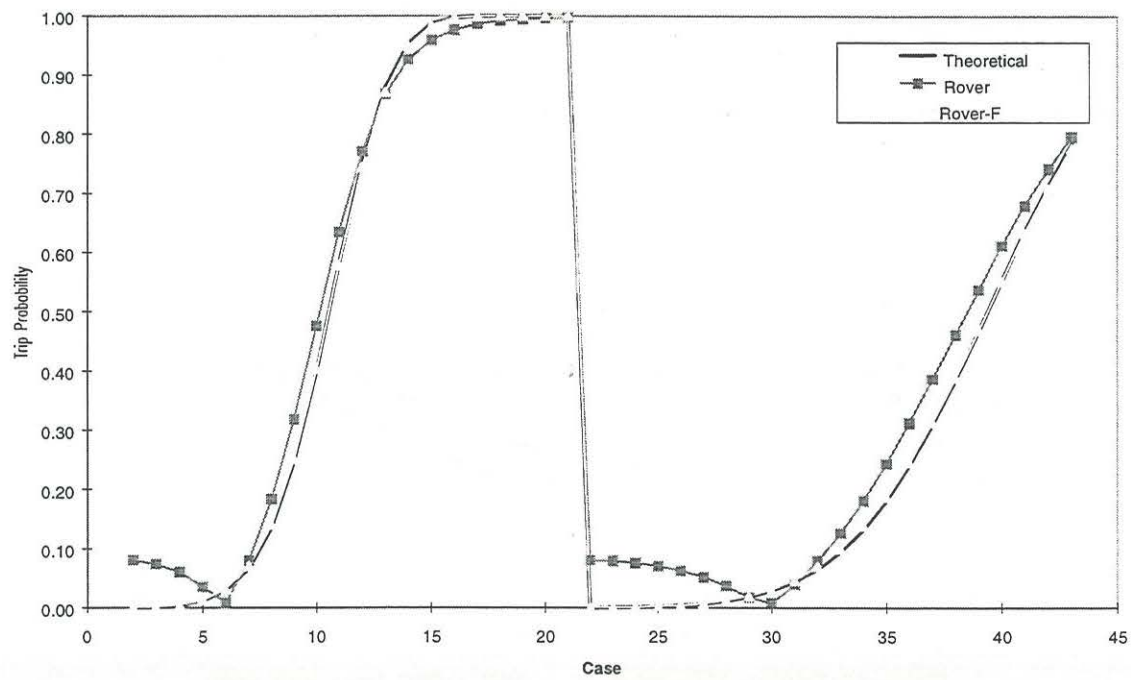


Figure 17: Detector Redundancy, 0.04% Integration Steps, Two Detectors in Channel D

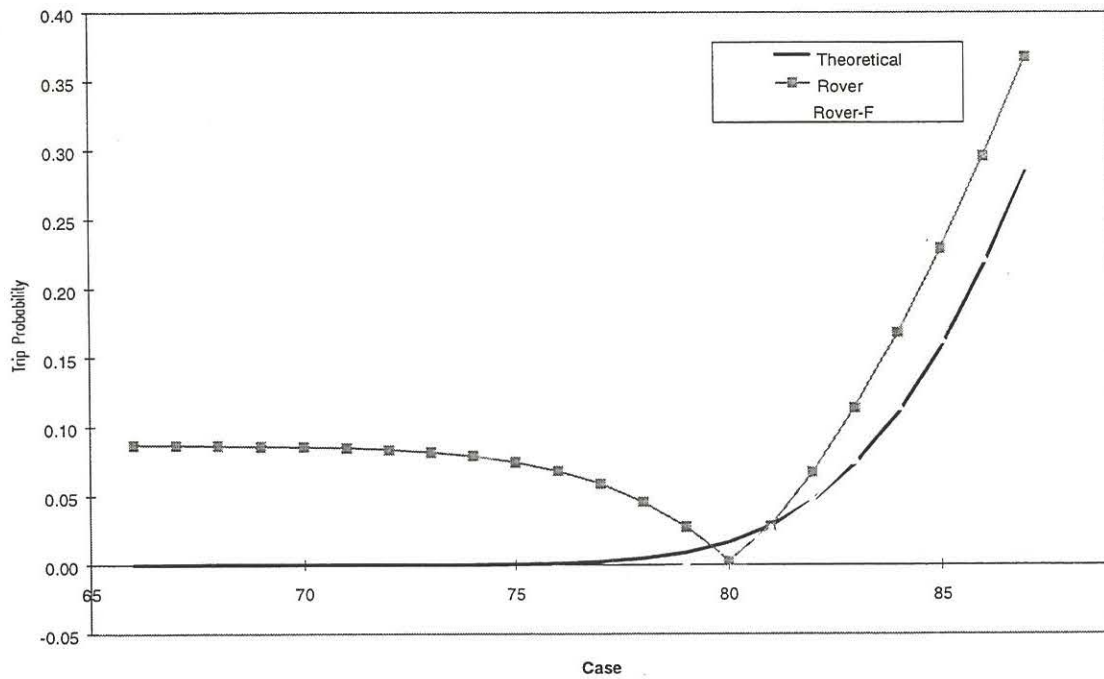


Figure 18: Detector Redundancy, 1% Integration Steps, One Detector Each in Channel D and E

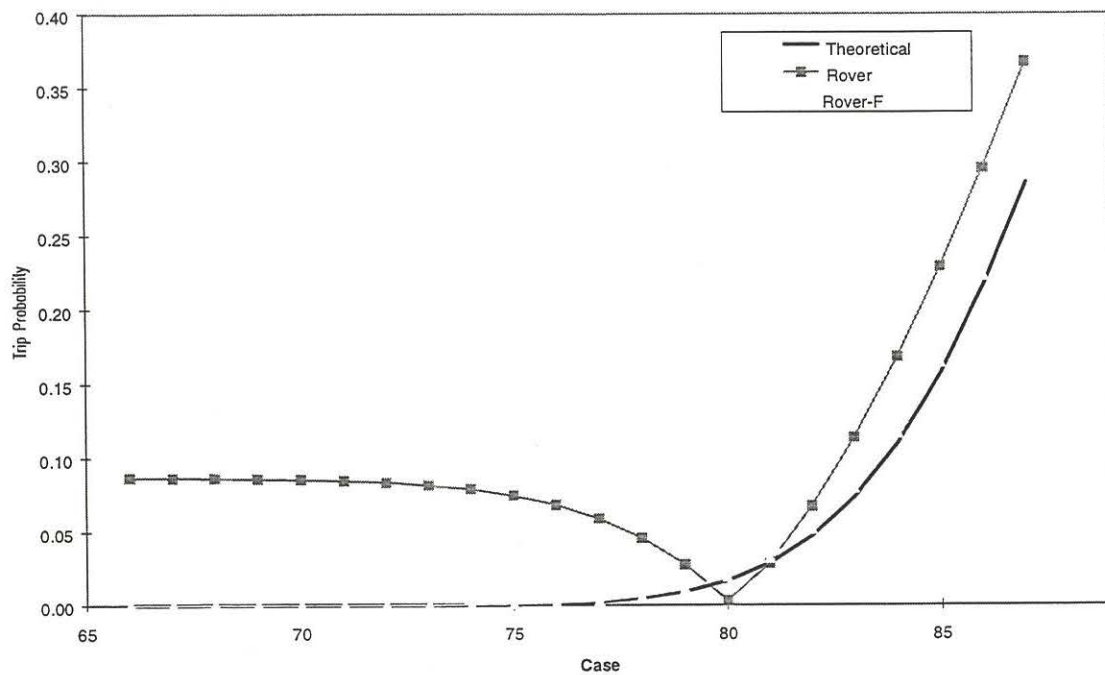


Figure 19: Detector Redundancy, 1% Integration Steps, One Detector Each in Channel D and E

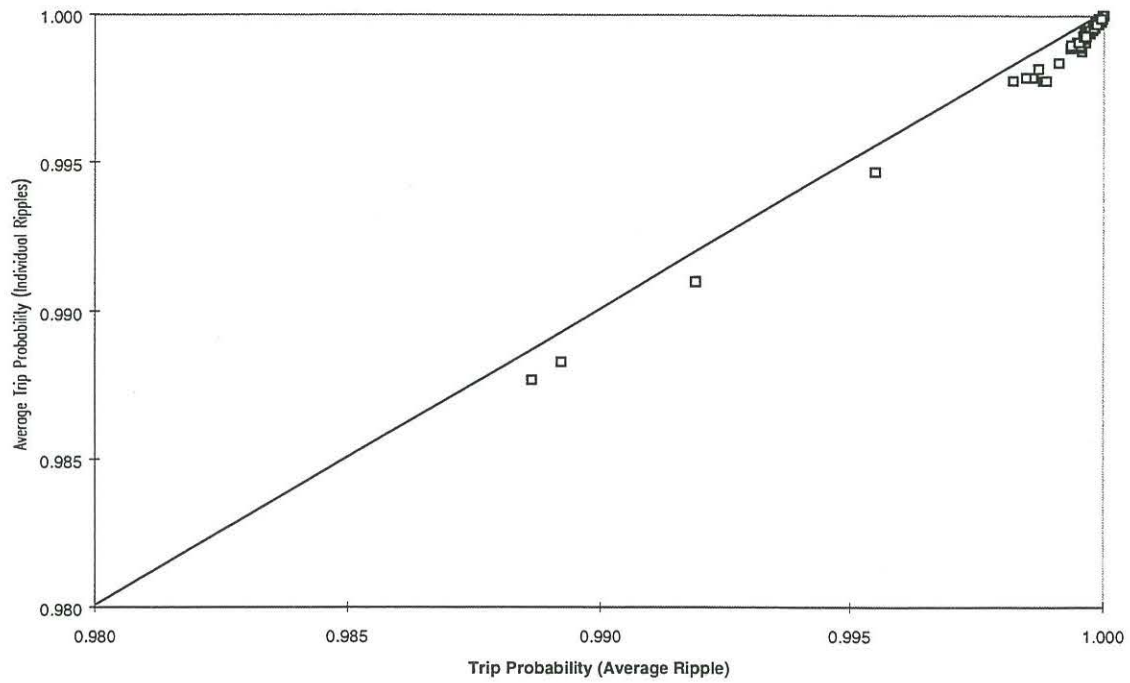


Figure 20: Ripple Averaging, 1% Integration Steps

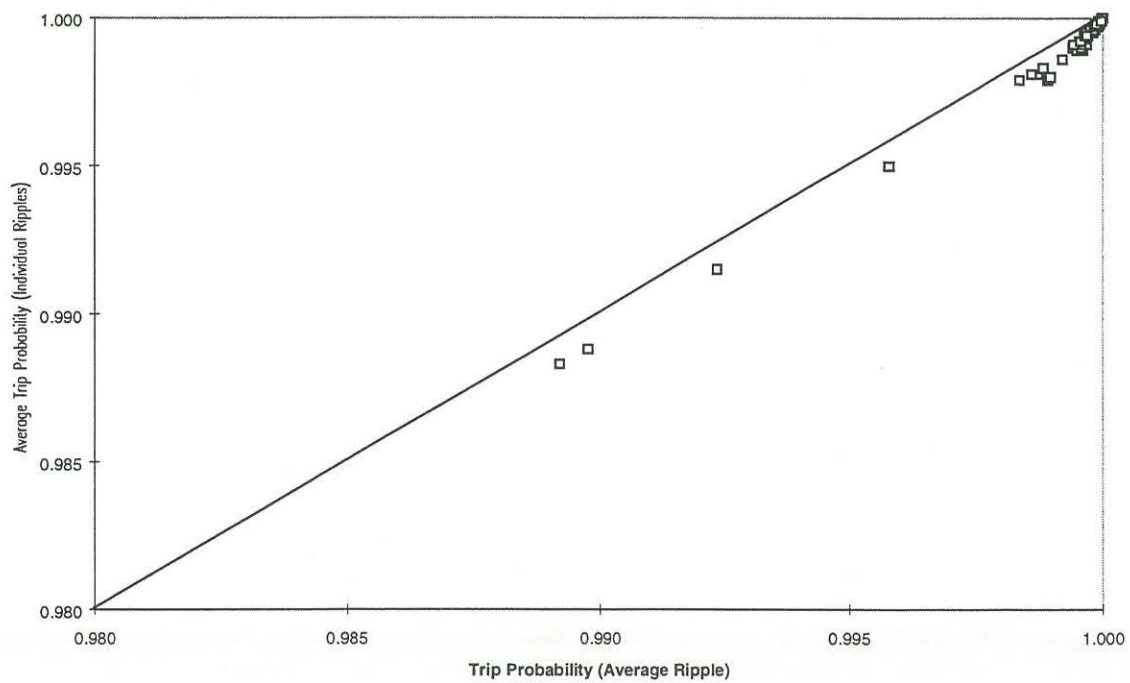


Figure 21: Ripple Averaging, 0.04% Integration Steps