

DEPARTMENT OF STATISTICS
STANFORD UNIVERSITY
STANFORD, CALIFORNIA

THE ANDERSON-DARLING STATISTIC

BY

MICHAEL A. STEPHENS

TECHNICAL REPORT NO. 39

OCTOBER 31, 1979

PREPARED UNDER GRANT
DAAG29-77-G-0031
FOR THE U.S. ARMY RESEARCH OFFICE

Reproduction in Whole or in Part is Permitted
for any purpose of the United States Government
Approved for public release; distribution unlimited.

DEPARTMENT OF STATISTICS
STANFORD UNIVERSITY
STANFORD, CALIFORNIA



THE ANDERSON-DARLING STATISTIC

By

Michael A. Stephens

TECHNICAL REPORT NO. 39

OCTOBER 31, 1979

Prepared under Grant DAAG29-77-G-0031

For the U.S. Army Research Office

Herbert Solomon, Project Director

Approved for public release; distribution unlimited.

DEPARTMENT OF STATISTICS
STANFORD UNIVERSITY
STANFORD, CALIFORNIA

Partially supported under Office of Naval Research Contract N00014-76-C-0475
(NR-042-267) and issued as Technical Report No. 278.

The findings in this report are not to be construed as an official Department of the Army position, unless so designated by other authorized documents.

THE ANDERSON-DARLING STATISTIC

1. Introduction.

The Anderson Darling Statistic is a member of the group of Goodness-of-Fit statistics which has come to be known as EDF statistics (Stephens, 1974) because they are based on a comparison of the empirical distribution function of a given sample with the theoretical distribution to be tested. It is designed to test that random variable X has a continuous cumulative distribution $F(x;\theta)$; θ is a vector of one or more parameters entering into the distribution function. Thus for the normal distribution, the vector $\theta = (\mu, \sigma^2)$.

The empirical distribution function (EDF) is defined as

$$F_n(x) = \frac{\text{number of sample values} \leq x}{n},$$

where the n values $x_1, x_2, \dots, x_n$ are assumed to be a random sample of X . From the x_i let $x_{(1)}, x_{(2)}, \dots, x_{(n)}$ be the order statistics, in ascending order. $F_n(x)$ is then defined by

$$F_n(x) = 0 \quad , \quad x < x_{(1)}$$

$$F_n(x) = i/n \quad , \quad x_{(i)} \leq x < x_{(i+1)}, \quad i = 1, \dots, (n-1)$$

$$F_n(x) = 1 \quad , \quad x_{(n)} < x.$$

Since $F_n(x)$ gives the proportion of a random sample $< x$, one might expect it to give a good estimate of $F(x;\theta)$, which is the probability of X less than x , and $F_n(x)$ is in fact a consistent estimator. It is therefore natural to test whether the sample appears to come from $F(x;\theta)$ by using a statistic based on the discrepancy between $F_n(x)$ and $F(x;\theta)$.

Many statistics of this type have been proposed, the most famous, and one of the oldest, being the Kolmogorov statistic D . This statistic is based on the largest vertical discrepancy between the two functions. An alternative measure is the Cramér-von Mises family, based on the squared integral of the difference between the EDF and the distribution tested:

$$W^* = \int_{-\infty}^{\infty} \{F_n(x) - F(x;\theta)\}^2 \psi(x) dx ; \quad (1)$$

the function $\psi(x)$ gives a weighting to the squared difference. One member of W^* is the Cramér-von Mises statistic itself, W^* with $\psi(x) = 1$.

The Anderson Darling statistic, the subject of this article, is W^* with

$$\psi(x) = [F(x;\theta)\{1 - F(x;\theta)\}]^{-1} .$$

This weight function counteracts the fact that the discrepancy between $F_n(x)$ and $F(x;\theta)$ is necessarily becoming smaller in the tails, since both approach 0 and 1 at the extremes. The weight function given weights the discrepancy by a factor inversely proportional to its variance, and has the effect of giving greater importance to observations in the tail than do most of the EDF statistics. Since tests of fit are often needed implicitly or explicitly to guard against wayward observations in the tails, the statistic is a recommended one, with, as we shall see, generally

good power properties over a wide range of alternative distributions when $F(x;\theta)$ is not the true distribution.

2. Computing Formula.

For practical purposes, the definition of the Anderson-Darling statistic given above needs to be turned to a computational formula.

This is done in the following sequence of steps:

(a) Calculate $z_i = F(x_{(i)};\theta)$, $i = 1, \dots, n$.

(b) The Anderson-Darling statistic is given by

$$A^2 = -\left\{ \sum_{i=1}^n (2i-1) [\ln z_i + \ln(1-z_{n+1-i})] \right\} / n - n. \quad (2)$$

Note that since the $x_{(i)}$ are in ascending order, the z_i will also be in ascending order, though the usual notation of order statistics has been omitted.

3. Goodness-of-fit test for a completely specified continuous distribution.

The formula for z_i above assumes that the tested distribution $F(x;\theta)$ is completely specified, i.e., the parameters in θ must be known. When this is the case we describe the situation as Case 0. The statistic A^2 was introduced by Anderson and Darling (1952, 1954), and for Case 0 they gave the asymptotic distribution and tables of percentage points. For testing purposes the upper tail of A^2 will be used; large discrepancies between the EDF and the tested distribution will indicate a bad fit. Later, Lewis (1961) demonstrated that the distribution of A^2 for a finite sample approaches the asymptotic distribution extremely quickly, so that for practical purposes only the asymptotic distribution is required for

sample sizes greater than 5. A table of percentage points is given in Table 1. To make the goodness of fit test, A^2 is calculated as in Equation (2) above, and compared with these percentage points; the null hypotheses that random variable X has the distribution $F(x;\theta)$ is rejected at level α if A^2 exceeds the appropriate percentage point at this level.

4. Asymptotic theory of the Anderson-Darling statistic.

The distribution of A^2 for Case 0 is the same for all distributions tested. This is because the probability integral transformation is made at step (a) and the values of z_i are ordered values from a uniform distribution with limits 0 and 1. A^2 is therefore a function of ordered uniform random variables. The asymptotic distribution theory for this special case can be found from the asymptotic theory of the EDF, or more specifically of the function

$$y_n(z) = \sqrt{n}(F_n(z) - z) ,$$

where $F_n(z)$ is the EDF of n uniform random variables as above. For a modern treatment of the empirical process given by $y_n(z)$ see Durbin (1973a, 1973b). When θ contains unknown components, the z_i given by the transformation (a) above, when an estimate $\hat{\theta}$ replaces θ , will not be ordered uniform random variables and the distribution theory of A^2 , as for all other EDF statistics, becomes substantially more difficult. In general, the distribution of A^2 , and of other EDF statistics, will depend on n and also on the values of the unknown parameters.

Fortunately, an important simplification occurs if unknown components of θ are location and scale parameters only; then the distribution of each EDF statistic, with an appropriate estimate for θ , will depend on the distribution tested, but not on the specific values of the unknown parameters. Thus for the test for the normal distribution, for example, with unknown μ and σ^2 , only one set of tables would be needed, of percentage points for each n . This simplification makes it worthwhile to calculate the asymptotic theory for A^2 and other EDF statistics, and this has been done for the normal case and the exponential case by Stephens (1974, 1976) and Durbin, Knott and Taylor (1975). Stephens, from Monte Carlo studies to find the distributions of EDF statistics for finite n , has calculated modifications of the basic statistics; these are functions of the statistic and of n which can be used with only the asymptotic percentage points. Thus only one line of percentage points is needed to make the test. The technique is set out below. Stephens (1977, 1978) has also done similar distribution theory to provide tests for the extreme value distribution (these can be used for the Weibull distribution also) and for the logistic distribution. Pettitt and Stephens (1976) have given tests for the Gamma distribution with unknown scale parameter but known shape parameter. From all these results, we set out the technique for goodness-of-fit testing as it applies to A^2 .

5. General procedure for any distribution with unknown location or scale parameter.

The first step in testing goodness-of-fit for any of these distributions is to estimate the unknown parameters. This should be done by

maximum likelihood, for the modifications and asymptotic theory to hold. Suppose that $\hat{\theta}$ is the vector of parameters, with any unknown parameters estimated as above. The steps continue as follows:

- (a) Calculate $z_i = F(x_{(i)}; \hat{\theta})$, $i = 1, \dots, n$.
- (b) Calculate A^2 from the formula (2) above.
- (c) Modify A^2 by the formula in the appropriate table below, and compare with the line of percentage points given.

6. Tests for different distributions.

It is worthwhile to set forth the practical details of these calculations, for each distribution separately.

6.1. Tests for the normal distribution. We here distinguish three cases:

Case 1: The mean μ is unknown and is estimated by $\bar{x}$, but σ^2 is known;

Case 2: The mean μ is known, and σ^2 is estimated by $\sum_i (x_i - \mu)^2 / n (= s_1^2, \text{ say})$;

Case 3: Both parameters are unknown and are estimated by $\bar{x}$ and

$$s^2 = \sum_i (x_i - \bar{x})^2 / (n-1).$$

For these cases the calculation of z_i is done in two stages. First w_i is found from

$$w_i = \frac{x_{(i)} - \bar{x}}{\sigma} \quad (\text{Case 1}); \quad w_i = \frac{x_{(i)} - \mu}{s_1} \quad (\text{Case 2}); \quad w_i = \frac{x_{(i)} - \bar{x}}{s} \quad (\text{Case 3});$$

then z_i is the cumulative probability of a standard normal distribution, to the value w_i , found from tables or computer routines. The value of A^2 is then calculated as described in Section 2. To make the test use the modification and percentage points given in Table 2, for the appropriate case.

Illustration. The following value of men's weights in pounds, first given by Snedecor, were used by Shapiro and Wilk [17] as an illustration of a test for normality: 148, 154, 158, 160, 161, 162, 166, 170, 182, 195, 236. The mean is 172 and the standard deviation 24.95. For a test for normality (Case 3), the values of w_i begin $w_1 = (148-172)/24.95 = -0.962$, and the corresponding z_1 is, from tables, 0.168. When all the z_i have been found, the formula in Section 2 gives $A^2 = 0.947$. Now to make the test, the modification in Table 2 first gives

$$A^* = A^2(1.0 + 0.75/11.0 + 2.25/121.0) = A^2(1.0868) = 1.029 ;$$

when this value is compared with the percentage points in Table 2, for Case 3, the sample is seen to be significant at approximately the 1 percent level.

6.2. Tests for the exponential distribution. The distribution tested is $F(x) = 1 - \exp(-x/\beta)$, $x > 0$, described as $\text{Exp}(x, \beta)$, with β an unknown positive constant. Maximum likelihood gives $\hat{\beta} = \bar{x}$, so that z_i are found from $z_i = 1 - \exp(-x_{(i)}/\bar{x})$, $i = 1, \dots, n$. A^2 is calculated as in Section 2, modified to give A^* by the formula in Table 3, and A^* is compared with the percentage points in Table 3.

For the more general exponential distribution given by $F(x) = 1 - \exp(-(x-\alpha)/\beta)$, $x > \alpha$, when both α and β are unknown, a convenient property of the distribution may be used to return the test situation to the case just described above. The distribution $y_{(i)} = x_{(i+1)} - x_{(1)}$ is made, for $i = 1, \dots, n-1$; the $n-1$ values of $y_{(i)}$ are then used to test that they come from $\text{Exp}(y; \beta)$ as just described. The substitution to $y_{(i)}$ reduces the sample size by one, but eliminates α very straightforwardly.

6.3. Tests for the extreme value distribution. The distribution tested is here $F(x;\theta) = \exp[-\exp\{-(x-\alpha)/\beta\}]$, $(-\infty < x < \infty)$, with $\theta = (\alpha, \beta)$; α and β are constants, β positive. As for the normal distribution we distinguish three cases:

Case 1: β is known and α is estimated;

Case 2: α is known and β is estimated;

Case 3: α and β are both unknown, and must be estimated.

Maximum likelihood estimates of α and β are given by solving equations:

$$\hat{\beta} = \sum_j x_j / n - \{\sum_j x_j \exp(-x_j / \hat{\beta})\} / \{\sum_j \exp(-x_j / \hat{\beta})\}, \quad \hat{\alpha} = -\hat{\beta} \log\{\sum_j \exp(-x_j / \hat{\beta}) / n\}.$$

The first equation is solved iteratively, and then $\hat{\alpha}$ can be found. In Case 1, β is known; then β replaced $\hat{\beta}$ in $\hat{\alpha}$. In case 2, α is known; suppose then that $y = x_i - \alpha, \hat{\beta}$ is given by solving

$$\hat{\beta} = \{\sum_j y_j - \sum_j y_j \exp(-y_j / \hat{\beta})\} / n.$$

These are then used in $F(x;\theta)$ to give z_i and hence A^2 . The modifications and percentage points for the different Cases are given in Table 4.

6.4. Tests for the Weibull distribution. The distribution tested, in its most general form, is

$$F(x;\theta) = 1 - \exp[-\{(x-\alpha)/\beta\}^\gamma], \quad (x > \alpha), \quad (3)$$

with $\theta = (\alpha, \beta, \gamma)$; β and γ must be positive. When α is known, the substitution $Y = -\ln(X-\alpha)$ gives, for the distribution function for Y , $F(y) = \exp[-\exp\{-(y-\alpha')/\beta'\}]$, ($y > \alpha'$), where $\beta' = 1/\gamma$ and $\alpha' = -\ln \beta$, so that Y has the extreme value distribution considered above. A test for the Weibull distribution, with α known but β, γ unknown therefore can be made as follows:

- (a) Find $y_{(i)} = -\ln(x_{(n+1-i)} - \alpha)$, $i = 1, \dots, n$.
- (b) Test that $y_{(i)}$ is a sample (now placed in ascending order by step (a)) from the extreme value distribution with two unknown parameters, as described in Section 6.3, Case 3.

Note also that if, in addition to α , γ is known in (3), the substitution $Y = -\ln(X-\alpha)$ gives an extreme-value distribution for Y with scale parameter β' now known (Case 1 of Section 6.3); if α and β are both known in (3), the substitution gives an extreme-value distribution for Y with location parameter α' known (Case 2 of Section 6.3)

6.5. Tests for the logistic distribution. The distribution tested is $F(x;\theta) = [1 + \exp\{-(x-\alpha)/\beta\}]^{-1}$, ($x > \alpha$), with $\theta = (\alpha, \beta)$; α, β are constants, with β positive. Again three cases are distinguished (Stephens, 1979):

Case 1: β is known, and α must be estimated;

Case 2: α is known, and β must be estimated;

Case 3: Both α and β are unknown and must be estimated.

Maximum likelihood estimates are given for Case 3 by the equations:

$$\sum_i [1 + \exp\{(x_i - \hat{\alpha})/\hat{\beta}\}]^{-1} = n/2 ;$$

$$\sum_i \left(\frac{x_i - \hat{\alpha}}{\hat{\beta}} \right) \left(\frac{1 - \exp\{(x_i - \hat{\alpha})/\hat{\beta}\}}{1 + \exp\{(x_i - \hat{\alpha})/\hat{\beta}\}} \right) = -n .$$

These may be solved iteratively, using, for example, $\bar{x}$ and $s\sqrt{3}/\pi$ as starting estimates of α and β . In Case 1, only the first equation is needed, with β replacing $\hat{\beta}$, and in Case 2 only the second equation is used, with α replacing $\hat{\alpha}$. In the transformation $z_i = F(x_{(i)}; \hat{\theta})$, the estimates $\hat{\alpha}$ and $\hat{\beta}$ are used in $\hat{\theta}$ as necessary, and A^2 is calculated from the formula (2). The modification to A^* , and the percentage points of A^* are given in Table 5.

6.6. Tests for the Gamma distribution with known shape parameter.

The density under test is $f(x; \theta) = \{\Gamma(m)\beta^m\}^{-1} x^{m-1} e^{-x/\beta}$, $x > 0$, and the distribution is $F(x; \theta) = \int_0^x f(t; \theta) dt$. The parameter vector $\theta = (m, \beta)$ contains m as shape parameter and β as scale parameter; note that the test involving an unknown location parameter is not considered. In the test which follow, we assume m is known; β is then estimated by $\hat{\beta} = m/\bar{x}$, where $\bar{x}$ is the sample mean. Estimated density $f(m; \hat{\theta})$ is $f(m; \theta)$ above with $\hat{\beta}$ replacing β , and $F(x; \hat{\theta})$ is defined in a similar way. Then for the goodness-of-fit test, values z_i are calculated from $z_i = F(x_{(i)}; \hat{\theta})$, and A^2 calculated as in Section 2. The modified form A^* , and tables of percentage points for A^* , are given for various m in Table 6.

7. Power of the Anderson-Darling Statistic.

As was described in the introduction, the Anderson-Darling statistic A^2 gives weight to observations in the tails of the distribution tested, whereas other statistics sometimes have the effect of giving less importance to these observations. A^2 can therefore be expected to be powerful in detecting alternatives which have high probability of giving observations in the tails. Several studies have been made, for example, on tests for the uniform distribution with limits at 0 and 1. This is the distribution of the z_i in Section 2, in the Case 0 situation where the tested distribution is completely specified. The type of alternative to uniformity generally considered has distribution function $F(x) = x^k$ or $F(x) = 1 - (1-x)^k$, $0 < x < 1$ with $k > 0$. These distributions produce points which are close to 1 or close to 0 respectively. Simple modifications of the distributions will produce densities with a peak at 0.5, or with the minimum value at 0.5 and high values in either tail. There is generally a clear difference in behavior of EDF statistics in detecting these alternatives. (Stephens, 1974; Quesenberry and Miller, 1977; Locke and Spurrier, 1978). The statistic A^2 will detect alternatives which produce observations towards 0 or 1, but other statistics of the EDF class are more suitable for alternatives which produce a cluster near 0.5 (Stephens, 1974).

In the important situations where parameters must be estimated, the differences in powers of the EDF statistic appear to level out; the opportunity to estimate parameters means that $F(x; \hat{\theta})$ is brought close

to the EDF of the sample, and the z_i values of Section 2 are super-uniform i.e., more regular than a genuine uniform sample. Even with an alternative distribution to the null, the z_i do not take very extreme departures from uniformity. In these circumstances the powers of the various EDF statistics are not so different among themselves as they are for the Case 0 situation; see, e.g., tests for normality reported in Stephens (1974). Nevertheless, because of the importance it gives observations in the tails, A^2 appears overall to be an effective EDF statistics in these situations. It compares very favorably with statistics also devised for testing for special distributions, e.g. the Shapiro-Wilk statistics for testing normality, or exponentiality (Shapiro and Wilk, 1965, 1972; Stephens, 1974, 1978). The statistic A^2 has the merit of being easy to calculate, and, using the modifications, easy to apply with only one line of percentage points for each test situation.

8. Related Topics.

This article has been concerned with the use of A^2 for testing goodness-of-fit of one sample, for a variety of distributions. There occur problems in which several samples, usually of small size, are available, and one wishes to combine the information in those samples to make an overall goodness-of-fit test. Pettitt (1977) has provided tables from which one can obtain the significance level p_i , $i = 1, \dots, k$, of k such tests, in the Case 3 situation of a test for the normal distribution (Section 6.1 above). The values p_i are then combined using Fisher's well known method.

Pettitt (1976) has also given a two sample version of the Anderson-Darling statistic; like the two sample versions of other EDF statistics, it is essentially a rank test. Pettitt adds some asymptotic power comparisons with other two sample rank tests, which show that A^2 compares very favorably.

REFERENCES

- Anderson, T. W. and Darling, D. A. (1952). Asymptotic theory of certain 'goodness of fit' criteria based on stochastic processes. Ann. Math. Statist. 23, 193-212.
- Anderson, T. W. and Darling D. A. (1954). A test of goodness of fit. J. Am. Statist. Assoc. 49, 765-769.
- Durbin, J. (1973a). Weak convergence of the sample distribution functions when parameters are estimated. Ann. Statist. 1, 279-290.
- Durbin, J. (1973b). Distribution theory for tests based on the sample distribution function. Philadelphia; SIAM.
- Durbin, J., Knott, M. and Taylor, C. C. (1975). Components of Cramer-von Mises statistics. II. J. R. Statist. Soc. B, 37, 216-237.
- Lewis, Peter A.W. (1961). Distribution of the Anderson-Darling statistic. Ann. Math. Statist. 32, 1118-1124.
- Locke, C. and Spurrier, J. D. (1978). On tests for uniformity. Commun. Statist.-Theor. Meth. A 7(3), 241-258.
- Pettitt, A. N. (1976). A two sample Anderson-Darling rank statistic. Biometrika, 63, 161-168.
- Pettitt, A. N. (1977). Testing the normality of several independent samples using the Anderson-Darling statistic. J. Royal Statist. Soc. A, 156-161.
- Pettitt, A. N. and Stephens, M. A. (1976). EDF statistics for testing for the Gamma distribution with applications to testing for equal variances. Unpublished.

- Quesenberry, C. P. and Miller, F. L. (1977). Power studies of some tests for uniformity. J. Statist. Comput. Simul. 5, 109-191.
- Shapiro, S. S. and Wilk, M. B. (1965). An analysis-of-variance test for normality (Complete Samples). Biometrika, 52, 591-611.
- Shapiro, S. S. and Wilk, M. B. (1972). An analysis of variance test for the exponential distribution. Technometrics, 14, 355-370.
- Stephens, M. A. (1974). EDF statistics for goodness-of-fit and some comparisons. J. Am. Statist. Assoc. 69, 730-737.
- Stephens, M. A. (1976). Asymptotic results for goodness-of-fit statistics with unknown parameters. Ann. Statist. 4, 357-369.
- Stephens, M. A. (1977). Goodness-of-fit for the extreme value distribution. Biometrika, 64, 583-588.
- Stephens, M. A. (1979). EDF tests of fit for the logistic distribution. To appear Biometrika, December 1979.

TABLES 1, 2 and 3

Modified forms A^* and upper tail percentage points for tests as follows.

Table 1: Test for any completely specified distribution (Case 0), Section 3.

Table 2: Tests for normality, Section 6.1.

Table 3: Tests for exponentiality, Scale parameter unknown, Section 6.2.

Table No.	Modified A^*	Upper tail percentage level α									
		.25	.20	.15	.10	.05	.025	.01	.005	.0025	
1.	For all $n \geq 5$			1.610	1.933	2.492	3.070	3.853			
2. Case 1	See below	.644		.782	.894	1.087	1.285	1.551	1.756	1.964	
Case 2	See below	1.072		1.430	1.743	2.308	2.898	3.702	4.324	4.954	
Case 3	$A^* = A^2(1.0+0.75/n+2.25/n^2)$	.472	.509	.561	.631	.752	.873	1.035	1.159	1.283	
3.	$A^* = A^2(1.0+0.3/n)$	.736	.816	.916	1.062	1.321	1.591	1.959	2.244	2.534	

Note. For Table 2 Cases 1 and 2, normal distribution tested, no modifications have been calculated. The percentage points given can be used with unmodified A^2 for $n > 20$.

TABLES 4 and 5

Modified forms A^* and upper tail percentage points for tests as follows.

Table 4: Tests for the extreme value distribution, Section 6.3.

Table 5: Tests for the logistic distribution, Section 6.5.

Table No.		Modified A^*	Upper tail percentage level α							
			.25	.20	.15	.10	.05	.025	.01	.005
4.	Case 1	$A^* = A^2(1.0+0.3/n)$	.736	.816	.916	1.002	1.321	1.591	1.959	2.244
	Case 2	None	1.060			1.725	2.277	2.854	3.640	2.534
	Case 3	$A^* = A^2(1.0+0.2/\sqrt{n})$	0.474			0.637	0.757	0.877	1.038	
5.	Case 1	$A^* = A^2(1.0+0.15/n)$	.615			.857	1.046	1.241	1.505	1.710
	Case 2	$A^* = (0.6nA^2 - 1.8)/(0.6n - 1.0)$	1.043			1.725	2.290	2.880	3.685	4.308
	Case 3	$A^* = A^2(1.0+0.25/n)$	.426			.563	.660	.769	.906	1.010

TABLE 6

Modified form A^* and upper tail percentage points for a test for the gamma distribution with known shape parameter m and unknown scale parameter, Section 6.6.

Note. For $m = 1$ (test for exponentiality) see Table 3.

Modified A^*	m	Upper tail percentage level α			
		.10	.05	.025	.01
For all $m \geq 2$: $A^* = A^2 + \frac{1}{n}(0.2 + \frac{0.3}{m})$	2	.989	1.213	1.441	1.751
	3	.959	1.172	1.389	1.683
	4	.944	1.151	1.362	1.648
	5	.935	1.139	1.346	1.627
	6	.928	1.130	1.335	1.612
	8	.919	1.120	1.322	1.595
	10	.915	1.113	1.314	1.583
	12	.911	1.110	1.310	1.578
	15	.908	1.106	1.304	1.570
	20	.905	1.101	1.298	1.502
	∞	.893	1.089	1.281	1.551

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER 39	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) THE ANDERSON-DARLING STATISTIC		5. TYPE OF REPORT & PERIOD COVERED TECHNICAL REPORT
		6. PERFORMING ORG. REPORT NUMBER
7. AUTHOR(s) MICHAEL A. STEPHENS		8. CONTRACT OR GRANT NUMBER(s) DAAG29-77-G-0031
9. PERFORMING ORGANIZATION NAME AND ADDRESS Department of Statistics Stanford University Stanford, CA 94305		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS P-14435-M
11. CONTROLLING OFFICE NAME AND ADDRESS U. S. Army Research Office Post Office Box 12211 Research Triangle Park, NC 27709		12. REPORT DATE OCTOBER 31, 1979
		13. NUMBER OF PAGES 18
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		15. SECURITY CLASS. (of this report) UNCLASSIFIED
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) APPROVED FOR PUBLIC RELEASE: DISTRIBUTION UNLIMITED.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES The findings in this report are not to be construed as an official Department of the Army position, unless so designated by other authorized documents. This report partially supported under Office of Naval Research Contract N00014-76-C-0475 (NR-042-267) and issued as Technical Report No. 278.		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Anderson-Darling Statistic, Tests of Fit, Goodness-of-Fit.		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) PLEASE SEE REVERSE SIDE		

DD FORM 1 JAN 73 1473

EDITION OF 1 NOV 53 IS OBSOLETE

S/N 0102-LF-014-6601

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

THE ANDERSON-DARLING STATISTIC

BY

Michael A. Stephens

ABSTRACT

The Anderson-Darling statistic A^2 is a goodness-of-fit statistic, based on the empirical distribution function. Its asymptotic distribution can be found for testing many important distributions when unknown parameters must be estimated from the data. Furthermore, A^2 can be easily adapted so that only the asymptotic points are needed for testing purposes. A^2 also is easy to calculate, and has overall good power properties. The report gives a review of A^2 and tables for testing the following distributions - normal, exponential, gamma, extreme-value and Weibull, and logistic; points are given also for testing any completely specified continuous distribution.

#278/#39