



Linear flaw detection in woven textiles using model-based clustering

J.G. Campbell ^{a,1}, C. Fraley ^{b,c,2}, F. Murtagh ^{a,*}, A.E. Raftery ^{b,3}

^a Faculty of Informatics, University of Ulster, Magee College, Londonderry BT48 7JL, Northern Ireland, United Kingdom

^b Department of Statistics, Box 354322, University of Washington, Seattle, WA 98195-4322, USA

^c MathSoft Inc., 1700 Westlake Avenue N., Suite 500, Seattle, WA 98109, USA

Received 9 September 1996; revised 7 August 1997

Abstract

We combine image-processing techniques with a powerful new statistical technique to detect linear pattern production faults in woven textiles. Our approach detects a linear pattern in preprocessed images via model-based clustering. It employs an approximate Bayes factor which provides a criterion for assessing the evidence for the presence of a defect. The model used in experimentation is a (possibly highly elliptical) Gaussian cloud superimposed on Poisson clutter. Results are shown for some representative examples, and contrasted with a Hough transform. Software for the statistical modeling is available. © 1997 Elsevier Science B.V.

Keywords: Model-based clustering; Pattern recognition; Bayesian cluster analysis; Machine vision; Industrial inspection; Hough transform

1. The flaw detection problem

Garment production can be divided into two distinct phases: manufacture of the textile fabric, followed by garment assembly. The two phases are often performed in different locations and by different organizations. Each phase in turn is made up of sub-phases, between which there are opportunities for inspection. Our interest is in the problem of product inspection after fabric manufacturing, before final assembly.

Typically fabric is produced by looms in two-meter wide rolls at a rate of about 10 mm per second. Although it might seem that product inspection could occur concurrently, the fabric is first packed into rolls and later unrolled for inspection. Reasons for this presumably include the slow speed of production, which is insufficient to keep an inspector occupied, and the relatively hostile working environment. This work is concerned with the replacement of manual inspection by an automatic procedure (Newman and Jain, 1995).

Two major obstacles to machine inspection of textile fabrics are the difficulty of characterizing defects, and the high data rate. The denim fabric considered here manifests the former problem in abundance: there are many defect types, some of

* Corresponding author. E-mail: fd.murtagh@ulst.ac.uk.

¹ E-mail: jg.campbell@ulst.ac.uk.

² E-mail: fraley@stat.washington.edu.

³ E-mail: raftery@stat.washington.edu.

which are quite subtle, due to the local texture irregularity that is one of its attractive features.

In the manual inspection process, the flaws are marked using chalk or metallic tape. At garment assembly, cutting into shapes is done on batches of approximately fifty layers. This layering is manually supervised, and the operators attempt to handle flawed regions via cutting and excising, or overlapping. In the context of automated manufacturing, there is clearly significant scope for introducing intelligence to these phases: if location of automatically detected flaws can be supplied to an automatic cutter, then an optimal cutting plan may be followed, i.e. flaws avoided with minimal wastage.

Obvious flaws, such as torn threads, can be captured by sizable deviations from the fixed background pattern. In this article, we assess a new and

powerful statistical methodology for less obvious flaws – those which present a subtle local pattern but which, due to their occurrence in an extended spatial pattern, are easily picked out by the human eye.

Fig. 1 shows a sample of faults. The torn thread “splurges” can be detected through thresholding and size of the contiguous area. The more difficult case of faint aligned flaws will be investigated in this article. Previous work on this data has included Campbell et al. (1995) which used discrete Fourier transform texture descriptors (the amplitude spectrum) in 32×32 subimage windows to provide input to a trainable classifier system. In this work we study one particular type of flaw only – a noticeable highly aligned pattern, associated with torn fabric or thread. In focusing on one type of flaw, we scan a larger area of the image – pixel dimensions of

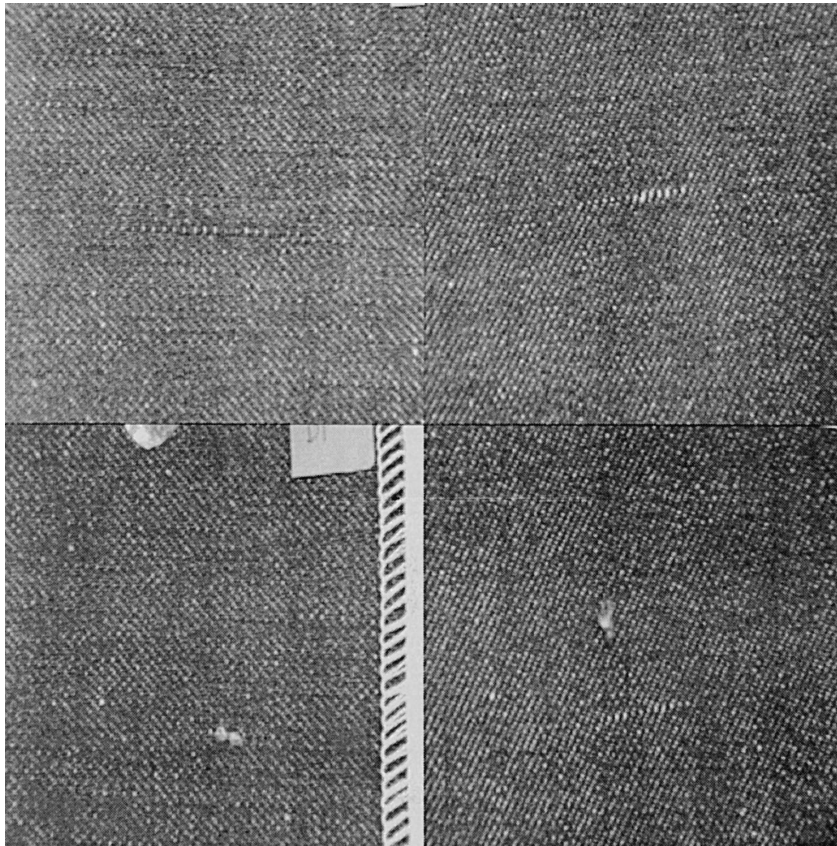


Fig. 1. Four different textile samples, rebinned to half their input dimensions and placed in the four quadrants seen here. Illustrated are unclear torn threads, label and edging corrupts in the bottom two images; and linear flaws in the top two samples.

around 500×500 are used in experiments. For this reason, this work aims at being both practical and task-specific.

The next section reviews the basis for the statistical cluster-finding and testing method. It assumes that a point pattern cluster is to be found in a background noise field. This section, Section 2, presents salient results from previous work in this area. Another viewpoint on these results may be found by perusing the software which implements this (details of availability in Section 3).

The preliminary image-processing steps are treated in a section on experimentation (Section 3). This involves thresholding and cleaning using mathematical morphology, followed by labeling and analysis of contiguous areas in order to provide point pattern data.

2. Model-based clustering

In our experimentation below, we will model the data as a (highly elliptical) Gaussian, subject to Poisson background clutter. The overall point pattern will be derived by thresholding and by morphological operators from the input image data. The data are 2-dimensional. To begin with, we discuss the modeling in the general context of distribution mixtures.

Consider data which are generated by a mixture of $(G - 1)$ bivariate Gaussian densities, $f_k(x; \theta) \sim \mathcal{N}(\mu_k, \Sigma_k)$, for clusters $k = 2, \dots, G$, and with Poisson background noise corresponding to $k = 1$. The overall population thus has the mixture density

$$f(x; \theta) = \sum_{k=1}^G \pi_k f_k(x; \theta),$$

where the mixing or prior probabilities, π_k , sum to 1, and $f_1(x; \theta) = \mathcal{A}^{-1}$, where $\mathcal{A}$ is the area of the data region. This is the basis for *model-based clustering* (Banfield and Raftery, 1993; Dasgupta and Raftery, 1995; Murtagh and Raftery, 1984; Banerjee and Rosenfeld, 1993).

The parameters, θ and π , can be estimated efficiently by maximizing the likelihood, sometimes also called the *mixture likelihood*, namely

$$L(\theta, \pi) = \prod_{i=1}^n f(x_i; \theta),$$

with respect to θ and π , where x_i is the i th observation.

In this work, we assume the presence of two clusters, one of which is Poisson noise, the other Gaussian. This yields the mixture likelihood

$$L(\theta, \pi) = \prod_{i=1}^n \left[\pi_1 \mathcal{A}^{-1} + \pi_2 \frac{1}{2\pi\sqrt{|\Sigma|}} \times \exp \left\{ -\frac{1}{2} (x_i - \mu)^T \Sigma^{-1} (x_i - \mu) \right\} \right],$$

where $\pi_1 + \pi_2 = 1$.

An iterative solution is provided by the expectation-maximization (EM) algorithm of Dempster et al. (1977). Let the “complete” (or “clean” or “output”) data be $y_i = (x_i, z_i)$ with indicator set $z_i = (z_{i1}, z_{i2})$ given by (1,0) or (0,1). Vector z_i has a multinomial distribution with parameters $(1; \pi_1, \pi_2)$. This leads to the *complete data log-likelihood*:

$$l(y, z; \theta, \pi) = \sum_{i=1}^n \sum_{k=1}^2 z_{ik} [\log \pi_k + \log f_k(x_k; \theta)].$$

The E-step then computes $\hat{z}_{ik} = E(z_{ik} | x_1, \dots, x_n, \theta)$, i.e. the posterior probability that x_i is in cluster k . The M-step involves maximization of the *expected complete data log-likelihood*:

$$l^*(y; \theta, \pi) = \sum_{i=1}^n \sum_{k=1}^2 \hat{z}_{ik} [\log \pi_k + \log f_k(x_i; \theta)].$$

The E- and M-steps are iterated until convergence.

For the 2-class case (Poisson noise and a Gaussian cluster), the complete-data likelihood is

$$L(y, z; \theta, \pi) = \prod_{i=1}^n \left[\frac{\pi_1}{\mathcal{A}} \right]^{z_{i1}} \left[\frac{\pi_2}{2\pi\sqrt{|\Sigma|}} \times \exp \left\{ -\frac{1}{2} (x_i - \mu)^T \Sigma^{-1} (x_i - \mu) \right\} \right]^{z_{i2}}.$$

The corresponding expected log-likelihood is then used in the EM algorithm. This formulation of the problem generalizes to the case of G clusters, of arbitrary distributions and dimensions.

In order to assess the evidence for the presence of a defect, we use the *Bayes factor* for the mixture model, M_2 , that includes a Gaussian density as well as background noise, against the “null” model, M_1 , that contains only background noise. The Bayes factor is the posterior odds for the mixture model against the pure noise model, when neither is favored a priori. It is defined as $B = p(x|M_2)/p(x|M_1)$, where $p(x|M_2)$ is the *integrated likelihood* of the mixture model M_2 , obtained by integrating over the parameter space. For a general review of Bayes factors, their use in applied statistics, and how to approximate and compute them, see (Kass and Raftery, 1995).

We approximate the Bayes factor using the *Bayesian Information Criterion* (BIC) (Schwarz, 1978). In the present context, this takes the form

$$2 \log B \approx \text{BIC}$$

$$= 2 \log L(\hat{\theta}, \hat{\pi}) + 2n \log \mathcal{N} - 6 \log n,$$

where $\hat{\theta}$ and $\hat{\pi}$ are the maximum likelihood estimators of θ and π , and $L(\hat{\theta}, \hat{\pi})$ is the maximized mixture likelihood. Any value of BIC greater than zero corresponds to evidence for a defect. Conventionally, BIC values between 0 and 2 correspond to weak evidence, values between 2 and 6 correspond to positive evidence, values between 6 and 10 correspond to strong evidence, and values greater than 10 correspond to very strong evidence (Kass and Raftery, 1995). The BIC criterion is prone to false positives but compared to other testing criteria performs very well (Titterton et al., 1985; Leroux, 1992), and this is backed up by experimental results.

The method described so far does not incorporate any explicit mechanism for linearity- or alignment-seeking. When there is only one flaw in the image, corresponding to a single Gaussian cluster, this does not seem to matter. The unconstrained Gaussian density tends to adapt to what is in the image, finding a feature that is highly linear (i.e. long and thin) if it is present. However, if there are several flaws, perhaps intersecting one another, a more explicit incorporation of linearity might be advantageous. We now indicate briefly how this can be done.

In model-based clustering, the covariance matrix Σ associated with a cluster is parametrized (Banfield

and Raftery, 1993) as $\Sigma = \lambda D A D^T$, where λ is the largest eigenvalue of Σ , D is the matrix of eigenvectors, and $A = \text{diag}\{1, \alpha\}$. Each of the three components of this decomposition of the covariance matrix corresponds to a geometric and visually intuitive property of the cluster that it describes. Thus, λ corresponds to the *volume* of the cluster, D to its *orientation*, and A (or equivalently here, α) to its *shape*. The value α is the ratio of second to first eigenvalues. For α close to 1, clusters will be spherical; while for values approaching 0, the clusters will be very linear (i.e. their members will be highly aligned).

The user (or program, e.g. using Bayes factors as described below), can set or determine values of λ to

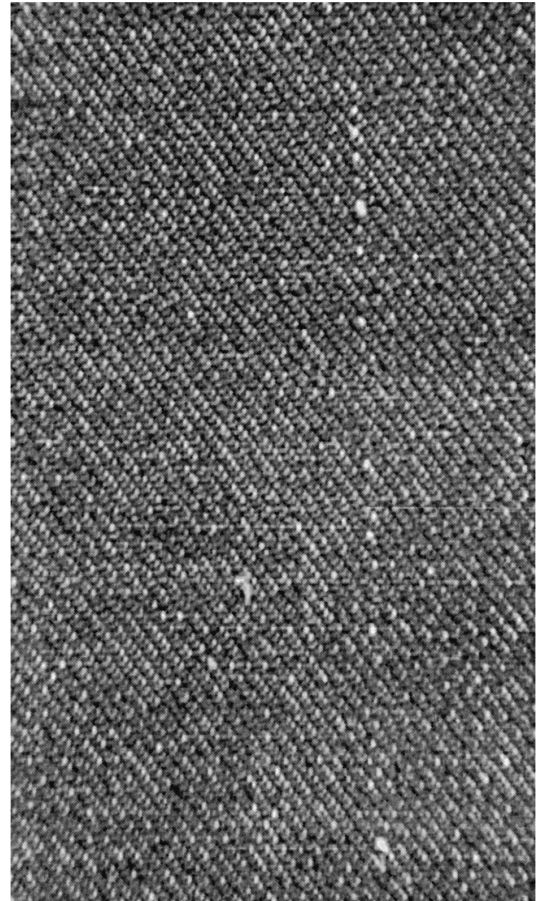


Fig. 2. Image used for experimentation referred to as d8.

control the cluster volume, D to control orientation, and A to control shape. By constraining some or all of λ , D and A to be equal across clusters, the finding of clusters of different types can be prioritized. In this work, we are interested only in letting the data determine the best value for A which amounts to determining the best value for α . Murtagh and Raftery (1984) assumed user-specification of α . A maximum likelihood estimate of the clusters (using EM) may additionally be used automatically to determine an optimal value of α in the following way.

Take a set of n points comprising a Gaussian cluster (n_1 points), with spatially homogeneous Poisson background (n_0 points), and let the sample covariance matrix for the cluster have spectral, or singular value, decomposition $\hat{\Sigma} = L\Omega L^T$. The maxi-

mized classification log-likelihood of the data, with α assumed known, is

$$\begin{aligned} 2l = & -(n - n_0)(2\log(2\pi) + 2(1 - \log 2)) \\ & + \log(|A|) - 2n_1\log(\text{tr}(\Omega_k A^{-1})/n_1) \\ & - 2n_0\log(\mathcal{A}). \end{aligned}$$

The “profile likelihood” with respect to α is then maximized. This results in a likelihood equation which reduces to the following simple expression for the estimate of α : $\hat{\alpha} = \omega_2/\omega_1$ i.e. the ration of eigenvalues (Dasgupta and Raftery, 1995). This reinforces the approach of Murtagh and Raftery (1984) by casting this problem in a likelihood framework.

To summarize, we seek a highly elliptical Gaussian cluster superimposed on a homogeneous Poisson

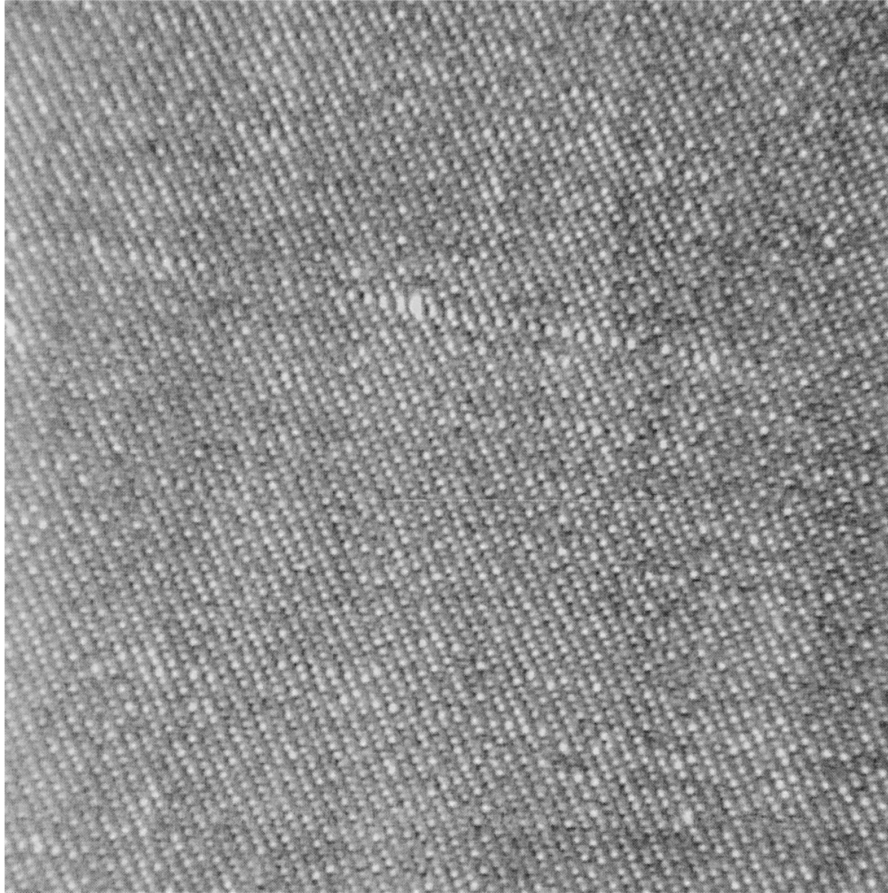


Fig. 3. Image used for experimentation referred to as d10.

background. Furthermore we use the BIC quality criterion for the fit of this 2-cluster mixture model to the data.

3. Sample processing of two images

Figs. 2 and 3 show the images used. The image shown in Fig. 2 was cropped of edging (to avoid undesired effects on thresholding and other operations) and, as shown, is of dimensions 415×501 . The image shown in Fig. 3 is of dimensions 512×512 . A simple thresholding using a 3-sigma detection limit (i.e. image mean value + 3 times the image standard deviation) was applied. A large number of thresholded pixel values remained. An opening (ero-

sion followed by dilation) was applied, with 3×3 structuring element, SE, $((0,1,0),(1,1,1),(0,1,0))$, i.e. a cross shape. A 3×3 SE of one-values had worked particularly well on Fig. 2, whereas a 2×2 SE of one-values had worked particularly well on Fig. 3, so the cross-shaped SE was chosen to cater for both cases. Fig. 4 shows the result of thresholding and applying the opening to Fig. 3. The contiguous thresholded regions were then labeled, and their centroids obtained. In this way a point pattern set was derived from these images.

To counteract difficulties in dealing with many unweighted points (for example, size-related weights were not investigated), we excluded from consideration all centroids associated with the (numerous) smallest contiguous regions. A lower limit of five

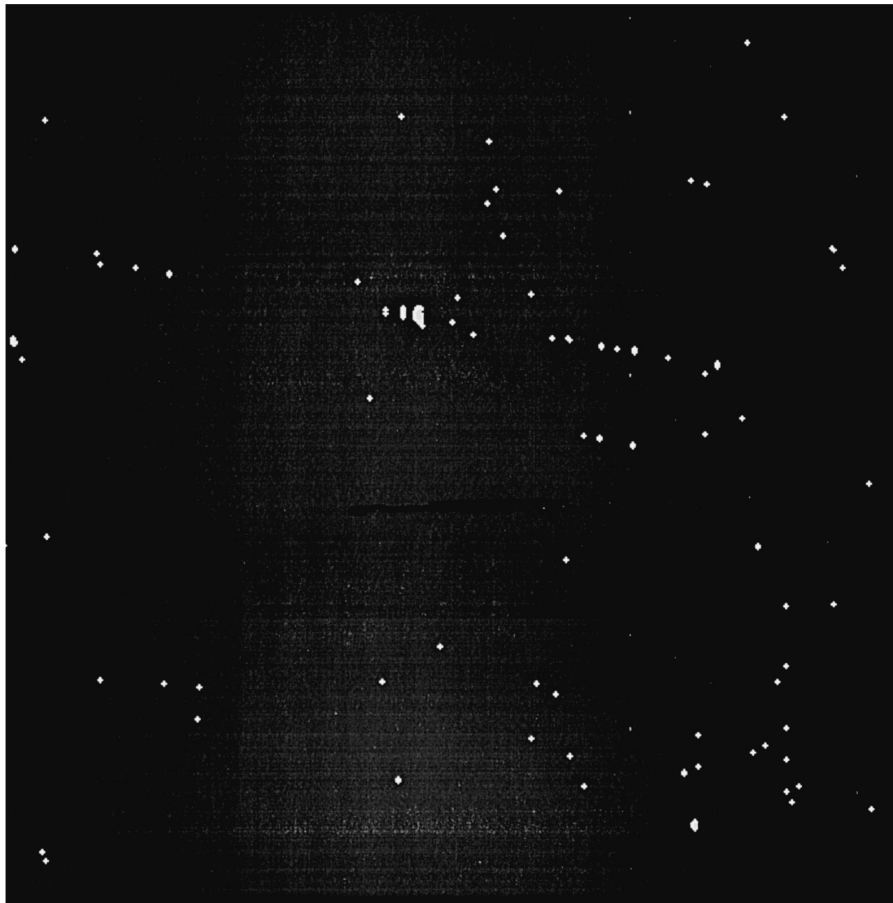


Fig. 4. Image d10 following thresholding and a morphological opening.

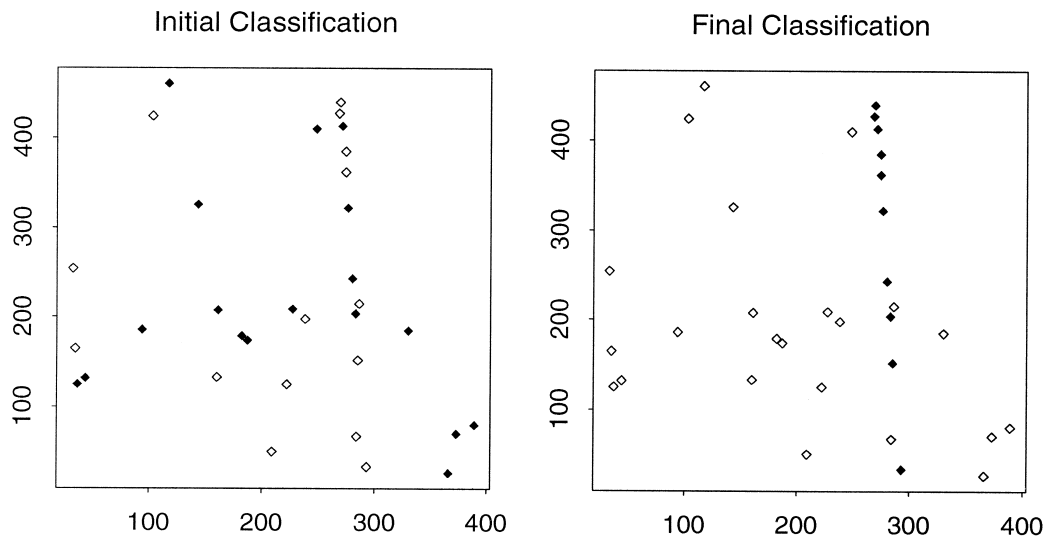


Fig. 5. Analysis of image d8. The point detections derived from the image shown in Fig. 2 are shown in pixel coordinates. The initial classification consists of random assignments. The final classification is an extremely elongated elliptical cluster (black) with a Poisson noise background (white).

pixels was imposed, as well as a rejection rule for labeled regions too close to the image boundary. Figs. 5 and 6 show the point sets used with the initial (random) and final configurations; dark points indicate those alleged to belong to the cluster. In each

case, the final cluster assumes an elliptical shape arising from the Gaussian model. A satisfactory solution was obtained (Figs. 5 and 6) with corresponding BIC values of 19.67 and 41.33, respectively. Thus, in each case, BIC correctly indicated strong evidence

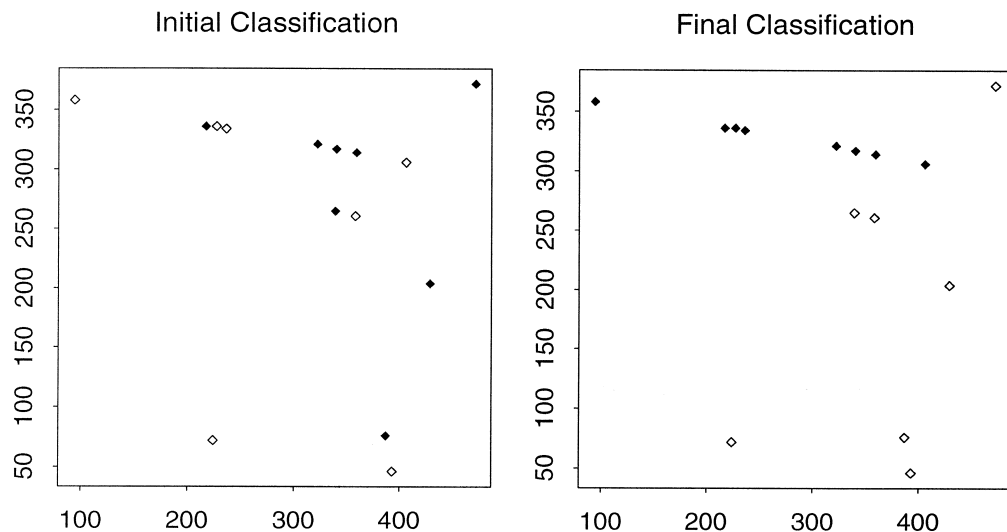


Fig. 6. Analysis of image d10. The point detections derived from the image shown in Fig. 3 are shown in pixel coordinates. As for the previous figure, the initial classification consists of random assignments; and the final classification is a fit of an elongated ellipse (black) with background noise (white).

for the presence of a defect (since $BIC > 10$). Fig. 5 also show that the model-based clustering method correctly identified where the defects were.

The clustering method can produce different results for different starting values (although the results are highly consistent: more than 80% of the time, in the case of examples discussed in this article). It is also the case that the ability of the method to detect flaws declines if the criterion for deriving the point pattern from the image is not stringent enough (e.g. if a lower limit of less than five pixels is used in images d8 and d10). However the BIC value was invariably reliable as an indicator of the quality of the solution. Good results corresponded to large BIC values (of the order of 10–40 for these examples), while bad ones yielded BIC

values that were either negative or fairly small in magnitude. The BIC criterion was also very reliable for treating data without any apparent aligned set of points: point sets consisting of Poisson-distributed data gave rise to small positive values of BIC for a few such simulations, but overwhelmingly these simulations gave rise to negative values. We see therefore that the BIC criterion provides a “safety net” in regard to starting configurations and in regard to the no cluster or no fault situation.

The initial image processing (thresholding, opening, labeling, object centroiding) was carried out in IDL. The analysis of the point patterns was carried out in S-Plus. The code used for the latter is available in the S archive at the Statlib server, <http://lib.stat.cmu.edu/S/mclustem.one> and also at

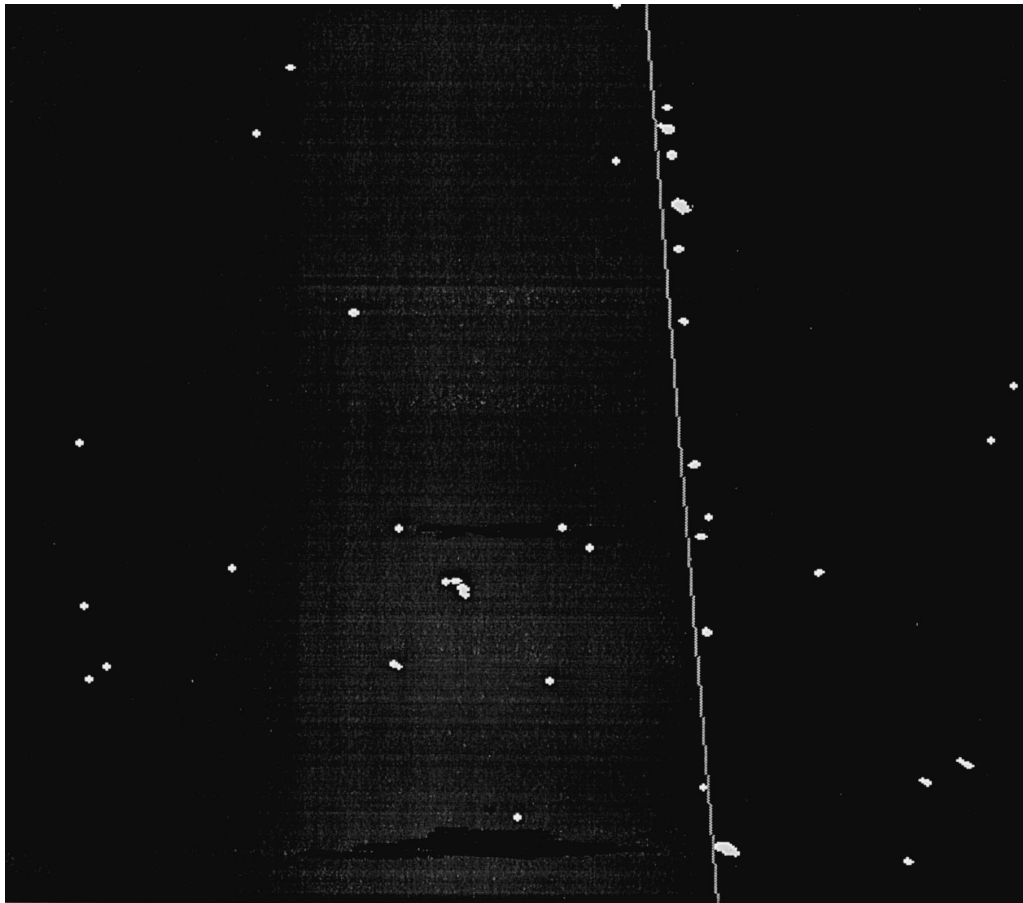


Fig. 7. A Hough transform result for image d8 (see Fig. 2), using threshold 160 on the accumulator array. This method for finding aligned points may be contrasted with the model-based clustering result shown in Fig. 5.

<http://www.stat.washington.edu/fraley/software.html>. From Statlib, the code can also be obtained by email, by sending a message of the form “send mclustem.one from S” to statlib@lib.stat.cmu.edu. For point patterns with numbers similar to those shown in the examples, the computational time required is insignificant.

The Hough transform is the traditional method for alignment-detection (Duda and Hart, 1973) and is also fast. Figs. 7 and 8 show two results for the image d8. Note that pixels were used for the Hough transform, so therefore the size of potential faulty regions of the fabric were taken into account. This is

advantageous (cf. discussion above on the generalization of the modeling method to account for weights). As opposed to this, the thresholding of the accumulator array is highly sensitive and Figs. 7 and 8 illustrate this. As seen in the latter, fitting multiple lines is possible. This is not necessarily advantageous in the case of the current problem since, as mentioned earlier, the thread-based faults are likely to be unique in a given area of fabric. With further processing of the Hough transform results (perhaps even with incorporation of a BIC criterion!), this method could provide very similar results to those provided by the modeling approach.

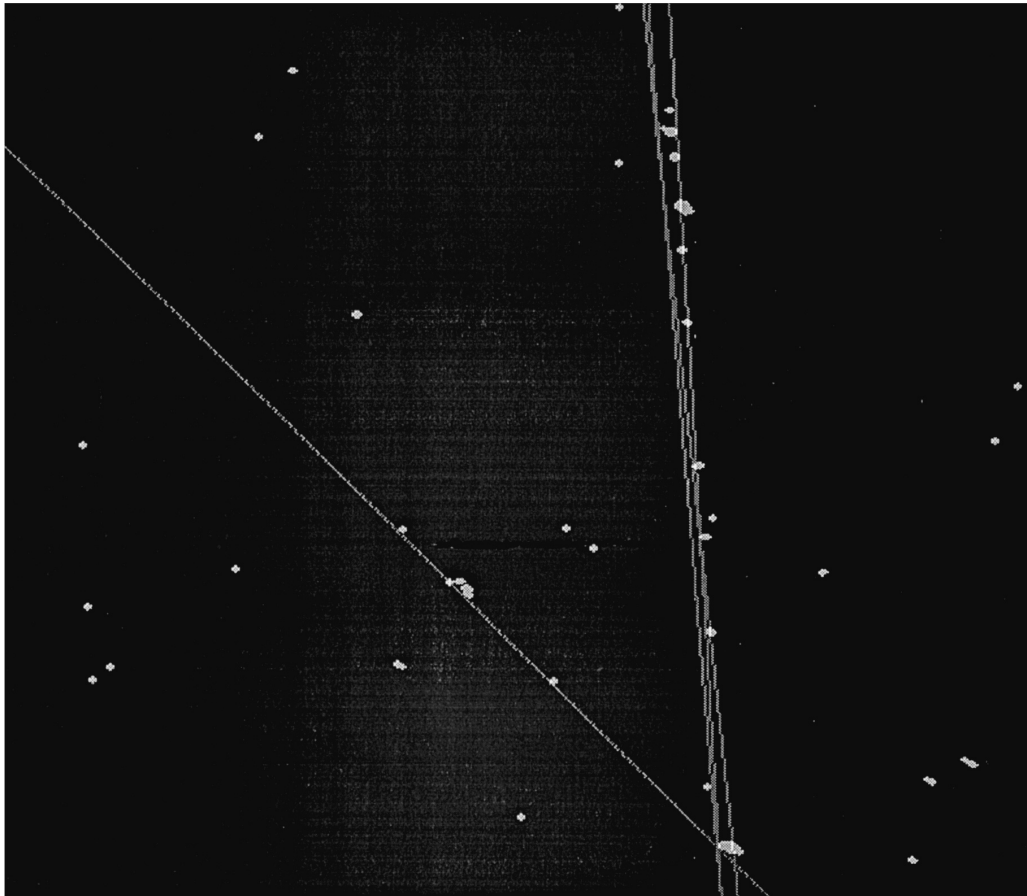


Fig. 8. A Hough transform result for image d8 (see Fig. 2), using threshold 100 on the accumulator array. This less stringent but typical threshold gives rise to the multiple lines. Compared to the model-based clustering result shown in Fig. 6, a conservative threshold selection such as shown here is susceptible to false alarms.

4. Conclusion

The question of whether or not a cluster is present in a data set is a recurrent one with a long history. The model-based methodology described here, with its attendant inferential detection criterion, is one which has worked well on other data. The model assumes a (possibly very linear, i.e. long and thin) Gaussian distribution for the cluster, and a uniform distribution for noise, assumed to arise from a spatial Poisson process. Here we apply this approach to the difficult problem of detecting relatively faint aligned faults in denim textiles. We have also devoted considerable attention to the processing chain which extends from the capture of the images. Finally, it should be noted that the operations carried out here are very fast, of the order of a second on Sparcstation 20 class platforms.

This work has focused on the aligned point pattern detection issue. A comprehensive fault detection system would need also to detect large discolored or differently textured areas, to distinguish heavily textured boundary areas, and various other special cases. The problem addressed here is quite an important one. We see its use as a component of a fault detection system, driven by a knowledge-based control module.

Acknowledgements

This work was supported by Office of Naval Research Grants N-00014-96-1-0192 and N-00014-

96-1-0330. We are grateful for the comments of three referees on an earlier version of this paper.

References

- Banerjee, S., Rosenfeld, A., 1993. Model-based cluster analysis. *Pattern Recognition* 26, 963–974.
- Banfield, J.D., Raftery, A.E., 1993. Model-based Gaussian and non-Gaussian clustering. *Biometrics* 49, 803–821.
- Campbell, J.G., Hashim, A.A., McGinnity, T.M., Lunney, T.F., 1995. Flaw detection in woven textiles by neural network. In: Keating, J.G. (Ed.), *Neural Computing: Research and Applications III*, Proc. 5th Irish Neural Network Conference, St. Patrick's College, Maynooth.
- Dasgupta, A., Raftery, A.E., 1995. Detecting features in spatial point processes with clutter via model-based clustering. Technical Report 295, Statistics Department, University of Washington (available at <http://www.stat.washington.edu/tech.reports/tr295.ps>).
- Dempster, A.P., Laird, N.M., Rubin, D.B., 1977. Maximum likelihood from incomplete data via the EM algorithm. *J. Roy. Statist. Soc. Ser. B* 39, 1–22.
- Duda, R.O., Hart, P.E., 1973. *Pattern Classification and Scene Analysis*. Wiley, New York.
- Kass, R.E., Raftery, A.E., 1995. Bayes factors. *J. Amer. Statist. Assoc.* 90, 773–795.
- Leroux, B.G., 1992. Consistent estimation of a mixing distribution. *Ann. of Statist.* 20, 1350–1360.
- Murtagh, F., Raftery, A.E., 1984. Fitting straight lines to point patterns. *Pattern Recognition* 17, 479–483.
- Newman, T.S., Jain, A.K., 1995. A survey of automated visual inspection. *Comput. Vision Image Understanding* 61, 231–262.
- Schwarz, G., 1978. Estimating the dimension of a model. *Ann. of Statist.* 6, 461–464.
- Titterton, D.M., Smith, A.F.M., Makov, U.E., 1985. *Statistical Analysis of Finite Mixture Distributions*. Wiley, New York.