



The evaluation of evidence relating to traces of cocaine on banknotes



Amy Wilson^a, Colin Aitken^{a,*}, Richard Sleeman^b, James Carter^c

^a School of Mathematics and Maxwell Institute, University of Edinburgh, Mayfield Road, Edinburgh EH9 3JZ, United Kingdom

^b Mass Spec Analytical Ltd., Building 20F, Golf Course Lane, Bristol BS34 7RP, United Kingdom

^c Clinical and Statewide Services Division, Queensland Health, PO Box 594, Archerfield, Queensland 4108, Australia

ARTICLE INFO

Article history:

Received 19 February 2013

Received in revised form 5 November 2013

Accepted 27 November 2013

Available online 17 December 2013

Keywords:

Autoregressive model

Banknotes

Cocaine

Evidence evaluation

Likelihood ratio

Nonparametric density estimation

ABSTRACT

Banknotes can be seized from crime scenes as evidence for suspected association with illicit drug dealing. Tandem mass spectrometry data are available from banknotes seized in criminal investigations, as well as from banknotes from general circulation. The aim of the research is to evaluate the support provided by the data gathered in a criminal investigation for the proposition that the banknotes from which the data were obtained are associated with a person who is associated with a criminal activity related to cocaine in contrast to the proposition that the banknotes are associated with a person who is not associated with a criminal activity involving cocaine. The data considered are the peak area for the ion count for cocaine product ion m/z 105. Previous methods for assessment of the relative support for these propositions were concerned with the percentage of banknotes contaminated or assume independence of measurements of quantities between adjacent banknotes. Methods which account for an association of the quantity of drug on a banknote with that on adjacent banknotes are described. The methods are based on an autoregressive model of order one and on two versions of a nonparametric approach. The results are compared with a standard model which assumes measurements on individual banknotes are independent; there is no autocorrelation. Performance is assessed using rates of misleading evidence and a recommendation made as to which method to use.

© 2014 Published by Elsevier Ireland Ltd.

1. Introduction

A novel approach to the evaluation of evidence of the quantity of cocaine on banknotes is described. The methods are applicable more generally and can be applied to measurements of quantities of other drugs. The novelty is the consideration of autocorrelation which, in this context, is a measure of the association between the quantities of drugs on adjacent banknotes. A previous method, described in [1], uses only one banknote and states that one cannot assume independence in the measurements between adjacent banknotes.

The evidence is evaluated through use of the likelihood ratio (LR). In this context, the ratio is that of the probability density function of the data¹ under each of two propositions,

- H_C : the banknotes are associated with a person who is associated with a criminal activity involving cocaine, and
- H_B : the banknotes are associated with a person who is not associated with a criminal activity involving cocaine.

There has been some previous use of LR in the area of drugs on banknotes. In [1], the likelihood ratio of the quantity of contamination of cocaine on a seized banknote was evaluated using a histogram. It is noted that calculating a LR for a set of multiple banknotes using this method is not possible without assuming independence (i.e., that the quantity of cocaine on a particular banknote is unaffected by the quantity of cocaine on any other banknote, such as an immediate neighbour). In [2] the likelihood ratio for the quantity of cocaine contamination on a set of banknotes is calculated using a univariate kernel density estimate. An assumption of independence is made and it is noted that this assumption may not be warranted. This assumption is not made in the models introduced in this paper. The results obtained from these models are compared with a model which assumes independence.

The data used for the analysis, $\mathbf{z} = (z_1, z_2, \dots, z_n)$, are the logarithms of the peak areas of cocaine on a set of n banknotes. The strength of the evidence of $\mathbf{z}$ in support of H_C or H_B is to be assessed. The logarithmic transformation is made to reduce skewness in the

* Corresponding author. Tel. +44 0131 650 4877

E-mail address: c.g.aitken@ed.ac.uk (C. Aitken).

¹ The ratio is known as a likelihood ratio as this is a technical phrase in statistical theory in which a probability density function for data given parameter values may be thought of as a likelihood of the parameter values given the data. The phrase has been transferred in the forensic statistic literature to refer to the ratio of the probability density functions given propositions. Note that as the data are continuous it is not possible to refer to the probability of the data.

data before fitting models. Training data are available from banknotes deemed to be associated with criminal activity involving cocaine and from banknotes deemed to be associated with general or background circulation. These training data are used to develop models associated with H_C and H_B , respectively.

The likelihood ratio LR associated with the propositions H_C and H_B is given by:

$$LR = \frac{f(\mathbf{z}|H_C)}{f(\mathbf{z}|H_B)},$$

where the function f is a probability density function for the measurements, conditional on H_C and on H_B in the numerator and the denominator, respectively. If this statistic is greater than one, then the evidence assigns more support to the proposition that the banknotes are associated with a person who is associated with drug crime involving cocaine. With an assumption of independence amongst the values in $\mathbf{z}$,

$$LR = \frac{\prod_{i=1}^n f(z_i|H_C)}{\prod_{i=1}^n f(z_i|H_B)}.$$

The interpretation of this LR is slightly different from the one used for comparison of possible sources of recovered and control evidence in [3]. Here there is only one set of evidence, the seized banknotes provided by the law enforcement agency. The LR provides a measure of support for one or other of the propositions as to whether the person with whom they are associated is himself associated or not with criminal activity involving cocaine.

The models described here give three approaches to the estimation of f in which there is no assumption of independence amongst the quantities on the n banknotes in the sample of unknown origin. One approach allows for an autocorrelation of lag one which models a dependency between measurements on adjacent banknotes. The other two approaches use a kernel density approach with multivariate conditional density functions instead of the univariate kernel density functions used in [2], again modelling a dependency between measurements from adjacent banknotes; one of these approaches has a fixed bandwidth, the other a bandwidth which varies with the density of the measurements.

2. Data

The method for acquisition of the data is described in [4]. Models were developed for the cocaine product ion m/z 105. For each sample of banknotes the data resemble a series of peaks, with each peak corresponding to a banknote. The height of the peak at any given scan number is given by the number of gas phase ion transitions giving rise to the cocaine product ion m/z 105 at that scan number. A peak detection algorithm was written in order to identify these peaks. The details are beyond the scope of this paper but are available from the corresponding author on request. Once identified, the area under each peak was measured and its logarithm used as a measure of the quantity of cocaine on each banknote.

Any sample for which the difference between the total number of recorded banknotes and the number of peaks detected by the peak detection algorithm was greater than 10% of the total number of recorded banknotes (either way) was removed as such a discrepancy meant there were difficulties with the peak detection algorithm. Any sample with fewer than twenty banknotes was removed as the statistical procedures were unreliable with such few banknotes. Any samples for which the information on the currency or the total number of banknotes were not available were also removed. Samples were often analysed in multiple runs and

there were some missing runs within samples. For these, the longest contiguous section of the sample was included, with the rest discarded. The data from each sample were plotted and any outlying data points were further investigated and removed if they were found to be incorrectly identified peaks.

Two sets of training data were formed from the samples analysed. One set, C , containing data $\mathbf{y}$, was formed to develop the model for H_C , and another set B containing data $\mathbf{x}$, was formed to develop the model for H_B .

2.1. Banknotes that have been associated with criminal activity involving cocaine

The training data $\mathbf{y}$ for models developed for H_C are obtained from banknotes in criminal cases in which the defendant was convicted of a drug crime involving cocaine. Each case consisted of multiple exhibits, which may have been found in different locations. There were 29 cases containing at least one exhibit with greater than 20 banknotes. The 29 cases consist of between one and six exhibits, and there were a total of 70 exhibits which are known to have been associated with a person who has been involved in drug crime relating to cocaine. For future reference, any set of banknotes used in the analyses discussed that is said to be associated with crime will be known as an exhibit.

The training data $\mathbf{y}$ of banknotes used to develop models associated with H_C may include exhibits with two different types of cocaine contamination.

- C(a) The banknotes have not been contaminated with cocaine any more than those banknotes in general circulation. The contamination detected on the banknotes is consistent with that typically detected on general circulation banknotes. This quantity of contamination could have arisen innocently, or because the banknotes were not contaminated in the course of a crime (perhaps no drug was present at an exchange of money) by the person with whom they are associated.
- C(b) The banknotes were contaminated through their use in an illegal drug-related activity involving cocaine or in the course of other, legal, drug-related activity. This activity could have been carried out by some person other than the person eventually convicted.

It is expected that the quantities of cocaine on banknotes in C(a) will be lower than the quantities of cocaine on banknotes in C(b). See Fig. 1 for an illustration of the overlap in mean quantities of cocaine from samples from general circulation and from exhibits associated with crime (case). The left mode of the crime exhibits is formed of those exhibits in set C(a), and the right hand mode is formed of those exhibits in set C(b). Note that C(a) and C(b) are not propositions but descriptions of the possible sources of cocaine on the banknotes used as training data for the development of the models associated with the proposition H_C .

2.2. General circulation banknotes

The training data $\mathbf{x}$ of banknotes in set B are very unlikely to contain whole samples that are all associated with crime (although individual banknotes within a sample may well have been involved in a crime). Thus the banknotes defined by B and by C(a) are likely to be similarly contaminated. The banknotes defined by C(b) have been contaminated through their involvement in criminal activity involving cocaine, and so are likely to have higher levels of contamination.

There were 193 general circulation samples of English or Scottish currency obtained from a variety of locations around the UK. As shown in Table 1, a large number of the general circulation

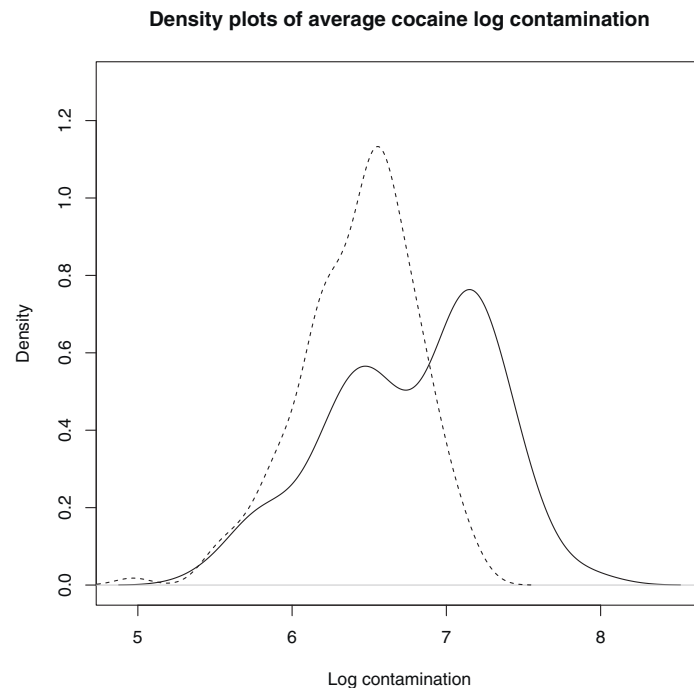


Fig. 1. Density plots of mean contamination of samples/exhibits. Dashed line – general circulation, solid line – crime/case.

samples were taken from the Bristol area. Tests have been carried out which did not find evidence that the region the banknotes in general circulation came from has an effect on the quantities of drug found [5]. For the calculation of likelihood ratios for seizures in a particular case it may be necessary to tailor the general circulation database for use in that case to the region in which the crime occurred or to refine the analysis with the use of regional factors (though these are not issues that are discussed further here). In addition, Table 2 shows that the majority of general circulation samples were taken from banks. This, again, could be a problem that requires further investigation if a defendant maintained that the banknotes had been acquired from some other type of location.

2.3. Banknotes for analysis

Data $\mathbf{z}$ used for testing with the likelihood ratio will generally have been provided by law enforcement agencies. This may be thought to place $\mathbf{z}$ in C , by definition. However, the definition of ‘association’ used here for the training set for C is that of conviction of a crime involving cocaine. Data from other cases brought by the law enforcement agencies have not been included in the analysis. This definition of a case is different from definitions used in previous work, when all seized banknotes were used as cases [1,4]. The data $\mathbf{z}$ are referred to also as the test set.

Table 1
Numbers of general circulation banknote samples in different police force areas.

Police force area	Number of samples	Police force area	Number of samples	Police force area	Number of samples
Avon and Somerset	78	West Midlands	4	Berkshire	1
Met	18	Nottinghamshire	3	Northamptonshire	1
West Yorkshire	11	Kent	3	Essex	1
Lancashire	9	Leicestershire	3	South Yorkshire	1
Hampshire	6	Hertfordshire	2	West Mercia	1
Northumbria	6	Merseyside	2	Dumfries and Galloway	1
Scotland (other/unknown)	6	Dorset	2	Wiltshire	1
Gloucestershire	4	Cleveland	2	Aberdeenshire	1
Gwent	4	Thames Valley	2	Flintshire	1
North Yorkshire	4	Humberside	2	Lincolnshire	1
Devon and Cornwall	4	Strathclyde	2		
South Wales	4	Unknown	2	Total	193

Table 2
Split of general circulation banknote samples between type of location.

Type of location	Number of samples	Type of location	Number of samples	Type of location	Number of samples
Bank	157	Sample provided by police	2	Casino	1
Newsagent	16	Cash point	2	Pub	1
Shopping Centre	4	Post office	1	Unknown	5
Bureau de Change	3	Sale of car	1	Total	193

3. Methods

3.1. Standardisation

The laboratory has four machines on which the banknotes may be analysed. The choice of which machine to use in a particular situation is made for operational reasons. The peak area measurements were standardised to allow for differences in response dependent on the machine used. In each run, a standard had been injected at the start of the analysis. The 193 general circulation samples were analysed, and where this standard could easily be identified (for 162 samples), the peak area was calculated. An average standard response for each of the four machines was calculated based on these 162 samples, and the ratio of each of the average standards in the second, third or fourth machines to the first machine was determined. The peak areas for each case exhibit (70 exhibits) and general circulation sample (193 samples) were then divided by the appropriate ratio, based on the machine which had been used to perform the analysis.

3.2. Exploratory data analysis

The aim of this research was to evaluate the likelihood ratio associated with the evidence relating to cocaine on banknotes for various models, given by the ratio of the likelihoods under H_C and H_B . Three models that account for autocorrelation were developed to model the quantity of cocaine on a sample of banknotes. The properties of the data that needed to be reflected by these models are discussed below.

3.2.1. Contamination on banknotes from general circulation

In previous studies of the quantities of cocaine on banknotes [1,4], it has been noted that using the percentage of contaminated banknotes within a sample as a statistic to distinguish between general circulation samples and crime exhibits is not generally possible due to the high frequency of contamination within samples from the general circulation, a conclusion supported by the current data. It is therefore not sufficient to focus on the proportion of contaminated banknotes; the quantity of contamination needs to be taken into account to differentiate between crime exhibits and general circulation samples. Fig. 1 shows the density plot of the mean quantity of banknote contamination of 193 general circulation samples with a dashed line against the mean contamination of the 70 crime exhibits with a solid line, where the means are determined from the individual quantities on the banknotes in the separate samples.

As can be seen, general circulation samples have mean quantities of contamination which, although generally lower than those of the positive case exhibits, are still high, and there is a substantial overlap between the two density plots. In addition, it can be seen that the means of the crime exhibits have a mode at around 6.5, in a similar position to the mode of the means of the general circulation samples, as well as a mode about 7.2. This is further evidence to support the suspicion that the set of crime exhibits C contains a large number of exhibits which are only contaminated in line with general circulation, and hence are in set $C(a)$.

3.2.2. Correlation

Banknotes in samples or small exhibits are analysed in their entirety whereas banknotes in larger exhibits are usually analysed in pre-selected groups, where the banknotes in each group were adjacent in the exhibit. Experiments carried out in [6] indicated that it was possible for drug traces to pass from one contaminated banknote to an adjacent one (although for heroin rather than cocaine). As a result, any transfer of drug that had occurred

between adjacent banknotes in the exhibit, as discussed in [6], would result in autocorrelation being present within the analysed exhibits. In addition, when banknotes are analysed, there is often no definitive end to the peak. Some of the ion counts occurring from cocaine on the previous banknote may be included in the reading for the next peak. This carry-over effect could also result in autocorrelated data.

When analysing the autocorrelation within the samples, it was found that about 90% of the 193 samples from general circulation and 80% of the 70 case exhibits had an autocorrelation of lag one. A sample or exhibit is said to have an autocorrelation of lag one if the approximate 95% confidence interval for the autocorrelation coefficient of lag one does not include zero, as described in p. 56 of [7]. The proportions of samples and exhibits with autocorrelations of higher order dropped to around 62% and 56% for samples and exhibits, respectively, at lag two, and 35% and 39% by lag five. The models described below account for lag one correlation but no more.

4. Models

The first model fits an autoregressive process of order one to the data. The second and third models use nonparametric conditional kernel density estimates to estimate the conditional densities of the peak areas of the banknotes, conditioning on the peak area of the previous banknote. One of these models uses a fixed bandwidth, the other uses a variable bandwidth. A fourth model, known as a *standard model*, assumes independence between measurements on separate banknotes and was used to compare the results with and without the independence assumption. The method for calculating the likelihood ratios using this model is given in section 5.

One method for the measurement of the relative performances of the four models is to consider rates of misleading evidence determined from testing the models on data of known sources. Misleading evidence can happen in one of two ways. Samples of banknotes in C can provide evidence to support H_B . Alternatively, samples of banknotes in B can provide evidence to support H_C .

Two sets of data, the training samples, are used for development of the models; these are

- $\mathbf{x} = \{x_{ij}; i = 1, \dots, m_B, j = 1, \dots, n_{B_i}\}$: the logarithms of the peak areas for cocaine on banknotes from general circulation as defined in Section 2.2; there are m_B samples with n_{B_i} banknotes in sample i .
- $\mathbf{y} = \{y_{ij}; i = 1, \dots, m_C, j = 1, \dots, n_{C_i}\}$: the logarithms of the peak areas of banknotes from criminal case exhibits for cocaine as defined in Section 2.1; there are m_C exhibits with n_{C_i} banknotes in exhibit i .

The questioned sample or test set is

- $\mathbf{z} = (z_1, z_2, \dots, z_n)$: the logarithms of the peak areas for cocaine of a sample of n banknotes of unknown origin. It may sometimes be known as the seized sample as it has been seized by a law enforcement agency.

Models are developed for H_B (using $\mathbf{x}$) and for H_C (using $\mathbf{y}$). The probability density function of $\mathbf{z}$ is then evaluated assuming separately H_B and then H_C and the likelihood ratio calculated.

4.1. Autoregressive models of order one

The form of the models for H_B and for H_C is the same, only the parameters are different. The model is described with generic notation here, with a general sample w substituting for x and y as appropriate.

The data of the logarithms of the peak areas of intensities of cocaine for a general sample are denoted $\mathbf{w} = (w_1, \dots, w_n)$. An autoregressive model AR(1) specifies the following relationship amongst the variables:

$$w_t - \mu = \alpha (w_{t-1} - \mu) + \epsilon_t \quad (1)$$

where $t = 2, \dots, n$; $\epsilon_t \sim N(0, \sigma^2)$ and $w_1 \sim N(\mu, \sigma^2)$, where $N(\mu, \sigma^2)$ is conventional notation denoting a normal distribution with mean μ and variance σ^2 . The autocorrelation coefficient α is a measure of the correlation between adjacent values (w_t, w_{t-1}) , in other words, the pairs of data $\{(w_2, w_1), \dots, (w_n, w_{n-1})\}$. The autocorrelation coefficient provides a measure of the association between the quantity of the drug on one banknote, indexed by t , with the quantity on the previous banknote, indexed by $t - 1$ indicating the order in which the banknotes were analysed. Like correlation, autocorrelation takes values between -1 and 1 . A value of zero indicates no association as can be seen by entering $\alpha = 0$ into (1). A value of one indicates that the value for banknote t only varies from the value for banknote $t - 1$ by random normal variation with variance σ^2 for $t = 2, \dots, n$.

4.1.1. Prior and posterior distributions

The training data are used in conjunction with prior distributions for the model parameters to determine posterior distributions for the parameters $\theta = (\mu, \sigma^2, \alpha)$ of the autoregressive model. The prior distributions used for the means and variances are similar to those used in [8,9], and a truncated normal prior is used for the autocorrelation parameter, as in [10], in order to provide compatibility with the development of other models that are beyond the scope of this paper.

The marginal prior distributions for the parameters are then given by:

- $\mu \sim N(\frac{1}{2}(\max(\mathbf{w}) + \min(\mathbf{w})), \text{range}(\mathbf{w})^2)$;
- $\sigma^2 \sim \text{IG}(2.5, \beta)$, where IG denotes the inverse gamma distribution and β is known as a hyperparameter; the form of the inverse gamma density function is given in Appendix A;
- $\beta \sim \Gamma(0.5, 4/\text{range}(\mathbf{w})^2)$.
- $\alpha \sim N(0, 0.25)$, with the autocorrelation restricted to lie between -1 and 1 .

The posterior distributions of the parameters μ , σ^2 and α were estimated using a Metropolis–Hastings sampler. Details of the sampler used for μ , σ^2 , α and β are given in Appendix A. (For more general information on Metropolis–Hastings samplers, see p. 289 of [11]).

The above procedure was used separately for each of the general circulation samples and each of the crime exhibits.

4.2. Nonparametric models

The autoregressive model assumes a normal distribution for the error terms. A nonparametric model, in which this assumption is dispensed with, is also used to fit the data. Use of the nonparametric model also means there is no need for prior distributions for parameters. As with autoregressive models, a generic notation is used. The point at which the probability density function is to be estimated is given by $\mathbf{w} = (w_1, \dots, w_n)$. The estimate is based on the i th sample $\mathbf{w}_i = (w_{i,1}, \dots, w_{i,n_{D_i}})$ from the appropriate training set where $i = 1, \dots, m_D$ and m_D is the number of samples ($D = B$) or exhibits ($D = C$) in the training set. The joint density function of $\mathbf{w}$ may be written as:

$$f_{D_i}(w_1, w_2, \dots, w_n) = f_{D_i}(w_1) f_{D_i}(w_2|w_1) \dots f_{D_i}(w_n|w_{n-1})$$

allowing for the autocorrelation of lag 1. The marginal density function $f_{D_i}(w_1)$ is estimated by a univariate kernel density estimate [12]. The conditional density function $f_{D_i}(w_t|w_{t-1})$, $t = 2, \dots, n$ for each $i \in (1, 2, \dots, m_D)$ can be estimated nonparametrically using kernel density estimation, at the point w_t , conditioned on the value of w_{t-1} , by:

$$\hat{f}_{D_i}(w_t|w_{t-1}) = \frac{\hat{g}_{D_i}(w_t, w_{t-1})}{\hat{r}_{D_i}(w_{t-1})} \quad (2)$$

The functions $\hat{g}_{D_i}$ and $\hat{r}_{D_i}$ are kernel density estimates based on sample i , given by:

$$\hat{g}_{D_i}(w_t, w_{t-1}) = \frac{1}{(n_{D_i} - 1)h_1 h_2} \sum_{j=2}^{j=n_{D_i}} K_1\left(\frac{w_t - w_{i,j}}{h_1}\right) K_2\left(\frac{w_{t-1} - w_{i,j-1}}{h_2}\right)$$

and

$$\hat{r}_{D_i}(w_{t-1}) = \frac{1}{(n_{D_i} - 1)h_3} \sum_{j=2}^{j=n_{D_i}} K_3\left(\frac{w_{t-1} - w_{i,j-1}}{h_3}\right),$$

Note that $\mathbf{w}$ refers to the observation at which a value for the density function is required.

Here, h_1 , h_2 and h_3 are bandwidths, and K_1 , K_2 and K_3 are kernel functions, see [13–15] for further details. See also [12] for earlier applications of kernel density estimation for independent observations in forensic science. In this analysis, the Gaussian kernel

$$K(s) = (2\pi)^{-1/2} \exp\left(\frac{-s^2}{2}\right)$$

is used for all three functions K_1 , K_2 and K_3 .

The functions $\hat{f}_{D_i}$ for each $i \in (1, 2, \dots, m_D)$ were calculated in R using the np package, [16]. This package sets $h_2 = h_3$ (the bandwidths that apply to the previous banknote in the numerator and denominator, respectively), and finds the optimal bandwidths h_1 and h_2 using cross-validation and maximising the estimated likelihood, a method which is described in [15]. As has already been seen, quantities of contamination vary considerably between different exhibits of banknotes. There is a need to evaluate the probability density function of a seizure of banknotes using each of the functions $\hat{f}_{D_i}$, each of which is based on a different exhibit of banknotes. Thus the probability density function may have to be evaluated where there are few data present. As a result, better results may be obtained by using a bandwidth which varies, depending on the amount of data present. Therefore, the functions $\hat{f}_{D_i}$ have been calculated using two different bandwidth types, for comparison. The first type is a fixed bandwidth, in which h_1 , h_2 and h_3 remain constant at all values of w_{it} and $w_{i,t-1}$ in the appropriate training sample. The second type is an adaptive nearest neighbour bandwidth, introduced in [17]. This type of bandwidth will vary, depending on the amount of data close by, becoming larger as the amount of nearby data reduces. The kernel density estimate of $\hat{r}_{D_i}(w_{t-1})$ becomes:

$$\hat{r}_{D_i}(w_{t-1}) = \frac{1}{n_{D_i} - 1} \sum_{j=2}^{j=n_{D_i}} \frac{1}{h_{3j}} K_3\left(\frac{w_{t-1} - w_{i,j-1}}{h_{3j}}\right)$$

where h_{3j} is the Euclidean distance from the point $w_{i,j-1}$ to the k th nearest sample point. See [18] for further details. Cross-validation is then used to select the value of k that maximises the estimated likelihood. The kernel density estimate of $\hat{g}_{D_i}(w_t, w_{t-1})$, with bandwidths h_1 and h_2 changes similarly.

5. Classification for a set of banknotes of unknown type

Details for the calculation of the LR for the autoregressive model of order one, the nonparametric models and the standard model are given here.

5.1. Autoregressive model

The probability density function $f(z_1, z_2, \dots, z_n | H_D)$ is given by:

$$\int_{\Theta_D} f(z_1 | \theta_D) f(z_2 | z_1, \theta_D) \dots f(z_n | z_{n-1}, \theta_D) f(\theta_D | \mathbf{w}) d\theta_D.$$

The distribution $f(\theta_D | \mathbf{w})$ is the posterior distribution of $\theta_D = (\mu_D, \sigma_D^2, \alpha_D)$ given the training set $\mathbf{w}$. It is estimated from a weighted average of the posterior distributions $f(\theta_{D_i} | \mathbf{w})$, details of the estimation of which are given in [Appendix A](#). The probability density function $f(z_1, z_2, \dots, z_n | H_D)$ is estimated using Monte Carlo integration. This involves evaluating $f(z_1 | \theta_D) f(z_2 | z_1, \theta_D) \dots f(z_n | z_{n-1}, \theta_D)$ for samples of θ_D , drawn from $f(\theta_D | \mathbf{w})$. The simple average of these evaluations gives an approximation to the probability density function $f(\mathbf{z} | H_D)$ for $D = C$ (the numerator of the LR) and $D = B$ (the denominator of the LR). Details of the procedure are contained in [\[19\]](#).

5.2. Nonparametric models

The probability density function of the data $\mathbf{z}$ assuming each of the prosecution and defence propositions has again to be calculated in order to calculate the likelihood ratio for the nonparametric models for fixed and adaptive bandwidths. For proposition H_D , the probability density function for $\mathbf{z}$ (substituted for $\mathbf{w} = (w_1, \dots, w_n)$ in Section 4.2) is then

$$\begin{aligned} f(z_1, z_2, \dots, z_n | H_D) &= f(z_1 | H_D) f(z_2 | z_1, H_D) \dots f(z_n | z_{n-1}, H_D) \\ &\simeq \sum_{i=1}^{m_D} v_i f_{D_i}(z_1 | H_D) f_{D_i}(z_2 | z_1, H_D) \dots f_{D_i}(z_n | z_{n-1}, H_D) \end{aligned} \quad (3)$$

where $f_{D_i}(z_t | z_{t-1})$ is the conditional density of z_t given z_{t-1} estimated using the equivalent conditional density function of sample or exhibit i , $i \in (1, \dots, m_D)$, for banknotes $t = 2, \dots, n$, $f_{D_i}(z_1)$ is the marginal density for banknote 1 and $v_i = n_{D_i} / \sum_{i=1}^{m_D} n_{D_i}$ is a weight assigned to each sample or exhibit i , with $\sum v_i = 1$.

The nonparametric method of estimating the functions f_{D_i} by $\hat{f}_{D_i}$ for each of $i \in (1, 2, \dots, m_D)$ has been described in Section 4.2. This is used to estimate the probability density function of $\mathbf{z}$ assuming the proposition, H_B or H_C .

5.3. Standard model

As a comparison to the three models introduced here, rates of misleading evidence were also calculated for the standard model for independent and normally distributed univariate data; i.e., measurements between adjacent banknotes are assumed independent. The method given in [\[12\]](#), which uses kernel density estimates for the between sample distribution of the mean was used, owing to the large variation in contamination on different samples and exhibits of banknotes. A slight adaptation of the method presented there is required, because the problem being considered here is a discrimination problem, not a comparison problem (comparing of a control and recovered item). The estimate of the between sample variance also has a slight adjustment, because the number of banknotes in each sample or exhibit varies between samples. The likelihood ratio for a seized sample $\mathbf{z}$ with n banknotes, for the discrimination problem, is given by

$$LR = \frac{\sum_{i=1}^{m_C} (m_C \sqrt{\tau_C^2 + n \lambda_C^2 s_C^2})^{-1} \exp[-n(\bar{\mathbf{z}} - \bar{\mathbf{y}}_i)^2 / 2(\tau_C^2 + n \lambda_C^2 s_C^2)]}{\sum_{i=1}^{m_B} (m_B \sqrt{\tau_B^2 + n \lambda_B^2 s_B^2})^{-1} \exp[-n(\bar{\mathbf{z}} - \bar{\mathbf{x}}_i)^2 / 2(\tau_B^2 + n \lambda_B^2 s_B^2)]} \quad (4)$$

where τ_C^2 and s_C^2 are respectively the within and between sample variances for the crime exhibits and τ_B^2 and s_B^2 are the within and between variances for general circulation samples. The bandwidths of the kernel density estimates for the between sample distributions of the mean are given by $\lambda_C s_C$ (crime exhibits) and $\lambda_B s_B$ (general circulation samples). The mean $\bar{\mathbf{y}}_i$ refers to the mean given by the equation

$$\bar{\mathbf{y}}_i = \sum_{j=1}^{n_{C_i}} \mathbf{y}_{ij}$$

with $\bar{\mathbf{x}}_i$ defined similarly.

The within sample variance for the crime exhibits is estimated by

$$\hat{\tau}_C^2 = \sum_{i=1}^{m_C} \sum_{j=1}^{n_{C_i}} \frac{(y_{ij} - \bar{\mathbf{y}}_i)^2}{N_C - m_C}$$

where N_C denotes the total number of datapoints so that

$$N_C = \sum_{i=1}^{m_C} n_{C_i}$$

and the between sample variance for the crime exhibits is estimated by

$$\hat{s}_C^2 = \sum_{i=1}^{m_C} \frac{n_{C_i} (\bar{\mathbf{y}}_i - \bar{\mathbf{y}})^2}{\tilde{n}_C (m_C - 1)} - \frac{\hat{\tau}_C^2}{\tilde{n}_C}$$

where the value of $\tilde{n}_C$ is given by

$$\tilde{n}_C = \frac{1}{m_C - 1} \left(N_C - \frac{\sum_{i=1}^{m_C} n_{C_i}^2}{N_C} \right).$$

The estimators used are the ANOVA estimators given on pages 19–21 of [\[20\]](#) The bandwidth parameter λ_C is selected using Silverman's rule of thumb [\[15\]](#), so that

$$\lambda_C = \left(\frac{4}{3m_C} \right)^{-1/5}$$

The between and within sample variances and the parameter λ_B for general circulation banknotes are estimated similarly by replacing y with x and C with B .

6. Results

The results are based on 70 exhibits and 193 general circulation samples. For the 70 exhibits, the sizes ranged from 20 to 1099 detected peaks and for the 193 samples, the sizes ranged from 21 to 257 detected peaks, each peak corresponding to an individual banknote.

6.1. Parameter estimation for the autoregressive model

Posterior distributions for θ_C and θ_B were derived for each crime exhibit (θ_C) and each general circulation sample (θ_B) using the procedures described in Section 4.1.

In general, the means of general circulation samples, μ_B were lower than the means of the crime exhibits μ_C . This is as expected:

Table 3

Rates of misleading evidence, estimated as (r/n) , where r is the number of exhibits/samples out of n analysed for which the likelihood ratio is misleading in each context.

	AR(1)	Nonparametric fixed bw	Nonparametric adaptive mn	Standard
Exhibit	0.37 (26/70)	0.27 (19/70)	0.26 (18/70)	0.50 (35/70)
General circulation	0.16 (30/193)	0.32 (62/193)	0.27 (52/193)	0.14 (26/193)

the crime banknotes should have higher quantities of contamination than general circulation banknotes.

The variances of both general circulation samples and crime exhibits were similar. The autocorrelation parameters varied between samples, but generally lay between 0.2 and 0.7. A few samples (both general circulation and crime exhibit) had an autocorrelation close to one. Such a value for the autocorrelation is indicative of a sample with very similarly contaminated banknotes.

6.2. Likelihood ratios

Rates of misleading evidence were calculated for each of the four models fitted: the autoregressive model of order one, the nonparametric model with fixed bandwidth, the nonparametric model with an adaptive bandwidth and the standard model. The results are displayed in Table 3.

In order to calculate these rates, each exhibit or sample was taken out in turn and treated as the seized sample of banknotes, z . The remaining crime exhibits (removing the chosen exhibit and the other exhibits from the same case) and general circulation samples (removing the chosen sample) were then treated as the measurements y and x , respectively, and used to calculate the likelihood ratio for the removed sample. If this likelihood ratio was greater than one, the evidence z was said to support H_C and if it was less than one, it was said to support H_B . The proportion of times that the evidence was said to support the proposition from which it did not originate is an estimate of the rate of misleading evidence.

The smallest rate of misleading evidence for the crime exhibits is achieved with the nonparametric model with adaptive bandwidth, followed by the nonparametric model with fixed bandwidth and then the AR(1) model. The rates are, however, quite large, with just over one-quarter of the evidence provided from crime exhibits being said to be misleading by the nonparametric model with adaptive bandwidth. These large numbers are caused by the large number of exhibits which are actually not contaminated any more than general circulation (and so are in set $C(a)$), and yet are still included in the test.

The rate of misleading evidence for the standard model for general circulation samples is the smallest of the rates, with slightly more than one in eight samples giving misleading evidence. The AR(1) model has slightly more than 1 in seven samples giving misleading evidence. Note that the rate of misleading evidence is used as a general measure of performance of a statistical model and is not a suggestion of a probability of error in a particular case. There are many other factors in addition to the output of a statistical model which will be considered in the assessment of the evidence in a particular case. One of the purposes of the research was to compare several models which take account of autocorrelation and assess their performance in comparison with a model which, incorrectly, assumed independence. The benefit of this work is the provision of a measure for performance for the models in contrast to assessments just based on 'general experience'.

7. Discussion

Consider the calculation of a likelihood ratio for a sample of banknotes for which $C(a)$ is the correct set. A model which is

accurate will provide a likelihood ratio of less than one, despite H_C being the correct proposition (though this is not known by the scientist). This is because samples of banknotes in set $C(a)$ have contamination in line with general circulation. This result will be interpreted as evidence that supports H_B . The model has provided misleading evidence; it supports a proposition that is not correct. This means that the approaches given here cannot be used to support proposition H_B because it is known to be wrong an unknown proportion of times even with an accurate model with the proportion depending on the proportion of exhibits in set $C(a)$. The approach can, however, be used to support the prosecution proposition.

Rates of misleading evidence are only one measure of performance. One of their shortcomings is that they do not give an indication of the size of the likelihood ratios. In practice, a likelihood ratio close to one does not provide strong evidence that a seized sample belongs to a particular class. A method which gives large likelihood ratios is of more value when being used in evidence evaluation. The Tippett plots in Fig. 2 help in the assessment of performance by illustrating likelihood ratios graphically.

The Tippett plots in Fig. 2 indicate the proportion of general circulation samples (dashed) and crime exhibits (solid) that have a log-likelihood ratio greater than the value on the horizontal axis of the plot. Fig. 2 shows that the autoregressive model of order one has log-likelihoods which are much closer to the neutral value of zero than the nonparametric models, meaning that it is less useful for discriminating between crime exhibits and general circulation samples. The two nonparametric models give larger likelihood ratios than the AR(1) model, but there are some large false positive values for these models. It is of concern that these large values could be erroneous. The Tippett plot for the standard model shows that most of the general circulation samples have a log-likelihood ratio very close to zero when this model is used. The absolute values of log-likelihood ratios for false positive general circulation samples are small, similar in size to those for the autoregressive models, with the exception of one general circulation sample which has a large erroneous value.

The Tippett plots for the nonparametric models have a lot of outliers. These outliers are thought to be caused by problems with the reliability of nonparametric estimation techniques in areas where there are few data. This causes problems when estimating likelihood ratios, and leads to results which are not very reliable.

Twelve of the 70 crime exhibits were declared as contaminated (to a degree that is unlikely to be observed in the general population) by forensic experts prior to their use in this analysis. The likelihood ratios obtained for these twelve exhibits using the four models were analysed. All twelve had likelihood ratios greater than one for all four models, suggesting that the models are correctly assigning support for H_C for these twelve exhibits. It was found that the five exhibits of these twelve with the fewest number of banknotes had much larger likelihood ratios when the standard model was used, than when the autoregressive model was used. This could suggest that by assuming independence (and using the standard model), there is a risk of overstating the likelihood ratio for small samples. A small group of highly contaminated banknotes which are close together in a small sample will influence the value of the mean, used in the standard model. This influence is reduced when autocorrelation is taken into account.

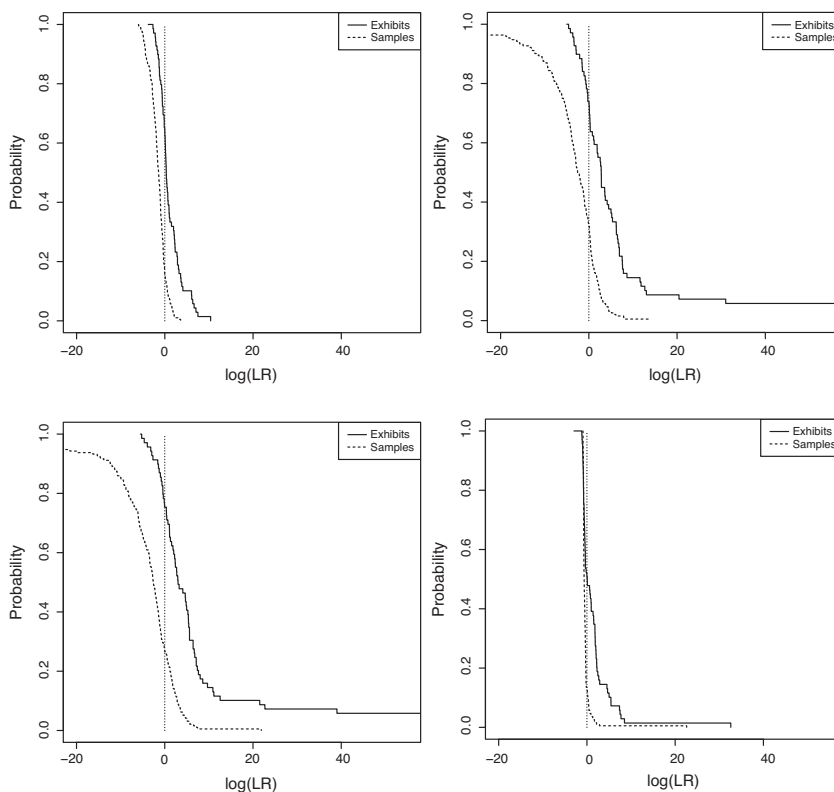


Fig. 2. Tippet plots of log likelihood ratio values to show the probability that the log likelihood ratio is greater than $\log(LR)$. Clockwise from top left – AR(1), fixed bandwidth, standard model, adaptive bandwidth. General circulation results are indicated with a dashed line and crime exhibit results with a solid line. A background value greater than zero is a false positive. A crime exhibit value less than zero is a false negative.

The problem of calculating likelihood ratios for samples in set C(a) means that we do not expect to obtain a low rate of misleading evidence for the crime exhibits, and so this is not a useful measure for evaluating the performance of the models. The most important desired outcome is that the rate of misleading evidence for general circulation samples is low, so as to avoid false support for proposition H_C (a false positive). In order to demonstrate that the method of analysing cocaine traces on banknotes combined with likelihood ratio evaluation with one of the statistical models described is useful, it is required that some of the crime exhibits have a likelihood ratio which is large, as this means that evidence can be provided to support H_C . It can be seen from the rates of misleading evidence in Table 3 that the autoregressive and standard models have the smallest rates of misleading evidence for general circulation samples. Also, as seen from Fig. 2, the Tippet plots for the AR(1) and standard models for crime exhibits show that log likelihood ratio values are large enough to provide support for H_C .

The recommendation is that the AR(1) model is used to analyse the log peak area data arising from the analysis of traces of cocaine on banknotes. The AR(1) model provides a low rate of misleading evidence for general circulation samples and provides sensible likelihood ratio values as a measure of support for proposition H_C . The standard model is included in the discussion for comparison purposes only and is not recommended for general use as it does not take account of autocorrelation, which can result in overstated likelihood ratios.

8. Conclusion

Previous methods described in [21] (though for heroin rather than cocaine) for the analysis of banknotes have used the percentage of contaminated banknotes in a collection as a measure

of contamination. However, many banknotes in general circulation are contaminated with cocaine so an approach based on such methods for cocaine is not discriminatory. Another method described in [1] uses intensity of contamination and calculates the likelihood ratio for just one banknote. Methods described in [2] assume independence between banknotes, an assumption which does not hold in the data analysed here. Conversely, the analytical method described in [21], using tandem mass spectrometry, enables contamination to be measured by the approximate quantity of drug on each individual banknote. Three statistical methods described here allow for correlation between measurements on neighbouring banknotes. Results are compared with a model assuming independence. There is an elementary extension for data with autocorrelations of lag greater than one.

The current paper has described initial work on univariate data, those of measurements of the log peak area of the cocaine product ion m/z 105. However, there are other variables that may be considered. Consideration of these would require a multivariate generalisation of the univariate method described here. For example, five drugs are analysed and data on log peak area and log peak height are available for each drug. Use of all these data requires a ten-dimensional model, development of which is the subject of further work.

A more sophisticated model, a hidden Markov model, is under development [19]. The hidden Markov model represents more directly, than do the models described above, the situation in which samples of banknotes may be a mixture of banknotes from general circulation and banknotes that are associated with crime, where there is a corresponding mix of quantities of cocaine on different banknotes within the same sample.

The methods described here provide a rigorous statistical analysis of the evaluation of evidence of quantities of cocaine on banknotes. However, they may also be used for evidence

evaluation for continuous data from other evidential types where there is autocorrelation between adjacent items.

Acknowledgements

The research is supported by an EPSRC CASE award, voucher number 009002219. The authors are grateful to the reviewers for their very helpful comments which have led to considerable improvements in the paper.

Appendix A

The methods for estimating the parameters $\theta = (\mu, \sigma^2, \alpha, \beta)$ from a single sample or exhibit are described here. For ease of notation, θ is presented without a subscript. A subscript needs to be added appropriately for application in the text. An iterative simulation procedure known as a Metropolis–Hastings (MH) sampler is used to obtain an estimated sample from a probability distribution. Consider a general set of data $\mathbf{w} = \{w_t, t = 1, \dots, n\}$ with parameters θ to be estimated. Here, an estimate of β is needed in order to estimate σ^2 , even though the resulting estimates of β are not needed to estimate the likelihood ratio. The MH sampler is used to obtain an estimated sample from the probability density function $f(\theta | \mathbf{w})$. This enables a posterior distribution of the parameters represented by θ to be obtained.

The likelihood $f(\mathbf{w} | \theta)$ is given by:

$$f(\mathbf{w} | \theta) = (2\pi\sigma^2)^{-\frac{n}{2}} \exp\left[-\frac{1}{2\sigma^2}(w_1 - \mu)^2\right] \times \exp\left[-\sum_{t=2}^n \left(\frac{1}{2\sigma^2}(w_t - \mu + \alpha\mu - \alpha w_{t-1})^2\right)\right] \quad (5)$$

The prior distributions of μ , σ^2 , α and β are taken to be normal, inverse gamma, truncated normal and gamma respectively, with parameters as follows:

$$\mu \sim N(\mu_0, V_\mu), \quad \sigma^2 \sim IG(\gamma, \beta), \quad \alpha \sim N(\alpha_0, V_\alpha)I(|\alpha| < 1), \\ \beta \sim \Gamma(g, h),$$

where μ_0 , V_μ , γ , α_0 , V_α , g and h take the values given in the body of the text and

- $IG(\gamma, \beta)$ denotes the inverse gamma distribution such that if $\sigma^2 \sim IG(\gamma, \beta)$, then $f(\sigma^2 | \gamma, \beta) = \frac{\beta^\gamma}{\Gamma(\gamma)} (\sigma^2)^{-(\gamma+1)} e^{-\beta/\sigma^2}$; $\beta > 0, \gamma > 0, \sigma > 0$, and $\Gamma(\gamma) = \int_0^\infty t^{\gamma-1} e^{-t} dt$.
- and $\Gamma(g, h)$ denotes the gamma distribution such that if $\beta \sim \Gamma(g, h)$, then $f(\beta | g, h) = \frac{h^g}{\Gamma(g)} \beta^{g-1} e^{-h\beta}$; $g > 0, h > 0, \beta > 0$,
- and $I(|\alpha| < 1)$ is the indicator function such that

$$I(|\alpha| < 1) = \begin{cases} 1 & \text{if } |\alpha| < 1, \\ 0 & \text{if } |\alpha| \geq 1. \end{cases}$$

The joint prior distribution, $f(\theta)$, is therefore given by:

$$f(\theta) = f(\mu) f(\sigma^2 | \beta) f(\beta) f(\alpha) \propto \exp\left[-\frac{1}{2V_\mu}(\mu - \mu_0)^2\right] \times \beta^\gamma \sigma^{-(2\gamma+2)} \exp\left[-\frac{\beta}{\sigma^2}\right] \beta^{g-1} \exp(-h\beta) \times \exp\left[-\frac{1}{2V_\alpha}(\alpha - \alpha_0)^2\right] I(|\alpha| < 1) \quad (6)$$

A MH sampler updates the parameters θ iteratively in steps, and accepts or rejects the update according to an acceptance probability which is defined below. Denote the parameters at step j by $\theta^{(j)} = (\mu^{(j)}, \sigma^{2(j)}, \alpha^{(j)}, \beta^{(j)})$ and the updated parameters as $\theta' = (\mu', \sigma'^2, \alpha', \beta')$, then the MH sampler updates the parameters as follows:

$$\mu' = \mu^{(j)} + \varepsilon_1, \quad \log(\sigma'^2) = \log(\sigma^{2(j)}) + \varepsilon_2, \\ \alpha' = \alpha^{(j)} + \varepsilon_3, \log(\beta') = \log(\beta^{(j)}) + \varepsilon_4.$$

Here, ε_k is a normally distributed random variable, with zero mean and variance V_k for $k \in \{1, 2, 3, 4\}$. It has been shown [22] that the V_k should be chosen so that the number of accepted updates is close to 25%.

The updated parameter θ' is accepted (meaning that $\theta^{(j+1)}$ is set to θ') if $U < \min(1, A)$, where U is drawn from a uniform distribution on the interval $[0, 1]$, and A is given by:

$$A = \frac{f(\mathbf{w} | \theta') f(\theta') \sigma'^2 \beta'}{f(\mathbf{w} | \theta^{(j)}) f(\theta^{(j)}) \sigma^{2(j)} \beta^{(j)}} \quad (7)$$

If $U > \min(1, A)$, the updated parameter values θ' are not accepted, and $\theta^{(j+1)}$ is set to $\theta^{(j)}$ (i.e. the parameter values do not change). The likelihoods $f(\mathbf{w} | \theta')$ and $f(\mathbf{w} | \theta^{(j)})$ in (7) can be calculated by substituting θ' and $\theta^{(j)}$ into the equation for the likelihood in (5). Similarly the values $f(\theta')$ and $f(\theta^{(j)})$ are calculated by substituting θ' and $\theta^{(j)}$ into the equation for the prior distribution in (6). The terms $\sigma'^2, \beta', \sigma^{2(j)}$ and $\beta^{(j)}$ in (7) are included to allow for the fact that the parameters σ^2 and β are updated via their logarithms. The prior distribution must therefore be transformed, and these extra terms are the Jacobian of that transformation.

By starting off the Metropolis–Hastings sampler with some initial values for step $j = 1$, we can then iterate through the process described, updating and then accepting or rejecting the updates, to obtain the required estimated samples from the probability density function $f(\theta | \mathbf{w})$.

In total, 250,000 samples were taken. In order to give the sampler a chance to move away from the initial values chosen, the first 50,000 of these 250,000 samples were discarded. As each update depends on the parameters at the previous step, there is autocorrelation between samples at small lag numbers (so samples with step numbers close together). To remove this autocorrelation, so that our samples are close to being independent, only every 25th sample was considered, discarding the rest. This left 8000 estimated samples from $f(\theta | \mathbf{w})$. Convergence to the distributions of the parameter estimates for use in the determination of the likelihood ratios was checked by visual inspection of the plots of the evolution of these samples over time.

References

- [1] L. Besson, Détection des stupéfiants par IMS, Master's thesis, University of Lausanne, 2004.
- [2] F. Taroni, S. Bozza, A. Biedermann, P. Garbolino, C.G.G. Aitken, Data Analysis in Forensic Science. A Bayesian Decision Perspective, John Wiley and Sons Ltd., Chichester, 2010.
- [3] D.V. Lindley, A problem in forensic science, *Biometrika* 64 (1977) 207–213.
- [4] S.J. Dixon, R.G. Brereton, J.F. Carter, R. Sleeman, Determination of cocaine contamination on banknotes using tandem mass spectrometry and pattern recognition, *Analytica Chimica Acta* 559 (2006) 54–63.
- [5] K.A. Ebejer, G.R. Lloyd, R.G. Brereton, J.F. Carter, R. Sleeman, Factors influencing the contamination of UK banknotes with drugs of abuse, *Forensic Science International* 171 (2007) 165–170.
- [6] K.A. Ebejer, J. Winn, J.F. Carter, R. Sleeman, J. Parker, F. Körber, The difference between drug money and a "lifetime's savings", *Forensic Science International* 167 (2007) 94–101.

- [7] C. Chatfield, *The Analysis of Time Series: An Introduction*, 6th ed., Chapman and Hall, London, 2004.
- [8] T. Rydén, EM versus Markov chain Monte Carlo for estimation of hidden Markov models: a computational perspective, *Bayesian Analysis* 3 (2008) 659–688.
- [9] S. Richardson, P.J. Green, On Bayesian analysis of mixtures with an unknown number of components (with discussion), *Journal of the Royal Statistical Society: Series B* 59 (1997) 731–792.
- [10] J.H. Albert, S. Chib, Bayes inference via Gibbs sampling of autoregressive time series subject to Markov mean and variance shifts, *Journal of Business and Economic Statistics* 11 (1993) 1–15.
- [11] A. Gelman, J.B. Carlin, H.S. Stern, D.B. Rubin, *Bayesian Data Analysis*, 2nd ed., Chapman and Hall, London, 2004.
- [12] C.G.G. Aitken, F. Taroni, *Statistics and the Evaluation of Evidence for Forensic Scientists*, John Wiley and Sons Ltd, Chichester, 2004.
- [13] J. Fan, Q. Yao, H. Tong, Estimation of conditional densities and sensitivity measures in nonlinear dynamical systems, *Biometrika* 83 (1996) 189–206.
- [14] P. Hall, J. Racine, Q. Li, Cross-validation and the estimation of conditional probability densities, *Journal of the American Statistical Association* 99 (468) (2004) 1015–1026.
- [15] B.W. Silverman, *Density Estimation for Statistics and Data Analysis*, Chapman and Hall, London, 1986.
- [16] T. Hayfield, J.S. Racine, Nonparametric econometrics: the NP package, *Journal of Statistical Software* 27 (5) (2008).
- [17] L. Breiman, W. Meisel, E. Purcell, Variable kernel estimates of multivariate densities, *Technometrics* 19 (1977) 135–144.
- [18] G.R. Terrell, D.W. Scott, Variable kernel density estimation, *Annals of Statistics* 20 (1992) 1236–1265.
- [19] A. Wilson, C.G.G. Aitken, R. Sleeman, J.F. Carter, The evaluation of evidence for autocorrelated data with an example relating to traces of cocaine on banknotes, 2013 (under revision).
- [20] T.A. Snijders, R.J. Bosker, *Multilevel Analysis: An Introduction to Basic and Advanced Multilevel Modeling*, 2nd ed., Sage, London, 2012.
- [21] K.A. Ebejer, R.G. Brereton, J.F. Carter, S.L. Ollerton, R. Sleeman, Rapid comparison of diacetylmorphine on banknotes by tandem mass spectrometry, *Rapid Communications in Mass Spectrometry* 19 (2005) 2137–2143.
- [22] A. Gelman, G.O. Roberts, W.R. Gilks, in: J.M. Bernardo, J.O. Berger, A.P. Dawid, A.F.M. Smith (Eds.), *Efficient Metropolis Jumping Rules*, *Bayesian Statistics 5*, Oxford University Press, New York, 1996, pp. 599–607.