

AD-A085 471 NAVAL RESEARCH LAB WASHINGTON DC
PROPERTIES OF CROSS-ENTROPY MINIMIZATION.(U)
APR 80 J E SHORE, R W JOHNSON

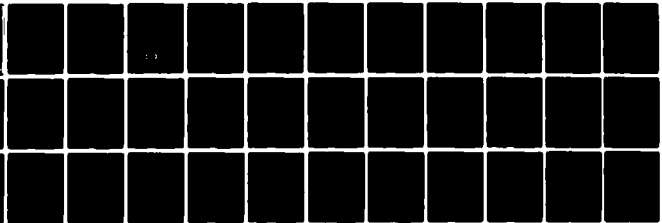
F/6 9/4

UNCLASSIFIED NRL-MR-4189

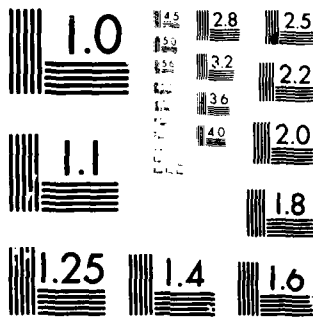
SBIE-AD-E000 431

NL

[16]
an
2000000



END
DATE
FILMED
7-80
DTIC



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A



SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER NRL Memorandum Report 4189	2. GOVT ACCESSION NO. AID-A085	3. RECIPIENT'S CATALOG NUMBER 477
4. TITLE (and Subtitle) PROPERTIES OF CROSS-ENTROPY MINIMIZATION		5. TYPE OF REPORT & PERIOD COVERED Interim report on a continuing NRL problem.
7. AUTHOR(s) J. E. Shore and R. W. Johnson		6. PERFORMING ORG. REPORT NUMBER
9. PERFORMING ORGANIZATION NAME AND ADDRESS Naval Research Laboratory Washington, DC 20375		8. CONTRACT OR GRANT NUMBER(s)
11. CONTROLLING OFFICE NAME AND ADDRESS Naval Research Laboratory Washington, DC 20375		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS 61153N, RR014-09-41 NRL Problem 75-0102-0-0
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		12. REPORT DATE April 1, 1980
		13. NUMBER OF PAGES 39
		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		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Information theory Entropy maximization Cross-entropy Minimum discrimination information Directed divergence Cross-entropy minimization		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) The principle of minimum cross-entropy (minimum directed divergence, minimum discrimination information) provides a general method of inference about an unknown probability density when there exists a prior estimate of the density and new information in the form of constraints on expected values. Previous work has shown that cross-entropy minimization can be characterized axiomatically by a small set of consistency axioms. Our purpose in this paper is to collect in one place and prove various fundamental properties of cross-entropy minimization. We include examples, and we discuss general analytic and computational methods of finding minimum cross-entropy probability densities. (Continues)		

DD FORM 1473
1 JAN 73

EDITION OF 1 NOV 65 IS OBSOLETE
S/N 0102-014-6601

1

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

20. ABSTRACT (Continued)

An interesting aspect of the results presented in this paper is the interplay between properties of cross-entropy minimization as an inference procedure and properties of cross-entropy as an information measure. Cross-entropy's well-known and unique properties as an information measure in the case of arbitrary probability densities are extended and strengthened when one of the densities involved is the result of cross-entropy minimization. For example, cross-entropy does not in general satisfy a triangle relation involving three arbitrary probability densities. But in certain important cases involving densities that result from cross-entropy minimization, cross-entropy satisfies triangle inequalities and triangle equalities.

CONTENTS

I.	INTRODUCTION.....	1
II.	DEFINITIONS AND NOTATION.....	2
III.	PROPERTIES GIVEN GENERAL CONSTRAINTS.....	5
IV.	PROPERTIES GIVEN EQUALITY CONSTRAINTS.....	19
V.	GENERAL DISCUSSION.....	25
VI.	ACKNOWLEDGEMENTS.....	25
	APPENDIX A--Mathematics of Cross-Entropy Minimization.....	26
	APPENDIX B--Remark on the Discussion of Property 12.....	34
	REFERENCES.....	35

DTIC
ELECTE
S JUN 17 1980 **D**
B

ACCESSION for		
NTIS	White Section	☒
DDC	Buff Section	☐
UNANNOUNCED		☐
JUSTIFICATION		
BY		
DISTRIBUTION/AVAILABILITY CODES		
Dist.	AVAIL. and/or	SPECIAL
A		-

I. INTRODUCTION

The principle of minimum cross-entropy provides a general method of inference about an unknown probability density $q^\dagger$ when there exists a prior estimate of $q^\dagger$ and new information about $q^\dagger$ in the form of constraints on expected values. The principle states that, of all the densities that satisfy the constraints, one should choose the posterior q with the least cross-entropy $H[q,p] = \int d\mathbf{x} q(\mathbf{x}) \log(q(\mathbf{x})/p(\mathbf{x}))$, where p is a prior estimate of $q^\dagger$. Cross-entropy minimization was first introduced by Kullback [1], who called it minimum directed divergence and minimum discrimination information. The principle of maximum entropy [2],[3] is equivalent to cross-entropy minimization in the special case of discrete spaces and uniform priors.

It is useful and convenient to view cross-entropy minimization as one implementation of an abstract information operator $\circ$ that takes two arguments --- a prior and new information --- and yields a posterior. Thus, we write the posterior q as $q = p \circ I$, where I stands for the known constraints on expected values. Recent work has shown that, if the operator $\circ$ is required to satisfy certain axioms of consistent inference, and if $\circ$ is implemented by means of functional minimization, then the principle of minimum cross-entropy follows necessarily [4].

Cross-entropy minimization satisfies a variety of interesting and useful properties beyond those expressed or implied by the axioms in [4]. Some of these just reflect well-known properties of cross-entropy [1],[5], but there are surprising differences as well. For example, cross-entropy does not in general satisfy a triangle relations involving three arbitrary probability densities. But in certain important cases involving densities that result

Note: Manuscript submitted January 23, 1980.

from cross-entropy minimization, cross-entropy satisfies triangle inequalities and triangle equalities. (See Properties 10, 12, and 13.)

It is the purpose the present paper to state and prove various fundamental properties of cross-entropy minimization. For completeness, we also restate the axioms from [4]. After introducing necessary definitions and notation in Section II, we consider first properties that are valid for both equality and inequality constraints on expected values (Section III) and then properties that are valid only for equality constraints (Section IV). We conclude with a brief discussion in Section V. We also include an Appendix in which we discuss general analytic and computational methods of finding minimum cross-entropy posteriors.

II. DEFINITIONS AND NOTATION

In this section, we introduce the same notation as in [4, Section II]. The discussion here places somewhat greater emphasis on mathematical questions relating to the existence of minimum-cross-entropy solutions. (See also the discussion following Property 1.)

We use lower-case boldface Roman letters for system states, which may be multidimensional, and upper-case boldface Roman letters for sets of system states. We use lower-case Roman letters for probability densities, and upper case script letters for sets of probability densities. Thus, let $\underline{x}$ be a state of some system that has a set $\underline{D}$ of possible states. Let $\mathcal{D}$ be the set of all probability densities q on $\underline{D}$ such that $q(\underline{x}) \geq 0$ for $\underline{x} \in \underline{D}$ and

$$\int_{\underline{D}} d\underline{x} q(\underline{x}) = 1 . \quad (1)$$

We use a dagger $\dagger$ to distinguish the system's unknown "true" state probability

density $q^* \in \mathcal{D}$. When $S \subseteq \mathcal{D}$ is some set of states, we write $q(x \in S)$ for the set of values $q(x)$ with $x \in S$.

New information takes the form of linear equality constraints

$$\int_{\mathcal{D}} dx \, q^*(x) a_k(x) = \bar{a}_k \quad (2)$$

and inequality constraints

$$\int_{\mathcal{D}} dx \, q^*(x) c_k(x) \geq \bar{c}_k \quad (3)$$

for known sets of functions a_k, c_k , and known values $\bar{a}_k, \bar{c}_k$. The probability densities that satisfy such constraints always comprise a convex subset $\mathcal{V}$ of $\mathcal{D}$. (A set $\mathcal{V}$ is convex if, given $0 \leq A \leq 1$ and $q, r \in \mathcal{V}$, it contains the weighted average $Aq + (1-A)r$.) We refer to the functions a_k, c_k as constraint functions and $\mathcal{V}$ as a constraint set. For a given constraint set there may of course be more than one set of constraint functions in terms of which it may be defined. We frequently suppress mention of a particular set of constraint functions, using the notation $I = (q^* \in \mathcal{V})$ to mean that q^* is a member of the constraint set $\mathcal{V} \subseteq \mathcal{D}$ and referring to I as a constraint. We use upper-case Roman letters for constraints.

Let $p \in \mathcal{D}$ be some prior density that is an estimate of q^* obtained, by any means, prior to learning I . We require that priors be strictly positive:

$$p(x \in \mathcal{D}) > 0 \quad (4)$$

(This restriction is discussed below.) Given a prior p and new information I , the posterior density $q \in \mathcal{V}$ that results from taking I into account is chosen by minimizing the cross-entropy $H[q, p]$ in the constraint set $\mathcal{V}$:

$$H[q, p] = \min_{q' \in \mathcal{V}} H[q', p] \quad (5)$$

where

$$H[q, p] = \int_{\mathcal{D}} dx \, q(x) \log(q(x)/p(x)). \quad (6)$$

We introduce an "information operator" $\circ$ that expresses (5) using the notation

$$q = p \circ I. \quad (7)$$

The operator $\circ$ takes two arguments --- a prior and new information --- and yields a posterior.

For some subset $\mathcal{S} \subseteq \mathcal{D}$ of states and $x \in \mathcal{S}$, let

$$q(x | x \in \mathcal{S}) = q(x) / \int_{\mathcal{S}} dx' \, q(x') \quad (8)$$

be the conditional density, given $x \in \mathcal{S}$, corresponding to any $q \in \mathcal{D}$. We use

$$q(x | x \in \mathcal{S}) = q * \mathcal{S} \quad (9)$$

as a shorthand notation for (8).

In making the restriction (4) we assume that $\mathcal{D}$ is the set of states that are possible according to prior information. We do not impose a similar restriction on the posterior $q = p \circ I$ since I may rule out states currently thought to be possible. If this happens, then $\mathcal{D}$ must be redefined before q is used as a prior in a further application of $\circ$. The restriction (4) does not significantly restrict our results, but it does help in avoiding certain technical problems that would otherwise result from division by $p(x)$. For more discussion, see [5].

When $\mathcal{D}$ is a discrete set of system states, densities are replaced by discrete distributions and integrals by sums in the usual way. In a more general setting for the discussion than we have chosen, $\mathcal{D}$ would be a measurable space, and p and q would be replaced by prior and posterior probability measures. By continuing to write in terms of probability

densities, we would then be implicitly assuming some underlying measure with respect to which the rest were absolutely continuous. Indeed such a measure certainly exists if we demand that no event with zero prior probability can have positive posterior probability, which in the present context we are in effect demanding by assuming (4).

III. PROPERTIES GIVEN GENERAL CONSTRAINTS

This Section concerns properties that apply in the case of both equality and inequality constraints (2)-(3). We follow the formal statement of each property with a brief discussion and then a proof or an appropriate reference. Throughout, we assume a system with possible states $\mathcal{D}$, probability density $q \in \mathcal{D}$, an arbitrary prior $p \in \mathcal{D}$, and arbitrary new information $I = (q \in \mathcal{A})$, where $\mathcal{A} \subseteq \mathcal{D}$ contains at least one density q such that $H(q, p) < \infty$.

Property 1 (Uniqueness): The posterior $q = p \circ I$ is unique.

Discussion: A solution to the cross-entropy minimization problem, if one exists, is unique provided only that $H[q, p]$ is not identically infinite as q ranges over the constraint set $\mathcal{A}$. To guarantee that a solution exists, a little more is required. One condition that suffices for existence is that, in addition to containing a density q with finite cross-entropy, the constraint set $\mathcal{A}$ be closed. (We call $\mathcal{A}$ closed if it contains every probability density q that is a limit of densities $q_i \in \mathcal{A}$. Limits are taken in the sense that $q_i \rightarrow q$ means $\int |q_i(x) - q(x)| dx \rightarrow 0$.) For $\mathcal{A}$ to be closed, it suffices in turn that the constraint functions be bounded. (And conversely, any closed convex set of probability densities can be defined by equality and inequality constraints (2), (3) with bounded constraint

functions, except that infinitely many may be required.) It is also possible to assert existence of $p \circ I$ under less stringent conditions, which do not imply that $\mathcal{A}$ is closed -- see Theorem 3.3 in [6] and Appendix A. This is fortunate, since a number of examples of practical importance involve unbounded constraint functions.

Proof of 1: See [6], [4, Section IV.E].

Property 2: The posterior satisfies $q = p \circ I = p$ if and only if the prior satisfies $p \in \mathcal{A}$.

Discussion: If one views cross-entropy minimization as an inference procedure, it makes sense that the posterior should be unchanged from the prior if the new information doesn't contradict the prior in any way. Consider the example of (A.10)-(A.12). If $a_k = \bar{x}_k$ for $k = 1, \dots, n$, then $q(\underline{x}) = p(\underline{x})$.

Proof of 2: Property 2 follows directly from the property of cross-entropy that $H[q, p] \geq 0$ with $H[q, p] = 0$ only if $q = p$ ([1, p. 14]).

Property 3 (idempotence): $(p \circ I) \circ I = p \circ I$.

Discussion: Taking the same information into account twice has the same effect as taking it into account once.

Proof of 3: Since $(p \circ I) \in \mathcal{A}$, idempotence follows from Property 2.

Property 4: Let I_1 be the information $I_1 = (q^+ \in \mathcal{A}_1)$ and let I_2 be the information $I_2 = (q^+ \in \mathcal{A}_2)$, for overlapping constraint sets $\mathcal{A}_1, \mathcal{A}_2 \subseteq \mathcal{D}$. If $(p \circ I_1) \in \mathcal{A}_2$ holds, then

$$p \circ I_1 = (p \circ I_1) \circ (I_1 \wedge I_2) = (p \circ I_1) \circ I_2 = p \circ (I_1 \wedge I_2) \quad (10)$$

holds.

Discussion: If the result of taking information I_1 into account already satisfies constraints imposed by additional information I_2 , taking I_2 into account in various ways has no effect. For example, let I_1 and I_2 be the constraints

$$\begin{aligned} \int_0^\infty dx x q^\dagger(x) &= a \\ \text{and} \quad \int_0^\infty dx x^2 q^\dagger(x) &= 2a^2, \end{aligned} \quad (11)$$

respectively. For an exponential prior $p(x) = r \exp(-rx)$, the posterior given I_1 is $q = p \circ I_1 = (1/a) \exp(-x/a)$ (see (A.10)-(A.12)). The second moment of q is just $2a^2$, so that q satisfies $q \in \mathcal{J}_2$, as well as $q = q \circ (I_1 \wedge I_2)$, $q = q \circ I_2$, and $q = p \circ (I_1 \wedge I_2)$. If the right side of (11) were anything but $2a^2$, the result of $p \circ (I_1 \wedge I_2)$ would be a truncated Gaussian or undefined and not an exponential [7, pp. 133-140]

Proof of 4: Since $(p \circ I_1) \in \mathcal{J}_1$ holds and, by assumption, $(p \circ I_1) \in \mathcal{J}_2$ also holds, it follows that $(p \circ I_1) \in (\mathcal{J}_1 \cap \mathcal{J}_2)$ holds. The first two equalities of (10) then follow directly from Properties 2 and 3. The last equality of (10) follows from $q = p \circ I_1$ having the smallest cross-entropy $H[q, p]$ of all densities in $\mathcal{J}_1$ and therefore in $\mathcal{J}_1 \cap \mathcal{J}_2$.

Property 5 (Invariance): Let $\tilde{\Gamma}$ be a coordinate transformation from $\underline{x} \in \mathcal{D}$ to $\underline{y} \in \mathcal{D}'$ with $(\tilde{\Gamma}q)(\underline{y}) = \tilde{J}^{-1} q(\underline{x})$, where $\tilde{J}$ is the Jacobian $\tilde{J} = \partial(\underline{y})/\partial(\underline{x})$. Let $\tilde{\Gamma}\mathcal{D}$ be the set of densities $\tilde{\Gamma}q$ corresponding to densities $q \in \mathcal{D}$. Let $(\tilde{\Gamma}\mathcal{J}) \subseteq (\tilde{\Gamma}\mathcal{D})$ correspond to $\mathcal{J} \subseteq \mathcal{D}$. Then

$$(\tilde{\Gamma}p) \circ (\tilde{\Gamma}I) = \tilde{\Gamma}(p \circ I) \quad (12)$$

and

$$H[\tilde{\Gamma}(p \circ I), \tilde{\Gamma} p] = H[p \circ I, p] \quad (13)$$

hold, where $\tilde{\Gamma} I = ((\tilde{\Gamma} q^*) \in (\tilde{\Gamma} \mathcal{A}))$.

Discussion: Eq. (12) states that the same answer is obtained when one solves the inference problem in two different coordinate systems, in that the posteriors in the two systems are related by the coordinate transformation. Moreover, the cross-entropy between the posteriors and the priors has the same value in both coordinate systems.

As an example, let y_1 and y_2 be the real and imaginary parts of a complex sinusoidal signal; let x_1 be the total power $x_1 = y_1^2 + y_2^2$, and let x_2 be the phase, so that

$$(y_1, y_2) = \tilde{\Gamma}(x_1, x_2) = (x_1^{1/2} \cos(x_2), x_1^{1/2} \sin(x_2)).$$

Then the Jacobian is constant:

$$\tilde{J} = \det \begin{bmatrix} \frac{1}{2} x_1^{-1/2} \cos(x_2) & -x_1^{1/2} \sin(x_2) \\ \frac{1}{2} x_1^{-1/2} \sin(x_2) & x_1^{1/2} \cos(x_2) \end{bmatrix} = 1/2.$$

Therefore, if the prior density $p(\underline{x})$ is uniform in some region in the $\underline{x}$ coordinate space, the transformed prior $(\tilde{\Gamma} p)(\underline{y})$ will be uniform on a corresponding region in the $\underline{y}$ coordinate space. For example, suppose

$$p(\underline{x}) = \begin{cases} 1/2\pi R^2, & (0 \leq x_1 \leq R^2, -\pi < x_2 \leq \pi) \\ 0, & \text{otherwise} \end{cases},$$

which makes p uniform in a certain rectangle. Thus, we find that

$$(\tilde{\Gamma} p)(\underline{y}) = \begin{cases} 1/\pi R^2, & (y_1^2 + y_2^2 \leq R^2) \\ 0, & \text{otherwise} \end{cases},$$

which makes Γp uniform on a certain disk. (Notice $1/\pi R^2 = \int_{-\pi}^{\pi} \int_0^R (1/2\pi R^2) r dr d\theta$.)

Now, let new information I specify the expected power

$$\int_0^{\infty} dx_1 \int_{-\pi}^{\pi} dx_2 x_1 q^{\dagger}(\underline{x}) = P.$$

The resulting posterior $q = p \circ I$ is exponential with respect to x_1 :

$$q(\underline{x}) = \begin{cases} A \exp[-\lambda x_1], & (0 \leq x_1 \leq R^2, -\pi < x_2 \leq \pi) \\ 0 & , \text{ otherwise} \end{cases}$$

for certain constants A and λ . The new information in the transformed coordinates, ΓI , is

$$\int dy_1 \int dy_2 (y_1^2 + y_2^2) q'^{\dagger}(\underline{y}) = P,$$

and the resulting posterior $q' = (\Gamma p) \circ (\Gamma I)$ has the form of a bivariate Gaussian inside the disk:

$$q'(\underline{y}) = \begin{cases} 2A \exp[-\lambda(y_1^2 + y_2^2)], & (y_1^2 + y_2^2 \leq R^2) \\ 0 & , \text{ otherwise} \end{cases}$$

The two posteriors q and q' are related by $q'(\underline{y}) = (\Gamma q)(\underline{y})$, as stated in (12).

Proof of 5: See [4, Section IV.E]. The proof of (12) follows directly from the fact that cross-entropy is transformation invariant. Eq. (13) is just a special case of this invariance.

Property 6 (System Independence): Let there be two systems, with sets $\mathcal{D}_1$ and $\mathcal{D}_2$ of states and probability densities of states $q_1 \in \mathcal{D}_1$ and $q_2 \in \mathcal{D}_2$. Let $p_1 \in \mathcal{D}_1$ and $p_2 \in \mathcal{D}_2$ be prior densities. Let $I_1 = (q_1^{\dagger} \in \mathcal{V}_1)$ and $I_2 = (q_2^{\dagger} \in \mathcal{V}_2)$ be new information about the two systems, where $\mathcal{V}_1 \subseteq \mathcal{D}_1$ and $\mathcal{V}_2 \subseteq \mathcal{D}_2$. Then

$$(p_1 p_2) \circ (I_1 \wedge I_2) = (p_1 \circ I_1)(p_2 \circ I_2) \quad (14)$$

and

$$H[q_1 q_2, p_1 p_2] = H[q_1, p_1] + H[q_2, p_2], \quad (15)$$

hold, where $q_1 = p \circ I_1$ and $q_2 = p \circ I_2$.

Discussion: Property 6 states that it doesn't matter whether one accounts for independent information about two systems separately or together in terms of a joint density. Whether or not the two systems are in fact independent is irrelevant: The property applies as long as there are independent priors and independent new information. Examples can easily be generated from the multivariate exponential and multivariate Gaussian examples in the Appendix.

Proof of 6: See [4, Section IV.E]

Property 7 (subset independence): Let $S_1, \dots, S_n$ be disjoint sets whose union is D . Let the new information I comprise information about the conditional densities $q^{\tau}_{S_i}$. Thus, $I = I_1 \wedge I_2 \wedge \dots \wedge I_n$, and $I_i = (q^{\tau}_{S_i} \in \mathcal{I}_i)$, where $\mathcal{I}_i \subseteq \mathcal{S}_i$ and $\mathcal{S}_i$ is the set of densities on S_i . Let $M = (q^{\tau} \in \mathcal{M})$ be new information giving the probability of being in each of the n subsets, where $\mathcal{M}$ is the set of densities q that satisfy

$$\int_{S_i} dx q(x) = m_i$$

for each subset S_i , where the m_i are known values. Then

$$(p \circ (I \wedge M)) \circ S_i = (p \circ S_i) \circ I_i \quad (16)$$

and

$$H[p \circ (I \wedge M), p] = \sum_i m_i H[q_i, p_i] + \sum_i m_i \log \frac{m_i}{s_i} \quad (17)$$

hold, where $p_i = p^*S_i$, $q_i = p_i \circ I_i$, and the s_i are the prior probabilities of being in each subset,

$$s_i = \int_{S_i} dx p(x) \quad . \quad (18)$$

Discussion: This property concerns situations in which the set of states D decomposes naturally into disjoint subsets S_i , in which the new information $I = I_1 \wedge I_2 \wedge \dots \wedge I_n$ comprises disjoint information about the conditional probability densities q^*S_i in each subset, and in which there is also new information M giving the total probability m_i of being in each subset S_i . Given this information, there are two ways to obtain posterior conditional densities for each subset: One way is to obtain a conditional posterior $(p^*S_i) \circ I_i$ from each conditional prior p^*S_i . Another way is to obtain a posterior $q = p \circ (I \wedge M)$ for the whole system and then to compute a conditional posterior q^*S_i . Property 7 states that the results are the same in both cases: it doesn't matter whether one treats an independent subset of system states in terms of a separate conditional density or in terms of the full system density.

To illustrate Property 7, suppose that a six-sided die was rolled a large number of times. The frequencies with which the different die faces turned up were not recorded individually, but the mean number of spots showing was determined separately for the odd results and for the even results. There was no prior reason to expect any face of the die to turn up more often than any other. Indeed, the probability for the number of spots showing to be odd turned out to be .5. However, the mean number of spots showing, given that the number was odd, was found to be 4; the mean number of spots showing, given

that the number was even, also was found to be 4. Given this information, we are asked to estimate the probability for each face of the die to turn up, as well as the conditional probability given whether the face is odd or even.

Let $S_1 = \{1, 3, 5\}$ and $S_2 = \{2, 4, 6\}$. We will first solve the problem on S_1 and S_2 separately and then solve it on $S_1 \cup S_2$.

In all cases, the prior is uniform. The prior p_1 on S_1 is $p_1(1) = p_1(3) = p_1(5) = 1/3$. The information I_1 giving the expected value for an odd number of spots is

$$\sum_{n \in S_1} n q_1(n) = 4 ;$$

therefore, we compute a posterior $q_1 = p_1 \circ I_1$ on S_1 by minimizing $H[q_1, p_1]$ subject to $q_1(1) + 3q_1(3) + 5q_1(5) = 4$. The result is

$$q_1(1) = 0.1162, \quad q_1(3) = 0.2676, \quad q_1(5) = 0.6162 . \quad (19)$$

Similarly, the prior p_2 on S_2 is $p_2(2) = p_2(4) = p_2(6) = 1/3$, the posterior q_2 is subject to the constraint I_2 ,

$$2q_2(2) + 4q_2(4) + 6q_2(6) = 4,$$

and the result of minimizing $H[q_2, p_2]$ is

$$q_2(2) = 1/3, \quad q_2(4) = 1/3, \quad q_2(6) = 1/3. \quad (20)$$

On $S_1 \cup S_2$, the prior p is $p(1) = p(2) = \dots = p(6) = 1/6$. The information I_1 , which concerns $q^\dagger * S_1$, may be expressed as

$$q^\dagger(1) + 3q^\dagger(3) + 5q^\dagger(5) = 4(q^\dagger(1) + q^\dagger(3) + q^\dagger(5)).$$

We therefore subject the posterior q to the constraint

$$- 3q(1) - q(3) + q(5) = 0. \quad (21)$$

Similarly, because of I_2 , we have the constraint

$$- 2q(2) + 2q(6) = 0. \quad (22)$$

Finally, because of the information M , we subject q to the constraint

$$q(1) - q(2) + q(3) - q(4) + q(5) - q(6) = 0, \quad (23)$$

since this is equivalent to $q(1) + q(3) + q(5) = .5 = q(2) + q(4) + q(6)$.

Upon minimizing $H[q, p]$ subject to the constraints (21)-(23), we find that

$q = p \circ (I_1 \wedge I_2 \wedge M)$ is given by

$$\begin{aligned} q(1) &= 0.0581, & q(2) &= 1/6, \\ q(3) &= 0.1338, & q(4) &= 1/6, \\ q(5) &= 0.3081, & q(6) &= 1/6. \end{aligned} \quad (24)$$

To find the conditional probabilities q^*s_1 and q^*s_2 , we divide both columns in this result by .5; the results agree with q_1 and q_2 as computed above ((19),(20)), and as stated in (16).

Proof of 7: See [4, Section IV.E].

Property 8 (weak subset independence): For the same definitions and notation as Property 7,

$$(p \circ I)^*s_i = (p^*s_i) \circ I_i \quad (25)$$

and

$$H[p \circ I, p] = \sum_i r_i H[q_i, p_i] + \sum_i r_i \log \frac{r_i}{s_i} \quad (26)$$

hold, where $p_i = p^*s_i$, $q_i = p_i \circ I_i$, the s_i are the prior

probabilities of being in each subset (18), and the r_i are the posterior probabilities of being in each subset,

$$r_i = \int_{S_i} dx q(x) , \quad (27)$$

for $q = p \circ I$.

Discussion: This property states that the two ways of obtaining the posterior conditional densities also lead to the same result in the case when one does not have information giving the total probability in each subset. Results for the full system posterior, however, will not in general be the same for the cases covered by Properties 7 and 8. That is, $q \circ I$ and $q \circ (I \wedge M)$ will not in general be equal.

To illustrate Property 8, we solve the example problem from Property 7, omitting the information M that the probability of an odd (or of an even) number of spots is .5. The separate solutions on S_1 and S_2 proceed exactly as before and yield the same posteriors q_1 and q_2 . The solution on $S_1 \cup S_2$ differs from the previous one only in that we minimize $H[q, p]$ subject to the constraints (21) and (22), but not subject to (23). The result, $q' = p \circ (I_1 \wedge I_2)$, is given by

$$\begin{aligned} q'(1) &= 0.0524, & q'(2) &= 0.1831, \\ q'(3) &= 0.1206, & q'(4) &= 0.1831, \\ q'(5) &= 0.2778, & q'(6) &= 0.1831, \end{aligned}$$

and differs from the previous result (24). Moreover, the subset probabilities r_1 and r_2 do not satisfy M : summing the two columns gives $r_1 = 0.4508$ and $r_2 = 0.5492$. However, dividing the two columns respectively by r_1 and

r_2 gives the same conditional probabilities as before: $q' * \underline{s}_1 = q_1$ and $q' * \underline{s}_2 = q_2$ (see (19), (20)).

Proof of 8: For $q = p \circ I$, let r_i be given by (27). Then let R be information $R = q^T \in \mathcal{R}$, where $\mathcal{R}$ is the set of densities satisfying (27). It follows from Property 4 that $p \circ I = p \circ (I \circ R)$ holds; (25) and (26) then follow from Property 7.

Property 9 (subset aggregation): Let $\underline{s}_1, \underline{s}_2, \dots, \underline{s}_n$ be disjoint sets whose union is $\underline{D}$. Let $\underline{\psi}$ be a transformation such that, for any $q \in \underline{\mathcal{D}}$, $q' = \underline{\psi} q$ is a discrete distribution with

$$q'(x_i) = \int_{\underline{s}_i} d\underline{x} q(\underline{x}),$$

where x_i is a discrete state corresponding to $\underline{x} \in \underline{s}_i$. Thus the transformation $\underline{\psi}$ aggregates the states in each subset $\underline{s}_i$. Suppose new information $I' = ((\underline{\psi} q^T) \in \underline{\mathcal{J}}')$ is obtained about the aggregate distribution $\underline{\psi} q^T$, where $\underline{\mathcal{J}}'$ is a convex set of discrete distributions. Then for any prior $p \in \underline{\mathcal{D}}$,

$$p * \underline{s}_i = (p \circ I) * \underline{s}_i, \quad (28)$$

$$(\underline{\psi} p) \circ I' = \underline{\psi}(p \circ I), \quad (29)$$

and

$$H[\underline{\psi}(p \circ I), \underline{\psi} p] = H[p \circ I, p] \quad (30)$$

all hold, where $I = \underline{\psi}^{-1} I'$ is the information I' expressed in terms of q^T instead of in terms of $\underline{\psi} q^T$. (That is, $I = (q^T \in (\underline{\psi}^{-1} \underline{\mathcal{J}}'))$, where $(\underline{\psi}^{-1} \underline{\mathcal{J}}') \subseteq \underline{\mathcal{D}}$ are the densities q such that $(\underline{\psi} q) \in \underline{\mathcal{J}}'$.)

Discussion: Note that (29) and (30), in which ψ is a many-to-one mapping, have the same form as the invariance property, which holds for one-to-one coordinate transformations $\tilde{\Gamma}$ (see (12)-(13)). Indeed, both invariance and subset aggregation can be viewed as special cases of a more general, measure-theoretic invariance. In mathematical terms, the operator $\circ$ is functorial.

Proof of 9: Let the information I' be a set of known expectations

$\sum_i g_{ki} q^{\dagger'}(x_i)$, for $k = 1, \dots, m$, or bounds on these expectations, where $q^{\dagger'} = \psi q^{\dagger}$. In terms of $q^{\dagger}$, this becomes a set of known or bounded expectations

$$\int_{\mathcal{D}} dx \, q^{\dagger}(x) f_k(x),$$

where $f_k(x \in S_i) = g_{ki}$ is constant in each subset S_i . The posterior $q = p \circ I$ has the form

$$q(x) = p(x) \exp \left(-\lambda_0 - \sum_{k=1}^m \lambda_k f_k(x) \right), \quad (31)$$

where some of the terms in the summation over k may be omitted in the case of inequality constraints (see (A.4)). Since f_k is constant on each subset, (31) has the form $q(x \in S_i) = A_i p(x \in S_i)$, where A_i is a subset dependent constant. This proves (28). Now, in general for any $q, p \in \mathcal{D}$, the cross-entropy $H[q, p]$ can be expressed [4] as

$$H[q, p] = \sum_i r_i H[q_i, p_i] + \sum_i r_i \log \frac{r_i}{s_i} \quad (32)$$

where $p_i = p|_{S_i}$, $q_i = q|_{S_i}$,

$$s_i = \int_{S_i} dx \, p(x), \text{ and } r_i = \int_{S_i} dx \, q(x).$$

In the present case we have $q_i = p_i$ from (28). Since $H[q_i, q_i] = 0$, (32) reduces to

$$\begin{aligned} H[q, p] &= \sum_i r_i \log \frac{r_i}{s_i} \\ &= H[\psi q, \psi p] . \end{aligned}$$

Minimizing the left side subject to I , yielding $q = p \circ I$, is equivalent to minimizing the right side subject to I' . This proves (29) and (30).

Property 10 (triangle relations): For any $r \in \mathcal{V}$,

$$H[r, p] \geq H[r, q] + H[q, p] , \quad (33)$$

where $q = p \circ I$. When I is determined by a finite set of equality constraints only, equality holds in (33).

Discussion: The triangle equality is important for applications in which cross-entropy minimization is used for purposes of classification and pattern recognition.

Proof of 10: We have

$$H[q, p] = \min_{q' \in \mathcal{V}} H[q', p] .$$

The densities $q' = (1-t)q + tr$ belong to $\mathcal{V}$ for all $t \in [0, 1]$ since $q \in \mathcal{V}$, $r \in \mathcal{V}$, and $\mathcal{V}$ is convex. For all such t we therefore have

$$H[(1-t)q + tr, p] \geq H[q, p] , \quad (34)$$

or $F(t) \geq F(0)$, where we have written $F(t)$ for the left side of (34). It follows that $F'(0) \geq 0$ (provided F is differentiable at 0). We therefore set

$$\frac{d}{dt} \left(\int d\tilde{x} [(1-t)q(\tilde{x}) + tr(\tilde{x})] \log \frac{(1-t)q(\tilde{x}) + tr(\tilde{x})}{p(\tilde{x})} \right) \Big|_{t=0} \geq 0$$

and differentiate under the integral sign. (For justification of this step and the existence of $F'(0)$, see Csiszár [6], who gives the proof in a more general measure-theoretic setting.) The result is

$$\left(\int d\tilde{x} \left\{ [r(\tilde{x}) - q(\tilde{x})] \log \frac{(1-t)q(\tilde{x}) + tr(\tilde{x})}{p(\tilde{x})} + [(1-t)q(\tilde{x}) + tr(\tilde{x})] \frac{r(\tilde{x}) - q(\tilde{x})}{(1-t)q(\tilde{x}) + tr(\tilde{x})} \right\} \right) \Big|_{t=0} \geq 0 ,$$

or

$$\int d\tilde{x} [r(\tilde{x}) - q(\tilde{x})] [1 + \log \frac{q(\tilde{x})}{p(\tilde{x})}] \geq 0 .$$

This implies

$$\int d\tilde{x} r(\tilde{x}) \log \frac{q(\tilde{x})}{p(\tilde{x})} \geq \int d\tilde{x} q(\tilde{x}) \log \frac{q(\tilde{x})}{p(\tilde{x})}$$

since $\int d\tilde{x} [r(\tilde{x}) - q(\tilde{x})] = 0$. Therefore,

$$\int d\tilde{x} r(\tilde{x}) \log \frac{r(\tilde{x})}{p(\tilde{x})} \geq \int d\tilde{x} r(\tilde{x}) \log \frac{r(\tilde{x})}{q(\tilde{x})} + \int d\tilde{x} q(\tilde{x}) \log \frac{q(\tilde{x})}{p(\tilde{x})} .$$

Consequently $H[r, p] \geq H[r, q] + H[q, p]$.

Now assume I is determined by finitely many equality constraints. Since $q = p \circ I$, $\log(q(\tilde{x})/p(\tilde{x}))$ assumes the form

$$\log \frac{q(\tilde{x})}{p(\tilde{x})} = -\lambda_0 - \sum_{k=1}^m \lambda_k f_k(\tilde{x})$$

(cf. (A.4)). But then

$$\int d\tilde{x} r(\tilde{x}) \log \frac{q(\tilde{x})}{p(\tilde{x})} = -\lambda_0 - \sum_{k=1}^m \lambda_k \bar{f}_k = \int d\tilde{x} q(\tilde{x}) \log \frac{q(\tilde{x})}{p(\tilde{x})} = H[q, p] ,$$

since r and q both satisfy the equality constraints. The equality

$$\int_{\mathcal{X}} r(\underline{x}) \log \frac{r(\underline{x})}{p(\underline{x})} d\underline{x} = \int_{\mathcal{X}} r(\underline{x}) \log \frac{r(\underline{x})}{q(\underline{x})} d\underline{x} + \int_{\mathcal{X}} r(\underline{x}) \log \frac{q(\underline{x})}{p(\underline{x})} d\underline{x}$$

then implies $H[r, p] = H[r, q] + H[q, p]$.

Property 11:

$$H[q^\dagger, p \circ I] \leq H[q^\dagger, p], \quad (35)$$

holds with equality if and only if $p \circ I = p$.

Discussion: This property states that the posterior $q = p \circ I$ is always closer to $q^\dagger$, in the cross-entropy sense, than is the prior p .

Proof of 11: Since $q^\dagger \in \mathcal{D}$ holds, (35) follows directly from (33) with $r = q^\dagger$.

IV. PROPERTIES GIVEN EQUALITY CONSTRAINTS

This Section concerns properties that apply when some of the new information is in the form of equality constraints (2) only. Throughout, we assume a system with possible states $\mathcal{D}$ and an arbitrary prior $p \in \mathcal{D}$.

Property 12. Let the system have a probability density $q^\dagger \in \mathcal{D}$, and let there be information $I = (q^\dagger \in \mathcal{D})$ that is determined by a finite set of equality constraints only. Then

$$H[q^\dagger, p] = H[q^\dagger, q] + H[q, p] \quad (36)$$

holds, where $q = p \circ I$.

Discussion: Since the difference $H[q^\dagger, p] - H[q^\dagger, q]$ is just $H[q, p]$, and since $H[q, p]$ is a measure [1] of the information divergence between q and p , Property 12 shows that $H[p \circ I, p]$ can be interpreted as the amount of information provided by I that was not already inherent in p . Stated differently,

$H[p \circ I, p]$ is the amount of information-theoretic distortion introduced if p is used instead of $p \circ I$. Since, for any prior p and any density $r \in \mathcal{D}$ with $H(r, p) < \infty$, there exists a finite set of equality constraints I_r such that $r = p \circ I_r$ (see appendix B), $H[r, p]$ is in general the amount of information needed to determine r when given p , or the amount of information-theoretic distortion introduced if r is used instead of p .

Proof of 12: Eq. (36) follows directly from (33), since $q^{\dagger} \in \mathcal{J}$ holds.

Property 13: Let the system have a probability density $q^{\dagger} \in \mathcal{D}$, and let there be information $I_1 = (q^{\dagger} \in \mathcal{J}_1)$ and information $I_2 = (q^{\dagger} \in \mathcal{J}_2)$, where $\mathcal{J}_1, \mathcal{J}_2 \subseteq \mathcal{D}$ are constraint sets with a non-empty intersection. Suppose that $\mathcal{J}_1$ is determined by a set of equality constraints (2) only. Then

$$(p \circ I_1) \circ (I_1 \wedge I_2) = p \circ (I_1 \wedge I_2) \quad (37)$$

and

$$H[q, p] = H[q, q_1] + H[q_1, p] \quad (38)$$

hold, where $q = p \circ (I_1 \wedge I_2)$ and $q_1 = p \circ I_1$.

Discussion: When I_1 is determined by equality constraints, (37) holds whether or not $(p \circ I_1) \in \mathcal{J}_2$ (compare with Property 4). Property 13 is important for applications in which constraint information arrives piecemeal, and states that intermediate posteriors can be used as priors in computing final posteriors without affecting the results. Like (33) and (36), the triangle equality (38) is important for applications in which cross-entropy minimization is used for purposes of classification and pattern recognition.

As an example of Property 13, we consider minimum cross-entropy spectral analysis [8]. If one describes a stochastic, band-limited, discrete-spectrum

signal in terms of a probability density $q^{\dagger}(\underline{x}) = q^{\dagger}(x_1, \dots, x_n)$, where x_k is the energy at frequency f_k , known values of the autocorrelation function can be expressed as expectations of $q^{\dagger}$, namely,

$$R_r = \int d\underline{x} \left(\sum_k 2x_k \cos(2\pi t_r f_k) \right) q^{\dagger}(\underline{x}),$$

where R_r is the autocorrelation sample at lag t_r . Let I_1 be a limited set of autocorrelations $R_1, \dots, R_m$. Then, for a prior p_W with a flat (white) power spectrum $P_k = \int d\underline{x} x_k p_W(\underline{x}) = P$, the power spectrum of the posterior $q_{LPC} = p_W^{\circ} I_1$ is just the m th-order maximum entropy or Linear Predictive Coding (LPC) spectrum [8]. Let I_2 be the set of autocorrelation samples $R_{m+1}, R_{m+2}, \dots$ that together with I_1 fully determine the power spectrum of $q^{\dagger}$. Then (37) yields $q_F = p_W^{\circ}(I_1 \wedge I_2) = q_{LPC}^{\circ}(I_1 \wedge I_2)$.

Proof of 13: The density q_1 has the form (A.4),

$$q_1(\underline{x}) = p(\underline{x}) \exp \left(-\lambda_0 - \sum_{k=1}^m \lambda_k a_k(\underline{x}) \right).$$

For an arbitrary density $q \in \mathcal{D}$, the cross-entropy with respect to q_1 satisfies

$$\begin{aligned} H[q, q_1] &= \int d\underline{x} q(\underline{x}) \log \frac{q(\underline{x}) \exp[\lambda_0 + \sum_k \lambda_k a_k(\underline{x})]}{p(\underline{x})} \\ &= H[q, p] + \lambda_0 + \int d\underline{x} q(\underline{x}) \sum_k \lambda_k a_k(\underline{x}) \end{aligned}$$

If q satisfies $q \in \mathcal{V}_1$, this becomes

$$H[q, q_1] = H[q, p] + \lambda_0 + \sum_k \lambda_k \bar{a}_k, \quad (39)$$

where λ_0 , λ_k , and $\bar{a}_k$ are constants. Since $H[q, q_1]$ and $H[q, p]$ differ by a constant on $\mathcal{V}_1$, it follows that they have the same minima on any subset

of $\mathcal{I}_1$. Since $(\mathcal{I}_1 \cap \mathcal{I}_2) \subseteq \mathcal{I}_1$ holds, this proves (37). Moreover, (39) and (A.5) yield (38), which is also a special case of (33).

Property 14. Suppose there are two underlying probability densities $q_1^\dagger$ and $q_2^\dagger$. Let I_1 and I_2 stand respectively for the sets of equality constraints

$$\int_{\mathcal{X}} f_i(\mathbf{x}) q_1^\dagger(\mathbf{x}) = F_i^{(1)} \quad (i = 1, \dots, m) \quad (40)$$

and

$$\int_{\mathcal{X}} f_i(\mathbf{x}) q_2^\dagger(\mathbf{x}) = F_i^{(2)} \quad (i = 1, \dots, s), \quad (41)$$

where $s \geq m$. Then

$$(p \circ I_1) \circ (I_2) = p \circ I_2 \quad (42)$$

holds. Moreover, if $\lambda_k^{(1)}$, $\lambda_k^{(12)}$, and $\lambda_k^{(2)}$ are the Lagrangian multipliers associated with $q_1 = p \circ I_1$, $q_{12} = q_1 \circ I_2$, and $q_2 = p \circ I_2$, respectively, then

$$\lambda_k^{(2)} = \lambda_k^{(1)} + \lambda_k^{(12)} \quad (k = 0, 1, \dots, m), \quad (43)$$

$$\lambda_k^{(2)} = \lambda_k^{(12)} \quad (k = m+1, \dots, s), \quad (44)$$

and

$$H[q_2, p] = H[q_2, q_1] + H[q_1, p] + \sum_{r=1}^m \lambda_r^{(1)} (F_r^{(1)} - F_r^{(2)}) \quad (45)$$

also hold.

Discussion: Property 10 can apply to situations in which $q_1^\dagger$ and $q_2^\dagger$ are system probability densities at different times and in which $q_1^\dagger$ or estimates of $q_1^\dagger$ are considered to be good estimates of $q_2^\dagger$. If I_2 is determined in

part by expectations of the same functions as I_1 , but with different expected values, then the results of taking I_1 into account are completely wiped out by subsequently taking I_2 into account. As an example, consider frame-by-frame minimum cross-entropy spectral analysis in which I_i is determined by autocorrelation samples in frame i at a fixed set of lags ($s = m$). Eq. (42) shows that the results for frame i are the same whether the assumed prior is an original prior p , the posterior from frame $i-1$, or some intermediate estimate. (However, there may be computational or bandwidth-reduction advantages to using $p \circ I_{i-1}$ as a prior in frame i .) Note that, if $s \geq m$ and $F_r^{(1)} = F_r^{(2)}$ for $r = 1, \dots, m$, Property 14 reduces to Property 13.

Proof of 14: From (A.4) we have

$$q_1(\underline{x}) = p(\underline{x}) \exp \left(-\lambda_0^{(1)} - \sum_{k=1}^m \lambda_k^{(1)} a_k(\underline{x}) \right),$$

where the $\lambda_k^{(1)}$ are chosen to satisfy the constraints (40). Similarly,

$$q_{12}(\underline{x}) = q_1(\underline{x}) \exp \left(-\lambda_0^{(12)} - \sum_{k=1}^s \lambda_k^{(12)} a_k(\underline{x}) \right),$$

holds. This is of the form $p(\underline{x}) \exp[-\lambda_0^{(2)} - \sum_k \lambda_k^{(2)} a_k(\underline{x})]$,

with $\lambda_k^{(2)} = \lambda_k^{(1)} + \lambda_k^{(12)}$ ($k = 1, \dots, m$) and $\lambda_k^{(2)} = \lambda_k^{(12)}$

($k = m+1, \dots, s$), and it is a probability density satisfying the constraints

(41); it is therefore equal to $p \circ I_2 = q_2$, which proves (43)-(44). Eq.

(45) follows from straightforward applications of (A.5).

Property 15 (expected value matching): Let I be the constraints

$$\int_{\mathcal{D}} d\mathbf{x} \, q^{\dagger}(\mathbf{x}) f_k(\mathbf{x}) = \bar{f}_k \quad (k = 1, \dots, m) \quad (46)$$

for a fixed set of functions f_k , and let $q = p \circ I$ be the result of taking this information into account. Then, for an arbitrary fixed density $q^* \in \mathcal{D}$, the cross entropy $H[q^*, q] = H[q^*, p \circ I]$ has a minimum value when the constraints (46) satisfy

$$\bar{f}_k = \bar{f}_k^* = \int_{\mathcal{D}} d\mathbf{x} \, q^*(\mathbf{x}) f_k(\mathbf{x}).$$

Discussion: This property states that, for a density q of the general form (A.4), $H[q^*, q]$ is smallest when the expectations of q match those of q^* . In particular, note that $q = p \circ I$ is not only the density that minimizes $H[q, p]$, but also is the density of the form (A.4) that minimizes $H[q^{\dagger}, q]$! Property 15 is a generalization of a property of orthogonal polynomials [10] that, in the case of speech analysis, is called the "correlation matching property" [9, Chapter 2].

Proof of 15: The cross-entropy $H[q^*, q]$ is given by

$$\begin{aligned} H[q^*, q] &= \int_{\mathcal{D}} d\mathbf{x} \, q^*(\mathbf{x}) \log(q^*(\mathbf{x})/q(\mathbf{x})) \\ &= \int_{\mathcal{D}} d\mathbf{x} \, q^*(\mathbf{x}) \log(q^*(\mathbf{x})/p(\mathbf{x})) + \int_{\mathcal{D}} d\mathbf{x} \, q^*(\mathbf{x}) (\lambda_0 + \sum_k \lambda_k f_k(\mathbf{x})) \\ &= \int_{\mathcal{D}} d\mathbf{x} \, q^*(\mathbf{x}) \log(q^*(\mathbf{x})/p(\mathbf{x})) + \lambda_0 + \sum_k \lambda_k \bar{f}_k^*, \end{aligned} \quad (47)$$

where we have used (A.4). Now, since the multipliers λ_k are functions of the expected values $\bar{f}_k$, variations in the expected values are equivalent to variations in the multipliers. Hence, to find the minimum of $H[q^*, q]$, we solve

$$\frac{\partial}{\partial \lambda_k} H[q^*, q] = 0 = \frac{\partial \lambda_0}{\partial \lambda_k} + \bar{f}_k^*,$$

where we have used (47). It follows from (A.9) that the minimum occurs when $\bar{f}_k = \bar{f}_k^*$.

V. GENERAL DISCUSSION

Property 1 and Eqs. (12), (14), and (16) are the inference axioms on which the derivation in [4] is based. It is important to recognize that it is these inference properties, and not the corresponding cross-entropy properties (Eqs. (13), (15), and (17)) that characterize cross-entropy minimization. For more information on this distinction, see [4, Section VI] and [5].

An interesting aspect of the results presented in this paper is the interplay between properties of cross-entropy minimization as an inference procedure and properties of cross-entropy as an information measure. Cross-entropy's well-known [1] and unique [5] properties as an information measure in the case of arbitrary probability densities are extended and strengthened when one of the densities involved is the result of cross-entropy minimization. (See the statement and discussion of Properties 10, 11, 12, 13, and 15.) Indeed, the resulting combined properties have led to a new information-theoretic method of pattern analysis and classification [11] that is a refinement of a method due to Kullback [1, p. 83].

VI. ACKNOWLEDGEMENTS

We thank Jacob Feldman for suggesting Property 9; we thank Bernard O. Koopman and a referee for drawing our attention to technical issues concerning the existence of minimum cross-entropy posteriors.

APPENDIX A

Mathematics of Cross-Entropy Minimization

We derive the general solution for cross-entropy minimization given arbitrary constraints, and we illustrate the result with the important cases of exponential and Gaussian densities. In general, however, it is difficult or impossible to obtain a closed-form, analytic solution expressed directly in terms of the known expected values rather than in terms of the Lagrangian multipliers. We therefore discuss a numerical technique for obtaining the solution, namely the Newton-Raphson method. This method is the basis for a computer program that solves for the minimum cross-entropy posterior given an arbitrary prior and arbitrary expected-value constraints.

Given a positive prior density p and a finite set of equality constraints

$$\int q(\underline{x}) d\underline{x} = 1 , \quad (A.1)$$

$$\int f_k(\underline{x}) q(\underline{x}) d\underline{x} = \bar{f}_k , \quad (k = 1, \dots, m) , \quad (A.2)$$

we wish to find a density q that minimizes

$$H[q, p] = \int q(\underline{x}) \log \frac{q(\underline{x})}{p(\underline{x})} d\underline{x} ,$$

subject to the constraints. For conditions that imply the existence of a unique minimum, see the discussion of Property 1 (uniqueness). One standard method for seeking the minimum is to introduce Lagrangian multipliers β and λ_k ($k = 1, \dots, m$) corresponding to the constraints, forming the expression

$$\int q(\underline{x}) \log \frac{q(\underline{x})}{p(\underline{x})} d\underline{x} + \beta \int q(\underline{x}) d\underline{x} + \sum_{k=1}^m \lambda_k \int f_k(\underline{x}) q(\underline{x}) d\underline{x} ,$$

and to equate the variation, with respect to q , of this quantity to zero:

$$\log \frac{q(\underline{x})}{p(\underline{x})} + 1 + \beta + \sum_{k=1}^m \lambda_k f_k(\underline{x}) = 0. \quad (\text{A.3})$$

Solving for q leads to

$$q(\underline{x}) = p(\underline{x}) \exp \left(-\lambda_0 - \sum_{k=1}^m \lambda_k f_k(\underline{x}) \right), \quad (\text{A.4})$$

where we have introduced $\lambda_0 = \beta + 1$.

In fact, the q , if it exists, that minimizes $H[q,p]$ has this form with the possible exception of a set $\underline{S}$ of points on which the constraints imply that q vanishes. (Such a situation would arise, for instance, if we had a constraint $\int q(\underline{x})f(\underline{x})d\underline{x} = 0$, where $f(\underline{x}) > 0$ when $\underline{x} \in \underline{S}$ and $f(\underline{x}) = 0$ when $\underline{x} \notin \underline{S}$.)

Informally, we could then imagine the Lagrangian multipliers becoming infinite in such a way that the argument of $\exp$ in (A.4) becomes $-\infty$ when $\underline{x} \in \underline{S}$.)

Conversely, if a density q is found that is of this form and satisfies the constraints, then the minimum-cross-entropy density exists and equals q [6], [1]. For simplicity in the following, we assume the set $\underline{S}$ is empty.

The cross-entropy at the minimum can be expressed in terms of the λ_k and the $\bar{f}_k$ by multiplying (A.3) by $q(\underline{x})$ and integrating. The result is

$$H[q,p] = -\lambda_0 - \sum_{k=1}^m \lambda_k \bar{f}_k. \quad (\text{A.5})$$

It is necessary to choose λ_0 and the λ_k so that the constraints are satisfied. In the presence of the constraint (A.1) we may rewrite the remaining constraints in the form

$$\int (f_k(\underline{x}) - \bar{f}_k) q(\underline{x}) d\underline{x} = 0 \quad (\text{A.6})$$

Now, if we find values for the λ_k such that

$$\int (f_i(\underline{x}) - \bar{f}_i) p(\underline{x}) \exp \left(- \sum_{k=1}^m \lambda_k f_k(\underline{x}) \right) d\underline{x} = 0, \quad (i = 1, \dots, m), \quad (A.7)$$

we are assured of satisfying (A.6); and we can then satisfy (A.1) by setting

$$\lambda_0 = \log \int p(\underline{x}) \exp \left(- \sum_{k=1}^m \lambda_k f_k(\underline{x}) \right) d\underline{x}. \quad (A.8)$$

If the integral in (A.8) can be performed, one can sometimes find values for the λ_k from the relations

$$-\frac{\partial}{\partial \lambda_k} \lambda_0 = \bar{f}_k. \quad (A.9)$$

The situation for inequality constraints is only slightly more complicated. Suppose we replace all the equal signs in (A.2) by $\leq$. (We lose no generality thereby: we can change inequalities with $\geq$ into inequalities with $\leq$ by changing the signs of the corresponding f_k and $\bar{f}_k$, and any equality constraint is equivalent to a pair of inequality constraints.) The q that minimizes $H(q, p)$ subject to the resulting constraints will in general satisfy equality for certain values of k in the modified (A.2), while strict inequality will hold for the rest. We can still use the solution (A.4), subjecting the Lagrange multipliers to the conditions $\lambda_k \leq 0$ for k such that equality holds in the constraint, and $\lambda_k = 0$ for k such that strict inequality holds in the constraint.

It unfortunately is usually impossible to solve (A.7) or (A.9) for the λ_k explicitly, in closed form; however, it is possible in certain important special cases. For example, consider the case in which the prior $p(\underline{x})$ is a multivariate exponential,

$$p(\underline{x}) = \prod_{k=1}^n (1/a_k) \exp[-x_k/a_k] , \quad (\text{A.10})$$

where $\underline{x} = (x_1, \dots, x_n)$ and the x_k each range over the positive real line, and in which the constraints are

$$\int \underline{dx} \, x_k q(\underline{x}) = \bar{x}_k , \quad (\text{A.11})$$

$k = 1, \dots, n$. Solving (A.9) in order to express the minimum cross-entropy posterior directly in terms of the known expected values $\bar{x}_k$ yields

$$q(\underline{x}) = \prod_k (1/\bar{x}_k) \exp[-x_k/\bar{x}_k] . \quad (\text{A.12})$$

Thus, the density remains multivariate exponential, with the prior mean values a_k being replaced by the newly learned values $\bar{x}_k$.

Now consider the case in which the x_k range over the entire real line, and in which the prior density is Gaussian,

$$p(\underline{x}) = \prod_k (2\pi b_k)^{-1/2} \exp[-(x_k - a_k)^2/2b_k] .$$

Suppose that the constraints are (A.11) and

$$\int \underline{dx} \, (x_k - \bar{x}_k)^2 q(\underline{x}) = v_k .$$

In this case the minimum cross-entropy posterior is

$$q(\underline{x}) = \prod_k (2\pi v_k)^{-1/2} \exp[-(x_k - \bar{x}_k)^2/2v_k] .$$

Thus, the density remains multivariate Gaussian, with the prior means and variances being replaced by the newly learned values.

Here is an example of a simple problem for which the solution of (A.7) cannot be expressed in closed form. Consider a discrete system with n states x_j and prior probabilities $p(x_j) = p_j$ ($j = 1, \dots, n$). The discrete form of (A.1) is

$$\sum_{j=1}^n q_j = 1, \quad (\text{A.13})$$

where $q_j = q(x_j)$. Suppose the only other constraint is that the mean m of the indices j is prescribed: $f(x_j) = j$, and

$$\sum_{j=1}^n j q_j = m. \quad (\text{A.14})$$

Then (A.4) becomes $q_j = p_j \exp[-\lambda_0 - \lambda j]$, which we write as $q_j = a p_j z^j$ by introducing the abbreviations $a = \exp[-\lambda_0]$ and $z = \exp[-\lambda]$. From (A.16) and (A.17) we then obtain

$$a = \left(\sum_{j=1}^n p_j z^j \right)^{-1}$$

and

$$\sum_{j=1}^n (j-m) p_j z^j = 0. \quad (\text{A.15})$$

The problem then reduces to finding a positive root of the polynomial in (A.15). As in the continuous case, there are special forms for the prior that lead to important particular solutions. But when $n > 5$, the roots of the polynomial (other than zero) cannot in general be written as explicit, closed-form expressions in the coefficients for arbitrary priors. Numerical methods of solution therefore become important. Our obtaining a polynomial

equation in the present example was an accidental consequence of the fact that the values of the constraint function f formed a subset of an arithmetic progression ($j = 1, 2, \dots$). Thus, for more general types of problems, numerical methods are even more important.

One such method is the Newton-Raphson method, which is for finding solutions for systems of equations that, like (A.7), are of the form

$$F_i(\lambda_1, \dots, \lambda_m) = 0, \quad (i = 1, \dots, m). \quad (\text{A.16})$$

The method starts with an initial guess at the solution, $\lambda^{(1)} = (\lambda_1^{(1)}, \dots, \lambda_m^{(1)})$, and produces further approximate solutions $\lambda^{(2)}, \lambda^{(3)}, \dots$ in succession. If the initial guess $\lambda^{(1)}$ is close enough to a solution of (A.16), if the F_i are continuously differentiable, and if the Jacobian $[\partial F_i / \partial \lambda_j]$ is nonsingular, then the $\lambda^{(r)}$ will converge to the solution in the limit as $r \rightarrow \infty$.

The method is based on the fact that, for small changes $\Delta \lambda^{(r)}$ in the arguments $\lambda^{(r)}$, we have the approximate equality

$$F_i(\lambda^{(r)} + \Delta \lambda^{(r)}) \approx F_i(\lambda^{(r)}) + \sum_{k=1}^m \frac{\partial F_i(\lambda^{(r)})}{\partial \lambda_k^{(r)}} \Delta \lambda_k^{(r)}$$

up to a term of order $o(\Delta \lambda^{(r)})$. We therefore take $\Delta \lambda^{(r)}$ to be a solution of the linear equation

$$\sum_{k=1}^m \frac{\partial F_i(\lambda^{(r)})}{\partial \lambda_k^{(r)}} \Delta \lambda_k^{(r)} = -F_i(\lambda^{(r)}) \quad (\text{A.17})$$

and set $\lambda^{(r+1)} = \lambda^{(r)} + \Delta \lambda^{(r)}$. In applying the Newton-Raphson method to cross-entropy minimization, we let $F_i(\lambda)$ be proportional to the discrete form of the left-hand side of (A.7); we set

$$F_i(\lambda^{(r)}) = \sum_{j=1}^n f_{ij} p_j \exp\left(-\sum_{u=1}^m \lambda_u^{(r)} f_{uj}\right), \quad (\text{A.18})$$

$$\frac{\partial F_i(\lambda^{(r)})}{\partial \lambda_k} = -\sum_{j=1}^n f_{ij} f_{kj} p_j \exp\left(-\sum_{u=1}^m \lambda_u^{(r)} f_{uj}\right), \quad (\text{A.19})$$

where $f_{ij} = f_i(x_j) - \bar{f}_i$, and we have removed a factor of $\exp[-\sum_u \lambda_u^{(r)} \bar{f}_u]$. With the abbreviation

$$g_j = p_j^{1/2} \exp\left(-\frac{1}{2} \sum_{u=1}^m \lambda_u^{(r)} f_{uj}\right),$$

we express the right-hand sides of (A.18) and (A.19) in matrix notation as $[\tilde{f} \text{diag}(\tilde{g}) \tilde{g}]_i$ and $[\tilde{f} \text{diag}(\tilde{g})^2 \tilde{f}^t]_{ik}$, respectively, where $\text{diag}(\tilde{g})$ is the diagonal matrix whose diagonal elements are the g_j , and $\tilde{f}^t$ is the transpose of $\tilde{f}$. The solution of (A.17) is then given by

$$\Delta \lambda^{(r)} = [(\tilde{f} \text{diag}(\tilde{g})^2 \tilde{f}^t)^{-1} \tilde{f} \text{diag}(\tilde{g})] \tilde{g}.$$

We remark that the quantity in brackets is the Moore-Penrose generalized inverse [12] of the matrix $\tilde{f} \text{diag}(\tilde{g})$. The approach just described has been made the basis for a computer program [13], written in APL, for solving cross-entropy minimization problems with arbitrary positive discrete priors p and equality constraints specified by matrices $\tilde{f}$. The approach is particularly convenient for programming in APL since the generalized inverse is a built-in APL primitive function [14]. To solve a minimum-cross-entropy problem with 500 states and 10 constraints, the program typically requires 15 seconds of CPU time when running under the APL SF interpreter on a DEC-10 system with a KI central processor.

Gokhale and Kullback [15] describe a somewhat different algorithm, also based on the Newton-Raphson method, that has been implemented in PL/I. Agmon, Alhassid, and Levine [16],[17] describe yet another cross-entropy minimization algorithm and a FORTRAN implementation. Tribus [7] presents programs in BASIC that compute singly and doubly truncated Gaussian distributions as maximum entropy distributions with prescribed means and variances.

APPENDIX B

Remark on the Discussion of Property 12

In the discussion of Property 12, it was stated that for any prior p and any density $r \in \mathcal{D}$ with $H(r, p) < \infty$, there exists a finite set of equality constraints I_r such that $r = p \circ I_r$. In fact, at most two are needed. Let

$$\begin{aligned} f_1(\underline{x}) &= \begin{cases} 0, & r(\underline{x}) \neq 0 \\ 1, & r(\underline{x}) = 0, \end{cases} \\ \bar{f}_1 &= 0, \\ f_2(\underline{x}) &= \begin{cases} \log(p(\underline{x})/r(\underline{x})), & r(\underline{x}) \neq 0 \\ 0, & r(\underline{x}) = 0, \end{cases} \\ \bar{f}_2 &= -H(r, p), \end{aligned}$$

and impose constraints

$$\begin{aligned} \int q(\underline{x}) f_1(\underline{x}) d\underline{x} &= \bar{f}_1, \\ \int q(\underline{x}) f_2(\underline{x}) d\underline{x} &= \bar{f}_2. \end{aligned}$$

The first constraint implies $(p \circ I)(\underline{x}) = 0$ where $r(\underline{x}) = 0$. On the complementary set, where $r(\underline{x}) \neq 0$, define $q(\underline{x})$ by (A.4) with all $\lambda_j = 0$ except $\lambda_2 = 1$; this gives a function q that satisfies the second constraint as well as the first and also agrees with r . Hence $r = q$ is the result of minimizing $H(q, p)$ with respect to (B.1) and (B.2).

REFERENCES

1. S. Kullback, Information Theory and Statistics, Wiley, New York, 1959.
2. E.T. Jaynes, "Information Theory and Statistical Mechanics I", Phys. Rev. 106, 1957, pp.620-630.
3. W.M. Elsasser, "On Quantum Measurements and the Role of the Uncertainty Relations in Statistical Mechanics," Phys. Rev. 52, (Nov. 1937), pp. 987-999.
4. J.E. Shore and R.W. Johnson, "Axiomatic Derivation of the Principle of Maximum Entropy and the Principle of Minimum Cross-Entropy," IEEE Trans. Inf. Theory, vol. IT-26, no. 1, Jan. 1980.
5. R.W. Johnson, "Axiomatic Characterization of the Directed Divergences and Their Linear Combinations," IEEE Trans. Inf. Theory, vol. IT-25, no. 6, pp. 709-716, Nov. 1979.
6. I. Csizsár, "I-Divergence Geometry of Probability Distributions and Minimization Problems," Ann. Prob., vol. 3, pp. 146-158, 1975.
7. M. Tribus, Rational Descriptions, Decisions, and Designs, Pergamon Press, New York, 1969.
8. J.E. Shore, "Minimum Cross-Entropy Spectral Analysis," NRL Memorandum Rept. 3921, Naval Research Laboratory, Washington, D.C. 20375, January, 1979.
9. J.D. Markel and A.H. Gray, Jr., Linear Prediction of Speech, New York, Springer-Verlag, 1976.
10. L. Geronimus, Orthogonal Polynomials, Consultants Bureau, New York, 1961.

11. J.E. Shore and R.M. Gray, "Optimal Information-Theoretic Pattern Classification and Coding," in preparation.
12. A.E. Albert, Regression and the Moore-Penrose Pseudoinverse, Academic Press, N. Y. 1972.
13. R.W. Johnson, "Determining Probability Distributions by Maximum Entropy and Minimum Cross-Entropy," Proceedings APL79, pp. 24-29.
14. M.A. Jenkins, "The Solution of Linear Systems of Equations and Linear Least Squares Problems in APL," Scientific Center Technical Report No. 320-2989, IBM, N.Y., June, 1970.
15. D.V. Gokhale and S. Kullback, The Information in Contingency Tables, Marcel Dekher, New York, 1978.
16. Y. Alhassid, N. Agmon, and R.D. Levine, "An Upper Bound for the Entropy and its Applications to the Maximal Entropy Problem," Chem. Phys. Lett. 53, no. 1, 1978, pp. 22-26.
17. N. Agmon, Y. Alhassid, and R.D. Levine, "An Algorithm for Finding the Distribution of Maximal Entropy," J. Comput. Phys. 30, 1979, pp. 250-258.

ILMED
8