

UNCLASSIFIED

AR-004-230

DEPARTMENT OF DEFENCE
DEFENCE SCIENCE AND TECHNOLOGY ORGANISATION
ELECTRONICS RESEARCH LABORATORY

TECHNICAL REPORT

ERL-0381-TR

A NEW APPROACH TO MULTITARGET TRACKING USING
PROBABILISTIC DATA ASSOCIATION

S.B. Colegrove

S U M M A R Y

This report develops the theory for a multitarget tracking algorithm based on Probabilistic Data Association with new selection rules for assigning sensor measurements to target tracks and for forming multitrack clusters. These new rules remove the requirement to form a gate about each target's predicted position for the selection of sensor measurements. The resultant algorithm is the same for all target tracks and clutter conditions. The algorithm adapts to the sensor measurements via probability terms which model the environment and sensor processing.



POSTAL ADDRESS: Director, Electronics Research Laboratory,
Box 2151, GPO, Adelaide, South Australia, 5001.

UNCLASSIFIED

TABLE OF CONTENTS

	Page
1. INTRODUCTION	1
2. OUTLINE OF TARGET, SENSOR AND UPDATE MODEL	2
3. NEW SELECTION AND CLUSTERING RULES	5
4. EVENT SPACE FOR SELECTED MEASUREMENTS	6
5. TARGET STATE ESTIMATE FOR PDA FILTER	12
6. EVENT PROBABILITIES FOR THE PDA FILTER	13
7. CONCLUSIONS	18
8. ACKNOWLEDGEMENTS	18
REFERENCES	19

LIST OF FIGURES

1. Sensor model to produce measurements in the region V_c
2. Example of two target tracks with sensor measurements
3. Sample spaces and mapping functions for the sensor model and one target
4. Illustration of Mapping between the unordered and ordered spaces

LIST OF APPENDICES

I DERIVATION OF THE PDA FILTER	21
II DERIVATION OF THE PDA FILTER EVENT PROBABILITIES	24
III EXAMPLE DERIVATION OF THE EVENT PROBABILITIES	26
TABLE III.1 ORIGIN LISTS FOR EVENT PROBABILITY β_0	26
TABLE III.2 ORIGIN LISTS FOR EVENT PROBABILITY β_1	27

1. INTRODUCTION

A solution to the problem of tracking a known object in a cluttered environment was addressed by Bar-Shalom and Tse(ref.2) who formulated a sub-optimal solution referred to as the probabilistic data association (PDA) filter. Their solution builds on the Kalman filter by taking into account the event that sensor measurements used for updating the filter could have originated from clutter and noise. In their analysis, and that of extensions to the PDA filter in references 3, 4 and 5, the sensor measurements for updating the filter are all those located inside a 'gate' which is centered on the target's predicted position. The differences between the predicted target position and these measurements are weighted by the probability of their having originated from the target. This leads to a weighted difference which is used to form the filtered estimate. This approach eliminates the need for complex testing procedures before the tracking step to resolve the target origins of measurements.

With the use of a gate, the number of measurements used to update a PDA filter can vary dramatically with time and clutter conditions. In the case of the Joint PDA filter(ref.3,4) for tracking multiple targets, the measurements within all the gates of neighbouring tracks are used for the probability computations. Neighbouring tracks are defined as those whose gates overlap and this group of tracks is referred to as a track cluster. In this case the probabilities of all origins of measurements from targets in the track cluster are computed. This results in a rapid growth of the number of possible origins with cluster size and number of measurements. Because of computer limitations, any implementation of these algorithms has to include upper bounds on the number of measurements and the cluster size.

To remove the above variability and thereby derive a filter algorithm which is independent of target and clutter conditions, new rules, which were foreshadowed in references 6 and 7, are introduced for selecting sensor measurements and for forming target clusters. In references 6 and 7 the gate selection approach for the PDA filter was replaced by an approach which involved selecting the 'm' nearest measurements with 'm' fixed. In this paper, the analysis in reference 7 is extended from the single target case to include track clusters of fixed size. In this case a cluster consists of a group of 't' tracks which may not necessarily be 'close' to one another. In addition, this paper also develops a sensor model and refines the analysis contained in reference 7. The incorporation of the event to account for track initiation and termination (see reference 7) is not included in this paper but will be treated separately. The emphasis of this paper is on the theoretical development of a PDA filter with these new selection rules.

The theoretical development is organized so that the first section gives an outline of the target, sensor and update models. This is followed by the definition of the new selection and clustering rules. The feasible events are then defined and this is followed by the formulation of the new algorithm based on these rules. The assumptions used in the development of the algorithm are numbered for ease of reference and to help emphasise the assumptions on which the analysis is based. As reference is made to Kalman filter theory which uses the symbol 'P' for covariance, the notation $P(.)$ is used to denote a covariance matrix, P and $Pr\{.\}$ denote probability. For continuous variables $p(.)$ denotes probability density while for discrete variables, $p(.)$ represents the probability mass function. In order to avoid unnecessary complication of the equations, the time index is only given when a variable or event is first defined. Thereafter the time index is only retained when required for clarity. In the case of scalars which are derived from time dependent vectors and other scalars, the time index is not included as this is implied.

2. OUTLINE OF TARGET, SENSOR AND UPDATE MODEL

The model used to represent the environment, target and sensor is illustrated in figure 1. Here the sensor receiver derives its input from clutter, noise and targets distributed throughout the sensor's surveillance region. This input can, in the case of a radar, be electromagnetic waves generated from the illumination of the sensor surveillance region with an appropriately modulated transmitter.

When propagation time between objects in the surveillance region and the sensor receiver is neglected, the receiver is equivalent to a bank of parallel receivers. Each of these receivers gives the output from one of the sensor's resolution cells in the surveillance region. The receiver number or index 'i' corresponds to the resolution cell number. The output from all of the receivers at time 'k', then consists of N values which are represented by the set $\{\zeta_1(k), \dots, \zeta_N(k)\}$, where N is the number of sensor resolution cells.

The detection process following the receiver, searches the receiver output to locate likely target signals. This process could, for example, involve comparing the receiver outputs against a threshold and those ζ 's which pass the threshold are then used for parameter estimation. This latter process normally involves selecting each ζ from the previous test and then comparing it with ζ 's in adjacent resolution cells to estimate the surveillance region coordinates of the source of the received signal. Each estimate from the detection process is referred to as a sensor measurement which is represented by the vector $z_i(k)$. This vector contains components such as range, bearing, etc and 'i' is an arbitrary measurement subscript. The measurements output by the detection process are represented by the set $\{z_1(k), \dots, z_{N'}(k)\}$ where N' is the number of measurement vectors from the detection process at time 'k'.

The clutter map process which follows the detection process selects those measurements in the clutter region V_c which is a subset of the surveillance region. The region V_c can be based on the predicted position of a track or on the distribution of the clutter and noise throughout the surveillance region. For the latter case, the surveillance region is partitioned into regions of approximately constant clutter and noise density. For either case, the clutter region is made so that the probability density function for clutter and noise is invariant. Measurements derived from clutter and noise are collectively referred to as clutter measurements and they satisfy the following assumption.

ASSUMPTION NO.1: Clutter measurements are assumed to be independent and uniformly distributed within a clutter region V_c which may or may not contain targets. The total number of measurements in V_c is known and represented by $n(k)$, where $n(k) \gg 1$.

The following analysis is based on targets in an arbitrary clutter region V_c with volume v_c . Extension to other clutter regions is by identical analysis.

The λ^{th} target's state at time index 'k' is given by $x_\lambda(k)$. The model used for the sensor measurement step and the equations for the target's state as a function of the time index 'k' are represented by:

ASSUMPTION NO.2(a): A sensor measurement $z(k)$ from target λ is related to the target's state $x_\lambda(k)$ by

$$z(k) = H(k)x_\lambda(k) + v(k) \quad (1)$$

where $H(k)$ is the measurement matrix and $v(k)$ is the measurement noise introduced by the above parameter estimation step in conjunction with the sensor parameters. This noise is represented by a zero mean and covariance $R(k)$.

ASSUMPTION NO.2(b): The target state evolves with time by

$$x_\lambda(k+1) = F(k)x_\lambda(k) + u(k) \quad (2)$$

where $F(k)$ is the state transition matrix and $u(k)$ is the process noise with zero mean and covariance $Q(k)$. The process noise accounts for target manoeuvres and model errors.

ASSUMPTION NO.2(c): The measurement matrix H , measurement error covariance R , target state transition matrix F and process noise covariance Q are known.

Note that from equation (1), the measurement vector does not contain a target subscript since the measurement origin is unknown. In this paper, multiple targets are considered and it is assumed that they satisfy the following assumption.

ASSUMPTION NO.3: A target's state is independent of all other targets and sensor measurements which originate from targets are independent of each other and of n .

Independent target states are based on minimum a priori knowledge. As a general rule this condition is satisfied, eg crossing aircraft. However in the case of aircraft flying in formation, their states are no longer independent. In this analysis no use is made of any correlation between target states. Independent sensor measurements normally apply unless the detection process is overloaded or the receiver's output from any resolution cell contains a significant contribution from more than one target. This latter case occurs when targets are of the order of a resolution cell apart and leads to a bias in the parameter estimation step. The closely spaced target case can also lead to a merged measurement. This has been analysed for the PDA filter in reference 4. In this analysis, the merged measurement is not considered.

The estimate of the λ^{th} target's state at time k , given data up to time j , is denoted by $\hat{x}_\lambda(k|j)$. The covariance associated with this estimate is denoted by $P_\lambda(k|j)$. At time ' k ' it is assumed that,

ASSUMPTION NO.4: For every target in the region V_c , there is a corresponding track which contains a state estimate $\hat{x}_\lambda(k|j)$ and covariance $P_\lambda(k|j)$.

For the special case when there is no uncertainty in the origin of sensor measurements used for updating a target's state estimate, the associated covariance is given by $P_{\lambda}^1(k|k)$. If Assumption No.2 is satisfied, the discrete-time Kalman filter(ref.1) applies and the updated state is given by,

$$\hat{x}_{i\lambda}(k|k) = \hat{x}_{\lambda}(k|k-1) + W_{\lambda}(k)\tilde{y}_{i\lambda}(k) \quad (3)$$

$$P_{\lambda}^1(k|k) = [I - W_{\lambda}(k)H]P_{\lambda}(k|k-1) \quad (4)$$

Where the innovation vector, innovation covariance, gain matrix and predicted covariance are respectively,

$$\tilde{y}_{i\lambda}(k) = z_i(k) - H\hat{x}_{\lambda}(k|k-1) \quad (5)$$

$$S_{\lambda}(k) = HP_{\lambda}(k|k-1)H^T + R \quad (6)$$

$$W_{\lambda}(k) = P_{\lambda}(k|k-1)H^TS_{\lambda}^{-1}(k) \quad (7)$$

$$P_{\lambda}(k|k-1) = FP_{\lambda}(k-1|k-1)F^T + Q \quad (8)$$

Here H^T and F^T denote H transpose and F transpose respectively.

When applying a Kalman filter to sensor measurements with unknown origins from a multitarget and cluttered environment, the measurements used for updating the λ^{th} filter are often chosen from those contained within a selection gate. This gate is scaled from the innovation covariance $S_{\lambda}(k)$, and placed about the predicted state $H\hat{x}_{\lambda}(k|k-1)$. If more than one measurement is contained within the gate or if the gates from two or more tracks overlap, the measurements are often assigned to tracks on the basis of nearest neighbour. Figure 2 illustrates the case of a two-dimensional region containing two targets with overlapping selection gates and measurements from clutter and targets. If measurements are assigned using a nearest neighbour rule, then one possible assignment is measurement 3 to track 1 and measurement 5 to track 2. Other assignments are possible but, regardless of the assignment, the filter update (see equations (3) and (4)) assumes that the assignment is correct. To reduce the occurrence of non-target measurements updating the filter and thereby minimise errors in the target state estimate, the sensor clutter measurement rate has to be maintained at a low level. In many applications, maximum sensitivity is required and therefore the clutter measurement rate is often adjusted to the maximum rate that can be tolerated by the tracking system. The following sections deal with an alternative approach to multitarget tracking in cluttered conditions based on a PDA filter which takes account of the clutter and other target origins for sensor measurements which update the filter.

3. NEW SELECTION AND CLUSTERING RULES

The new sensor measurement selection rules and track clustering rules introduced in this paper replace the 'gate' concept. These rules are based on a nearest neighbours approach where more than one neighbour is selected. Here the 'm' nearest measurements and the 't-1' nearest tracks to each reference track are selected with 'm' and 't' constant. In the following development, a track is considered with respect to the 't-1' surrounding tracks and the track subscript is re-ordered so that $\lambda=1$ represents the reference track. To simplify notation, the track subscript is omitted when dealing with the reference track. The rules for selecting the measurements and the track clusters for each reference track can be written as;

- (a) Select the next track as reference and set its track subscript $\lambda=1$,
- (b) Select the 't-1' nearest target tracks based on the distance,

$$\tau_{\lambda} = \tilde{x}_{\lambda}^T S^{-1} \tilde{x}_{\lambda} \quad (9)$$

where $\tilde{x}_{\lambda} = \hat{Hx}_{\lambda}(k|k-1) - H\hat{x}_1(k|k-1)$. The track subscript ' λ ' is re-ordered so that $\lambda=t$ is the most distant track.

- (c) Select the 'm' nearest measurements to the reference track based on the distance,

$$d_i = \tilde{y}_i^T S^{-1} \tilde{y}_i \quad (10)$$

The measurement subscript 'i' is re-ordered so that $i=1$ is the nearest and the track subscript ' λ ' in equation (10), which for this case is '1', has been omitted.

Thus 'm' measurements are always selected for each track and all tracks are in a cluster of 't' tracks. The following analysis is based on $t \leq m < n$. The 'm' nearest measurements to the reference track represent an ordered set derived from the set of 'n' measurements in the region V_c . The ordered measurements are denoted by $y_i(k)$, $i=1, \dots, m$ with the ordered set defined by,

$$Y(k) = \{y_1(k), \dots, y_m(k); d_1 \leq \dots \leq d_m\} \quad (11)$$

For each track in the cluster, the past data up to time 'k-1' are denoted by $T_{\lambda}(k-1)$. All the track data in the cluster of 't' tracks are represented by the set,

$$T(k-1) = \{T_{\lambda}(k-1)\}_{\lambda=1, \dots, t} \quad (12)$$

From these definitions, the same set of 'm' measurements are not necessarily selected when the other 't-1' tracks in the cluster are used as the reference track. This is not considered to result in any penalty because the selected measurements form the group of most probable measurements for each reference track. Also this new procedure requires there to be 't' tracks, ie,

ASSUMPTION NO.5: There are at least 't' targets in the clutter region V_c . Measurements from targets which are not included in any cluster of 't' tracks are equivalent to clutter measurements.

To ensure that there are always 't' tracks, dummy tracks can be included with parameters such that they do not interfere.

4. EVENT SPACE FOR SELECTED MEASUREMENTS

The measured values at time 'k' from each of the process outputs in the previous section represent the values of random variables. Throughout this section, a symbol without a time index, eg ζ_i , denotes a random variable while a symbol with a time index is the value of the random variable. Thus $\zeta_i(k)$ is the value of the random variable ζ_i at time 'k'. Now the N receiver outputs from the surveillance region represent a random sample $\{\zeta_1, \dots, \zeta_N\}$. Each ζ originates from either clutter or target plus clutter in the sensor's resolution cell. To denote the origin of the i^{th} receiver output, an origin index ℓ_i is introduced. This origin index has the range $0, 1, \dots, t$ with $\ell_i = 0$ used to denote a clutter origin and with $\ell_i > 0$ used to denote a target plus clutter origin associated with track $\lambda = \ell_i$. A target plus clutter origin is referred to as a target origin. Therefore the pair (ζ_i, ℓ_i) denotes the i^{th} receiver output with origin defined by ℓ_i . Because the numbers of the non-zero origins are identical to the re-ordered track subscripts, an origin can be used for a track subscript.

The elements $\mathbf{r}$ of the receiver output space $\mathbf{R}$ consist of N receiver outputs with N origins and past target data T. By grouping the ζ and ℓ variables separately, the element $\mathbf{r}$ is given by $\mathbf{r} = [(\zeta_1, \dots, \zeta_N), (\ell_1, \dots, \ell_N), T]$. The output space $\mathbf{R}$ is then constructed from all possible values for the element $\mathbf{r}$. These values are obtained by allowing T and the N variables ζ and ℓ to vary through all permissible values. Remember that the origin value of each component in the elements of the sample space is known. The valid lists of origins contained in each element $\mathbf{r}$ have a maximum of one occurrence of each target plus clutter origin. Now the measurements with a clutter origin are samples from a common distribution function whereas each measurement with a target origin is from a unique distribution function which is indexed by the origin value. From the definition of the event space, there are elements with the values $\zeta_1(k), \dots, \zeta_n(k)$ and associated origins $\ell_1(k), \dots, \ell_n(k)$ which have identical ζ and ℓ values but the locations of the clutter measurements re-arranged. These elements have the same joint distribution function on $\mathbf{R}$. No similar re-arrangement of target origin subscripts exists because each element $\mathbf{r}$ has a maximum of one measurement for each $\ell_i > 0$. Remember that each arrangement of clutter subscripts represents a different element $\mathbf{r}$.

The detection process which follows the receiver (see figure 1) is represented as a mapping from the receiver output space $\mathbf{R}$ to the detection output space $\mathbf{U}$. A detection space element is given by $\mathbf{u}=[(z_1, \dots, z_{N'}), (\ell_1, \dots, \ell_{N'}), T]$ where the range of N' is $m+1, \dots, N$. The mapping is illustrated in figure 3 where the vector function $\mathbf{h}(\mathbf{r})$ is used to denote the detection process. From Section 2, the detection process operates on a number of receiver outputs for the parameter estimation step. Therefore each measurement with origin attached is given by $(z_j, \ell_j) = h_j(\{\zeta, \ell\}_i)$ where $\ell_j = \ell_i$ and $\{\zeta, \ell\}_i$ is the set of receiver outputs used for parameter estimation. The components of the elements of the sample space $\mathbf{U}$ are obtained from those receiver outputs where $h(\{\zeta, \ell\}_i) \neq \emptyset$. When a receiver output is obtained with a target origin, the target is detected. Detection is represented by the event D_ℓ which is defined by,

$$D_\ell = \{\mathbf{u} = \mathbf{h}(\mathbf{r}); (z_i, \ell_i) = h_i(\mathbf{r}) \text{ for some } i=1, \dots, N', \text{ and } \forall \mathbf{r} \in \mathbf{R}\} \quad \ell=0, 1, \dots, t \quad (13)$$

Where $h_i(\mathbf{r})$ is the i^{th} component of the vector function $\mathbf{h}(\mathbf{r})$. From equation (13), D_ℓ and $\bar{D}_\ell$, for each $\ell > 0$, partition $\mathbf{U}$ and figure 3 gives the case for one target only. In the case of D_0 , every element of $\mathbf{U}$ has at least one component with a clutter origin. Therefore $D_0 = \mathbf{U}$. In this analysis there are 't' targets and their events D_ℓ are independent of each other (see Assumption No.3). For an arbitrary element $\mathbf{u} \in \mathbf{U}$, then $\mathbf{u} \in D_\ell$ or $\bar{D}_\ell$ for each $\ell > 0$. All the valid combinations of D_ℓ or $\bar{D}_\ell$ for the 't' targets can be derived from a binary number of 't' bits. This binary number is represented by the set

$$B_\kappa(t) = \{(b_1, \dots, b_t); b_i \in (0, 1) \text{ and } \kappa = 1 + \sum b_i 2^{(i-1)}\} \quad \kappa=1, \dots, 2^t \quad (14)$$

By equating the bit b subscript 'i' to the event D subscript ' ℓ ', the bit values then denote D and $\bar{D}$ where $b_i = 1$ denotes D_i and $b_i = 0$ denotes $\bar{D}_i$. For the case of $t=2$, there are 4 partitions on $\mathbf{U}$ where each partition is the intersection of

κ	$B_\kappa(2)$	Event on $\mathbf{U}$
1	0 0	$\bar{D}_1 \cap \bar{D}_2$
2	1 0	$D_1 \cap \bar{D}_2$
3	0 1	$\bar{D}_1 \cap D_2$
4	1 1	$D_1 \cap D_2$

Note that through the action of the clutter map process and the selection process, these partitions map onto the output sample spaces of these processes. The elements contained in the mapped partitions are not the same but, to avoid a burgeoning notation problem, the same symbol is used. The sample space associated with a mapped partition is apparent from its context. Furthermore, the origin accompanying the variable from each process is

represented by the same symbol ℓ . It is to be remembered that the origin value $\ell_i(k)$ linked with say, the receiver output value $\zeta_i(k)$, is not the same origin for the component i from say, the detection process.

The clutter map process which follows the detection process, selects those measurements in the region V_c and is represented by the function $v(z_i, \ell_i)$. This function maps from U to Z as illustrated in figure 3. The elements of the clutter map output space Z are given by $z = \{(z'_1, \dots, z'_n), (\ell_1, \dots, \ell_n), T\}$ where n has the range $m+1, \dots, N'$ and $z = v(u)$. As in the detection process, the clutter map does not pass all measurements and the measurement subscripts in any element z are arbitrarily assigned. When a measurement with origin ℓ is output, the event V_ℓ occurs and is defined as

$$V_\ell = \{z = v(u); (z'_i, \ell_i = \ell) = v_i(u) \text{ for some } i=1, \dots, n, \\ \text{and } \forall u \in D_\ell\} \quad \ell=0, 1, \dots, t \quad (15)$$

Where $v_i(u)$ is the i^{th} component of the vector function $v(u)$. From the above, for origin $\ell > 0$, each element $z \in D_\ell \cap V_\ell$ has $\ell = \ell_i$ for some i . The complement also exists which contains those elements z where the target origin $\ell \neq \ell_i$ for all i , ie $z \in \overline{D_\ell \cap V_\ell}$. Therefore the binary number $B_\kappa(t)$ can also be used to obtain all possible combinations of these events on Z .

The clutter map output space Z is the input to the selection process which selects the ' m ' nearest z' values using the scalar distance d in equation (10). The selection process which is represented by the vector mapping function $\psi(z)$, involves a partial ordering of Z into elements with the ' m ' smallest d values. The remaining $n-m$ z' values have d values greater than the ' m ' smallest and are discarded while the value of n is retained. Thus the selection process output space Y has the elements $y = [((y_1, \dots, y_m), (\ell_1, \dots, \ell_m), n, T); d_1 \leq \dots \leq d_m]$.

For any element $z \in Z$, origin lists ' L ', ' A ' and ' $\bar{A}$ ' are introduced. The origin list L contains the origins associated with the measurements which have the ' m ' smallest d values. The origin list A contains target origins not in L but contained in the element z . The origin list $\bar{A}$ contains the target origins in $LU \bar{A}$. In addition to these lists, an arrangement list ' A ', which contains the measurement subscripts associated with the origins in the list L is included. Thus A provides the link between the origins in L and the measurements in z . The lists L , A , $\bar{A}$ and $\bar{L}$ are then defined as

$$L \equiv \{(\ell_1, \dots, \ell_m); d_1 \leq \dots \leq d_m\} \\ A \equiv \{(\ell_1, \dots, \ell_m); d_1 \leq \dots \leq d_m\} \\ \bar{A} \equiv \text{the unordered list of target origins in } \bar{L} \cap \{\ell_1, \dots, \ell_n\} \\ \bar{L} \equiv \text{the unordered list of target origins in } \overline{LU \bar{A}} \quad (16)$$

From the above, there are many lists $(L, A, \bar{A}, \bar{A})$. The primary index for these lists is based on the knowledge that the Kalman filter provides the filtered estimate for sensor measurements from a target. Because of this and the fact that more than one target is considered, a second index or subscript is also added for referencing a particular list L , ie L_{ij} . The first subscript 'i' denotes the position of the origin with $\ell_i = 1$ in L . Thus $i=2$ represents the case of the second smallest d value originating from the reference target. The range of the subscript 'i' is $i=0, \dots, m$ where the list L_{0j} contains no origin equal to the reference target index. The second subscript 'j' denotes the permutation of the remaining $m-1$ origins with $\ell_i > 1$. Thus 'j' represents an index for track origin arrangements and has the range $j=1, \dots, \xi_i$. Examples of the track origin arrangements L_{0j} and L_{1j} for $m=2$ and $t=2$ are,

$$\begin{aligned} L_{01} &= \{0,0\} & L_{11} &= \{1,0\} \\ L_{02} &= \{0,2\} & L_{12} &= \{1,2\} \\ \bar{L}_{03} &= \{2,0\} \end{aligned} \tag{17}$$

For each list L_{ij} the following symbols are defined.

$$\begin{aligned} \phi_{ij} &\equiv \text{number of clutter origins in } L_{ij} \\ \bar{L}_{ij} &\equiv \text{set of target origins not in } L_{ij}, \quad j=1, \dots, \xi_i \end{aligned} \tag{18}$$

Note that elements z associated with a list L_{ij} have components with some of the origins in $\bar{L}_{ij}$. These elements have the d value associated with an origin $\ell \notin L_{ij}$, greater than d_{i_m} , ie not in the m smallest. Therefore all elements z which have the d value of a measurement with target origin ℓ in the m smallest are represented by the event M_ℓ . By adding the measurement origin to the d value, the event M_ℓ is defined by,

$$M_\ell = \{z; (d_i, \ell_i = \ell) \leq (d_{i_m}, \ell_{i_m}) \text{ and } \forall z \in D_\ell \cap V_\ell\} \quad \ell=0, 1, \dots, t \tag{19}$$

Therefore from equation (19), all elements in $D_\ell \cap V_\ell \cap \bar{M}_\ell$ have the d value of the measurement with origin ℓ outside the m smallest but contained in the element z .

From the preceding, each element z corresponding to the origins in the set L_{ij} has $n-m$ origins which contain none or any one of the possible target origin

combinations selected from the origins in $\bar{L}_{ij}$. Each target combination is contained in the lists $(\Lambda, \bar{\Lambda})$. When an element z has an origin $\ell \in \Lambda$ and $\ell \in z$ then $z \in D_\ell \cap V_\ell \cap \bar{M}_\ell$. Conversely when the origin $\ell \in \bar{\Lambda}$ and $\ell \notin z$ then $z \in D_\ell \cap V_\ell$. To represent these possibilities, the binary number defined in (14) is utilized. In this case the number of bits is $t-m+\phi_{ij}$ which is the number of target origins in $\bar{L}_{ij}$. Therefore the binary number is $B_\kappa(t-m+\phi_{ij})$. The members of Λ_{ijk} are now obtained from the direct product of the members of $B_\kappa(t-m+\phi_{ij})$ and $\bar{L}_{ij}$. Clutter origins, ie zero entries, are not retained in $\bar{\Lambda}_{ijk}$. The members of Λ_{ijk} are therefore obtained from complementing the binary number before performing the direct product with $\bar{L}_{ij}$. These sets and the associated symbols to describe them are defined by

$$\begin{aligned}\Lambda_{ijk} &\equiv \text{target origins from the direct product of the} \\ &\quad \text{members of } \bar{L}_{ij} \text{ and } B_\kappa(t-m+\phi_{ij}), \\ \delta_{ijk} &\equiv \text{number of bits with } b_i=0 \text{ in } B_\kappa(t-m+\phi_{ij}), \\ \bar{\Lambda}_{ijk} &\equiv \text{the } \delta_{ijk} \text{ target origins in } \overline{L_{ij} \cup \Lambda_{ijk}}, \quad \kappa=1, \dots, \partial_{ij}\end{aligned}\quad (20)$$

where $\partial_{ij} = 2^{t-m+\phi_{ij}}$.

For each set of lists $(L_{ij}, \Lambda_{ijk}, \bar{\Lambda}_{ijk})$, there are many arrangement lists A to account for all possible permutations of clutter subscripts. These permutations are indexed by a fourth subscript 'h', ie A_{ijkh} . The range of h , which is represented by η_i , is found from the number of clutter origins in L_{ij} and the total number of clutter origins from n and the lists $(L_{ij}, \Lambda_{ijk}, \bar{\Lambda}_{ijk})$.

From the definition of the lists $(L, A, \Lambda, \bar{\Lambda})$ and event M_ℓ , the event space Z consists of the partitions

$$\begin{aligned}\Omega_{ijk} &= \{z; (L, A) = (L_{ij}, A_{ijkh}) \text{ for } h=1, \dots, \eta_{ijk} \text{ and } \forall z \in D_{ijk}\} \\ &\quad i=0, \dots, m, \quad j=1, \dots, \xi_i \text{ and } \kappa=1, \dots, \partial_{ij}\end{aligned}\quad (21)$$

$$\text{where } D_{ijk} = \left[\bigcap_{\ell \in L_{ij}} (D_\ell \cap V_\ell \cap M_\ell) \right] \left[\bigcap_{\ell \in \Lambda_{ijk}} (D_\ell \cap V_\ell \cap \bar{M}_\ell) \right] \left[\bigcap_{\ell \in \bar{\Lambda}_{ijk}} (\overline{D_\ell \cap V_\ell}) \right]$$

Figure 4 represents diagrammatically these partitions on the sample space Z . The D, V and M events on Z are not shown because of the complexity of the partitions as illustrated in equation (21).

From equation (21), each partition Ω_{ijk} contains η_{ijk} permutations of subscripts. Each permutation is represented as a sub-partition* ω_{ijkh} on Ω_{ijk} . The range of the subscript h depends on the D, V and M events in which an element z is contained. From the list definitions for partition Ω_{ijk} , the number of clutter measurements in L_{ij} is ϕ_{ij} and the total number of clutter measurements from which these are selected is $n-t+\delta_{ijk}$. Thus the number of clutter index permutations, and hence the number of sub-partitions is $(n-t+\delta_{ijk})!/(n-t+\delta_{ijk}-\phi_{ij})!$.

From the above, the mapping from Z to Y by the function $\psi(z)$ (see figure 4) involves a re-arrangement of subscripts and is a many-to-one mapping[10]. The inverse from each element $y \in \theta_{ijk}$ consists of η_{ijk} branches onto Ω_{ijk} . Each branch is a one-to-one from an element $y \in \theta_{ijk}$ to an element z of a sub-partition ω_{ijkh} . The partitions or events θ_{ijk} on the ordered space Y are defined as

$$\begin{aligned} \theta_{ijk} &= \{y = \psi(z); L = L_{ij} \text{ and } \forall z \in \Omega_{ijk}\} \\ i &= 0, \dots, m, \quad j = 1, \dots, \xi_i \text{ and } \kappa = 1, \dots, \partial_{ij} \end{aligned} \quad (22)$$

The union of θ_{ijk} over all i, j and κ gives the ordered space Y . The union of θ_{ijk} over all j and κ gives the partitions (or events) θ_i defined by,

$$\begin{aligned} \theta_i &= \bigcup_{j=1}^{\xi_i} \bigcup_{\kappa=1}^{\partial_{ij}} \theta_{ijk} & i &= 0, \dots, m \\ &= \{y; \ell_i = 1\} & i &= 1, \dots, m \\ &= \{y; \text{all } \ell_i \neq 1\} & i &= 0 \end{aligned} \quad (23)$$

* This sub-partition is only strictly correct when the ordering does not include the equal operator, ie $d_i < d_j$. In this paper, the ordering defined in equation (8) is used and, since the probability of two or more d values being equal is effectively zero, each arrangement is considered to form a sub-partition.

The events θ_i , $i=1, \dots, m$ then contain all elements y which have the i^{th} nearest measurement originating from the reference track. The event θ_0 contains all those elements y which do not have measurements from the reference track in the m nearest.

5. TARGET STATE ESTIMATE FOR PDA FILTER

From Section 3, the values output by the sensor at time 'k' for the reference track are $(Y(k), n(k), T(k-1))$. From the definition of the sample space Y , many elements can be found which have components with these values. Also the elements which contain these values are distributed throughout the events or partitions $\theta_{-1}, \dots, \theta_m$ on Y . The state estimate for the target associated with the reference track is then given by reference 2.

$$\begin{aligned}\hat{x}(k|k) &= E\{x(k)|Y(k), n(k), T(k-1)\} \\ &= \sum_{i=-1}^m \beta_i E\{x(k)|\theta_i(k), Y(k), n(k), T(k-1)\}\end{aligned}\quad (24)$$

where $\beta_i = \Pr\{\theta_i(k)|Y(k), n(k), T(k-1)\}$

That is, the state estimate is the weighted sum of the estimates of the target's state conditioned on all the target data. The weighting term is the probability that the i^{th} selected measurement is from the target. At this stage it is assumed that the event probabilities β_i are available. Note that the density associated with the estimate involves the sum of weighted normal densities. These are obtained by track split back to the start of the track(ref.9). To avoid the complexity of this procedure, the following assumption is introduced.

ASSUMPTION NO.7: An approximate sufficient statistic for the past target data $T_\ell(k-1)$ at time 'k' is mean $\hat{x}_\ell(k|k-1)$ and covariance $P_\ell(k|k-1)$.

This is the key assumption for the PDA filter. Using this assumption, Appendix I gives the derivation of the target's state and the associated covariance. The reference target's state is given by,

$$\hat{x}(k|k) = \hat{x}(k|k-1) + W\tilde{y} \quad (25)$$

where

$$\tilde{y} = \sum_{i=1}^m \beta_i \tilde{y}_i$$

and the associated covariance is given by,

$$P(k|k) = \beta_0 P(k|k-1) + [1-\beta_0]P^1(k|k) + W \left[\sum_{i=1}^m \beta_i \tilde{y}_i \tilde{y}_i^T - \tilde{y} \tilde{y}^T \right] W^T \quad (26)$$

The covariance in equation (26) consists of the sum of the covariance when all the measurements are clutter and the covariance when the correct measurement updates the filter. The final term in the sum reflects the uncertainty of the weighted innovation y by adding its associated covariance.

6. EVENT PROBABILITIES FOR THE PDA FILTER

The derivation of an expression for the event probabilities β_i (see equation (24)) requires expressions for the density function of an arbitrary y , the probabilities associated with the events D , V , and M , and the probability associated with the partition of $\mathbf{Y}$ by the value $n(k)$. This section gives an outline of the derivation of these probabilities.

From Assumption No.1, the density of y_i for a clutter measurement in the region $\mathbf{V}_c$ is $1/v_c$. From Assumption No.7, the probability density of a measurement from target ' ℓ ' at time ' k ' using data from time ' $k-1$ ' is normal. With y_i unconstrained, ie not restricted to being in the region $\mathbf{V}_c$, the density for an arbitrary clutter measurement and an arbitrary target measurement is then given by,

$$\begin{aligned} f_0(y_i) &= p(y_i, \ell_i=0) \\ &= 1/v_c && \text{if } y_i \in \mathbf{V}_c \\ &= 0 && \text{otherwise} \end{aligned}$$

$$\begin{aligned} f_\ell(y_i) &= p(y_i, \ell_i=\ell | T_\ell) \\ &= \exp(-\frac{1}{2} \tilde{y}_{i\ell}^T S_\ell^{-1} \tilde{y}_{i\ell}) / [(2\pi)^{\mu/2} \sqrt{|S_\ell|}] \quad \ell=1, \dots, t \end{aligned} \quad (27)$$

where μ = number of components of vector y .

The probability for event M , which is presented later, requires the probability $\Pr\{d_i \leq d\}$ for an unordered d_i . For clutter measurements, this probability is,

$$\begin{aligned} G_0(d) &= \Pr\{(d_i, \ell_i=0) \leq d\} \\ &= v(\mathbf{V}_d \cap \mathbf{V}_c) / v_c \quad 0 \leq d \leq r \end{aligned} \quad (28)$$

where $v(\cdot)$ denotes volume, V_d is the multidimensional ellipsoid $\{\tilde{y}^T S^{-1} \tilde{y} \leq d\}$ with volume $v_d = \pi^{\mu/2} d^{\mu/2} |S|^{1/2} / \Gamma(\mu/2 + 1)$ where μ is the number of elements in the measurement vector y . χ is the largest d which is contained within V_c .

In the variable ' d ', when $V_d \subset V_c$, the density for the reference target, $\ell=1$, is chi-squared with μ degrees of freedom. The probability $\Pr\{d_i \leq d\}$ for a target measurement is then

$$\begin{aligned} G_1(d) &= \Pr\{(d_i, \ell_i=1) \leq d | T_1\} \\ &= \int_0^d \frac{A_d u^{\mu/2-1} e^{-u/2} du}{2^{\mu/2} \Gamma(\mu/2)} \end{aligned} \quad (29)$$

where the term A_d is a scaling factor to account for the intersection of the V_d shell with V_c . In many applications this does not occur and therefore $A_d=1$. For series approximations to the integral in equation (29) see reference 8.

For other targets in the cluster, ie $\ell > 1$, no similar elegant function is available for $\Pr\{d_{i\ell} \leq d\}$. Therefore it is assumed that

ASSUMPTION NO.8: The probability function $\Pr\{(d_{i\ell}, \ell_i=\ell) \leq d | T_\ell\}$, for $\ell > 1$ is available and represented by $G_\ell(d)$.

An expression for $G_\ell(d)$ involves the integral of the multivariate normal density function $f_\ell(y)$ over the offset region V_d . An approximate expression for this probability is obtained by assuming that the covariances of track ' ℓ ' and the reference track are related by $S_\ell^{-1} = c^{-1} S^{-1}$, where c is a scalar. Under these conditions, the region V_d can be made circular by scaling the y coordinate space so that track ' ℓ ' has unit variance. Thus the density of d/c is non-central chi squared with μ degrees of freedom and non-centrality parameter $\sqrt{t_\ell}/c$. Series approximations for this distribution can be found in reference 8.

The event probabilities for D, V and M can now be given. When these events are assumed to be independent of each other and of n (see Assumption No's 1 and 3), they can be written as

$$P_{d\ell} = \Pr\{D_\ell(k) | T_\ell(k-1)\}$$

$$P_{v\ell} = \Pr\{V_\ell(k) | D_\ell(k), T_\ell(k-1)\}$$

$$P_{m\ell} = \Pr\{M_\ell(k) | D_\ell(k), V_\ell(k), T_\ell(k-1)\} \quad \ell=0, \dots, t \quad (30)$$

If any of the above events are not independent of n , then n has to be included on the LHS of the conditioning and the expression for the event probability then includes n .

From the definition of the above events, the probabilities for origin value $\ell=0$, ie P_{d0} , P_{v0} and P_{m0} are unity. The probability of target detection, $P_{d\ell}$, is application dependent because it is a function of the detection processing illustrated in figure 2. Therefore the assumption is made that

ASSUMPTION NO.9: The event probability $P_{d\ell}$ for $\ell=1, \dots, t$ is known.

The value for the probability of a measurement from the target being in the region V_c , ie $P_{v\ell}$, is obtained from the integral of $f_\ell(y_i)$ over the region V_c . This probability can be obtained by series approximations and from the dimensions of V_c .

The probability of a measurement from target ℓ having a d value in the m nearest, $P_{m\ell}$, is equal to $G_\ell(d)$ with the value $d=d_m$ (see equations (28), (29) and Assumption No. 8). Thus $P_{m\ell}=G_\ell(d_m)$.

The last remaining probability required to obtain a solution for the event probabilities deals with the partition of θ_i by the number of measurements $n(k)$. From Section 4, the elements y have n ranging from $m+1$ to N' . The partition caused by $n(k)$ is expressed in terms of the probability of obtaining 'r' clutter measurements in the region V_c given $T(k-1)$. This probability is represented by $\Phi(r|T(k-1))$ and

ASSUMPTION NO.10: The a priori probability of having n measurements from clutter is equal to that for $n-\phi$ measurements from clutter, ie $\Phi(n|T(k-1))=\Phi(n-\phi|T(k-1))$ where $\phi=1, \dots, t$.

An alternative to this assumption is to adopt the approach taken in reference 3 where a Poisson distribution with known clutter measurement density is used. This approach is not taken here because n is assumed to be large.

From the events defined in Section 4, the preceding density and probability functions, the event probabilities can be derived as shown in Appendix II, and are given by

$$\beta_i = b_i / \sum_{j=0}^m b_j \quad i=0, \dots, m \quad (31)$$

The b_i terms are obtained by evaluating,

$$b_i = \sum_{j=1}^{\xi_i} \left\{ \left[\prod_{\ell_1 \in L_{ij}} f_{\ell_1}(y_1) P_{d\ell_1} \right] \sum_{\kappa=1}^{\partial_{ij}} \{ a_{ij\kappa} (1-P_{m0})^{(\delta_{ij\kappa}-\phi_{ij})} \right. \\ \left. \left[\prod_{\ell \in \Lambda_{ij\kappa}} P_{d\ell} (P_{v\ell} - P_{m\ell}) \right] \prod_{\ell \in \bar{\Lambda}_{ij\kappa}} \alpha_{\ell} \right\} \quad (32)$$

where

$$a_{ij\kappa} = \frac{(n-t+\delta_{ij\kappa})!}{(n-t+\delta_{ij\kappa}-\phi_{ij})!}$$

$$\alpha_{\ell} = (1-P_{d\ell} P_{v\ell}) / P_{d\ell} P_{v\ell}$$

Equation (32) contains a density which is proportional to the product of the joint density for the targets and clutter measurements in the list L_{ij} times the probability of having the other $n-m$ clutter and target measurements outside the m nearest. This joint density is summed over all j origin lists. Note that from equation (32), the time dependent functions are n , $f_{\ell}(y_1)$, $P_{d\ell}$, $P_{v\ell}$ and $P_{m\ell}$. An example of the derivation of the b_i terms for the case $m=2$, $t=2$ is given in Appendix III. The validity of these results can be checked by setting $P_{d2}=0$ for the b_i terms in Appendix III. When the b_i terms are all multiplied by $P_{v1} P_{d2}$ and common terms cancelled, the following is obtained.

$$b_0 = n\alpha_1 P_{v1}/v_c + (n-2)(P_{v1} - P_{m1})/[v_c(1-P_{m0})]$$

$$b_1 = f_1(y_1) \quad (33)$$

$$b_2 = f_1(y_2)$$

This is the same result as that given in reference 7 when track initiation is not included and when $P_{v1}=1$. When $P_{d2}=1$ and $f_2(y)=0$, the b_0 term changes to

$$b_0 = (n-1)\alpha_1 P_{v1}/v_c + (n-3)(P_{v1} - P_{m1})/[v_c(1-P_{m0})] \quad (34)$$

This is the correct result since in this case the number of clutter measurements is reduced by one.

From the above it can be seen that if there are not 't' tracks in the region V_c , dummy tracks with a low $P_{d\ell}$ (ie 0.001) will not result in any interference with the valid tracks.

As can be seen from the development given in Appendix III, the algorithm complexity rapidly increases with 'm' and 't'. No general expression has been found for deriving the number of events from 'm' and 't'. Therefore the events are written down from the origin lists (see definitions (16), (18) and (20)). For most practical applications, $t=3$ and $m=3$ are reasonable upper limits. For this case, when $i=0$, $\xi=13$ and when $i>0$, $\xi=7$. The intermediate case of $t=2$ and $m=3$ has ξ values of 4 and 3 respectively. While this complexity exists, the algorithm is the same for all tracks and there is no on-line requirement for the computation of the feasible events as in reference 3,4. The algorithm adapts to changing clutter conditions via the terms n and $P_{m\ell}$.

7. CONCLUSIONS

The use of gates to both select sensor measurements and locate tracks which are likely to interfere results in a PDA algorithm whose number and complexity of computations change with conditions. This paper has developed the theory for a PDA algorithm based on alternative rules for selecting measurements and interfering tracks. These rules replace the requirement for a gate by a nearest neighbour rule where more than one neighbour is selected. While the algorithm is more complex than a PDA algorithm with tracks in isolation, it remains unaltered with clutter and target conditions.

Because tracks are all updated with the same computations, an array processor can be used to achieve the required processing power for the application of this algorithm to real-time tracking on H.F. and microwave radar systems. The algorithm presented in this paper is to be extended in a subsequent paper to include non-uniform clutter density, track initiation and termination.

8. ACKNOWLEDGEMENTS

The author is indebted to Dr A.W. Davis, Division of Mathematics and Statistics, CSIRO for his guidance with the probability theory developed in this paper. Thanks is also due to Mr J.K. Ayliffe, Radar Division, ERL for his help in clarifying many of the ideas introduced in this paper.

REFERENCES

- | No. | Author | Title |
|-----|---|--|
| 1 | Kalman, R.E | "A New Approach to Linear Filtering and Prediction Problems".
Journal of Basic Engineering, Vol.82,
pp.35-45, March 1960 |
| 2 | Bar-Shalom, Y. and
Tse, E. | "Tracking in a Cluttered Environment
With Probabilistic Data Association".
Automatica, Vol.11, pp.451-461,
September 1975 |
| 3 | Fortmann, T.E.,
Bar-Shalom, Y., and
Scheffe, M. | "Sonar Tracking of Multiple Targets
Using Join Probabilistic Data
Association".
IEEE Journal of Oceanic Engineering,
Vol.OE-8, No.3, pp.173-184, July 1983 |
| 4 | Chang, K., and
Bar-Shalom, Y. | "Joint Probabilistic Data Association
for Multitarget Tracking with Possibly
Unresolved Measurements and Maneuvers".
IEEE Transactions on Automatic Control,
Vol.AC-29, No.7, pp.585-594, July 1984 |
| 5 | Gauvrit, M. | "Bayesian Adaptive Filtering for
Tracking with Measurements of Uncertain
Origin".
Automatica, Vol.20, No.2, pp.217-224,
1984 |
| 6 | Colegrove, S.B. and
Ayliffe, J.K. | "An Extension of Probabilistic Data
Association to Include Track Initiation
and Termination".
Digest of the 20th IREE International
Convention, pp.853-856, September 1985 |
| 7 | Colegrove, S.B.,
Davis, A.W. and
Ayliffe, J.K. | "Track Initiation and Nearest
Neighbours Incorporated into
Probabilistic Data Association".
Journal of Electrical and Electronics
Engineering, Australia, Vol.6, No.3,
pp.191-198, September 1986 |
| 8 | | "Handbook of Mathematical Functions".
Dover Publications Inc., New York, 1965 |
| 9 | Singer, R.A.,
Sea, R.G. and
Housewright, K.B. | "Derivation and Evaluation of Improved
Tracking Filters for Use in Dense
Multitarget Environments".
IEEE Transactions on Information
Theory, Vol.IT-20, No.4, pp.423-432,
July 1974 |
| 10 | Zehna, P.W. | "Probability Distributions and
Statistics".
Allyn and Bacon, Inc., Boston,
pp.310-313, 1970 |

APPENDIX I

DERIVATION OF THE PDA FILTER

An expression for the state estimate of each reference target requires the mean $E\{x|\theta_i, Y, n, T\}$. This mean is considered separately for the cases of $i=0$ and $i=1, \dots, m$. From Assumption No.7, the Normal density function applies and it is abbreviated in this appendix as $N\{\text{mean}, \text{variance}\}$.

For $i=0$, the mean given the event that all measurements Y are not from the reference track and past track data, is

$$\begin{aligned} E\{x|\theta_0, Y, n, T\} &= E\{x|T_1\} \\ &= \hat{x}(k|k-1) \end{aligned} \quad (I.1)$$

Note that in this case, no use can be made of data Y and the current target state is independent of other targets (see Assumption No.3). The density associated with the estimate (see Assumption No.7) is

$$N\{\hat{x}(k|k-1), P(k|k-1)\} \quad (I.2)$$

For $i=1, \dots, m$, the mean given that the i^{th} measurement is from the reference track and all other measurements are not, is as before based only on the past data for the reference track. Thus from Assumption No's 1, 2 and 4, and the Kalman equations (3) to (8), the mean is given by;

$$\begin{aligned} E\{x|\theta_i, Y, n, T\} &= E\{x|y_i, T_1\} \\ &= \hat{x}(k|k-1) + W\tilde{y}_i \\ &= \hat{x}_i(k|k) \end{aligned} \quad (I.3)$$

For equation (I.3) the target index in equation (3) has been omitted because $l=1$. The density associated with this estimate is that for the Kalman filter and is represented by

$$N\{\hat{x}_i(k|k), P^1(k|k)\}, \quad (I.4)$$

where the covariance is obtained from equation (4). Note that the covariance is independent of the measurement 'i' because it is conditioned on this measurement being from the target.

From equation (23), the event probabilities sum to unity. Therefore with equations (I.1) and (I.3) substituted into equation (24), the state estimate

of the reference target is given by;

$$\begin{aligned}\hat{x}(k|k) &= \beta_0 \hat{x}(k|k-1) + \sum_{i=1}^m \beta_i \hat{x}_i(k|k) \\ &= \hat{x}(k|k-1) + W \tilde{y}\end{aligned}\quad (I.5)$$

where

$$\tilde{y} = \sum_{i=1}^m \beta_i \tilde{y}_i$$

Now from equation (24), the covariance associated with this estimate is given by;

$$\begin{aligned}P(k|k) &= \int [x - \hat{x}(k|k)] [x - \hat{x}(k|k)]^T p(x|Y, n, T) dx \\ &= \sum_{i=0}^m \beta_i \int [x - \hat{x}(k|k)] [x - \hat{x}(k|k)]^T p(x|\theta_i, Y, n, T) dx\end{aligned}\quad (I.6)$$

By substituting $[(x - \hat{x}_i(k|k)) + (\hat{x}_i(k|k) - \hat{x}(k|k))]$ for $[x - \hat{x}(k|k)]$ in equation (I.6) and by using the density defined in expression (I.4), the integral in equation (I.6) for $i=1, \dots, m$ becomes,

$$\begin{aligned}&\int [x - \hat{x}_i(k|k)] [x - \hat{x}_i(k|k)]^T N\{\hat{x}_i(k|k), P^1(k|k)\} dx \\ &+ [\hat{x}_i(k|k) - \hat{x}(k|k)] [\hat{x}_i(k|k) - \hat{x}(k|k)]^T \int N\{\hat{x}_i(k|k), P^1(k|k)\} dx \\ &+ [\hat{x}_i(k|k) - \hat{x}(k|k)] \int [x - \hat{x}_i(k|k)]^T N\{\hat{x}_i(k|k), P^1(k|k)\} dx \\ &+ \int [x - \hat{x}_i(k|k)] N\{\hat{x}_i(k|k), P^1(k|k)\} dx [\hat{x}_i(k|k) - \hat{x}(k|k)]^T\end{aligned}\quad (I.7)$$

The first integral in equation (I.7) evaluates to $P^1(k|k)$, the second is unity, and the third and fourth are zero. Similar results are obtained for $i=0$. Therefore equation (I.6) can be written as,

$$P(k|k) = \beta_0 P(k|k-1) + \sum_{i=1}^m \beta_i P^1(k|k) + \sum_{i=1}^m \beta_i \hat{\tilde{x}}_i(k|k) \hat{\tilde{x}}_i^T(k|k) - \hat{\tilde{x}}(k|k) \hat{\tilde{x}}^T(k|k) \quad (I.9)$$

From the definition of $\hat{\tilde{x}}(k|k)$ and $\hat{\tilde{x}}_i(k|k)$ in equations (I.5) and (I.3), The covariance of the state estimate can be written as,

$$P(k|k) = \beta_0 P(k|k-1) + [1-\beta_0] P^1(k|k) + W \left[\sum_{i=1}^m \beta_i \tilde{y}_i \tilde{y}_i^T - \tilde{y} \tilde{y}^T \right] W^T \quad (I.10)$$

APPENDIX II

DERIVATION OF THE PDA FILTER EVENT PROBABILITIES

The event probabilities β_i for the set of events defined by equation (24) can be expanded using Baye's rule to give:

$$\begin{aligned} \beta_i &= \Pr\{\theta_i | Y, n, T\} \\ &= \frac{p(Y, \theta_i, n | T)}{\sum_{i=0}^m (\text{numerator})} \end{aligned} \quad (\text{II.1})$$

The joint density in the numerator of equation (II.1) is the probability density for the ordered values in Y having the i^{th} value from the reference track and there being n measurements in V_c given the past target data T . From equation (23), θ_i is the union of θ_{ijk} over all j and κ , therefore the numerator expands to:

$$p(Y, \theta_i, n | T) = \sum_{j=1}^{\xi_i} \sum_{\kappa=1}^{\partial_{ij}} p(Y, \theta_{ijk}, n | T) \quad (\text{II.2})$$

From the definition of θ_{ijk} in equation (22), the density on the RHS of equation (II.2) is for the values in Y with the origins defined by L_{ij} with n measurements in V_c and with target measurements in the partition D_{ijk} .

The joint density of the unordered y values in Y with the origins defined in L_{ij} and in the partition $\bigcap_{\ell \in L_{ij}} (D_\ell \cap V_\ell \cap M_\ell)$ is derived on Z using the values in Y , ie,

$$\prod_{\ell_1 \in L_{ij}} f_{\ell_1}(y_{\ell_1}) P_{d\ell_1} \quad (\text{II.3})$$

From the definition of D_{ijk} in equation (21), the probability associated with the target origins in Λ_{ijk} and $\bar{\Lambda}_{ijk}$ on Z is

$$\left[\prod_{\ell \in \Lambda_{ijk}} P_{d\ell} (P_{v\ell} - P_{m\ell}) \right] \prod_{\ell \in \bar{\Lambda}_{ijk}} (1 - P_{d\ell} P_{v\ell}) \quad (II.4)$$

The probability of the remaining $n-t+\delta_{ijk}-\phi_{ij}$ measurements being from clutter and outside the m nearest is

$$\Phi(n-t+\delta_{ijk}-\phi_{ij})(1-P_{m0})^{(n-t+\delta_{ijk}-\phi_{ij})} \quad (II.5)$$

From the definition of the ordered space which is illustrated by the mapping functions in figure 4, there are $(n-t+\delta_{ijk})!/(n-t+\delta_{ijk}-\phi_{ij})!$ branches from the ordered space Y to the unordered space Z . Thus the density on the RHS of equation (II.2) is obtained from the product of this number of branches with expressions (II.3), (II.4) and (II.5) to give,

$$p(Y, \theta_{ijk}, n|T) = a_{ijk} c_{ijk} \left[\prod_{\ell_1 \in L_{ij}} f_{\ell_1}(y_1) P_{d\ell_1} \right] (1-P_{m0})^{(\delta_{ijk}-\phi_{ij})} \quad (II.6)$$

$$\left[\prod_{\ell \in \Lambda_{ijk}} P_{d\ell} (P_{v\ell} - P_{m\ell}) \right] \prod_{\ell \in \bar{\Lambda}_{ijk}} \alpha_{\ell}$$

where

$$a_{ijk} = \frac{(n-t+\delta_{ijk})!}{(n-t+\delta_{ijk}-\phi_{ij})!}$$

$$c_{ijk} = \Phi(n-t+\delta_{ijk}-\phi_{ij})(1-P_{m0})^{n-t} \prod_{\ell=1}^t P_{d\ell} P_{v\ell}$$

$$\alpha_{\ell} = (1-P_{d\ell} P_{v\ell})/P_{d\ell} P_{v\ell}$$

The solution to the β_i 's now follows by substituting equation (II.6) into equation (II.2) and then into equation (II.1). From Assumption No.10, $\Phi(.)$ is constant and therefore c_{ijk} cancels to give the result in equations (31) and (32).

APPENDIX III

EXAMPLE DERIVATION OF THE EVENT PROBABILITIES

This appendix gives the derivation of the event probabilities for the case $m=2$ and $t=2$. Equation (17) contains the origin lists L_{0j} and L_{1j} . The remaining lists L_{2j} follow from L_{1j} . Therefore only the derivation of the probabilities β_0 and β_1 are given below.

For $i=0$, the target origin lists Λ , $\bar{\Lambda}$ and parameters ϕ and δ are given in Table III.1. The values for a_k which are derived from these parameters are also tabulated.

TABLE III.1 ORIGIN LISTS FOR EVENT PROBABILITY β_0

	Λ	$\bar{\Lambda}$	δ	a_k
$L = \{0,0\}$	$\emptyset$	1,2	2	$n(n-1)$
$\bar{L} = \{1,2\}$	2	1	1	$(n-1)(n-2)$
$\phi = 2$	1	2	1	$(n-1)(n-2)$
	1,2	$\emptyset$	0	$(n-2)(n-3)$
$L = \{0,2\}$	$\emptyset$	1	1	$(n-1)$
$\bar{L} = \{1\}$	1	$\emptyset$	0	$(n-2)$
$\phi = 1$				
$L = \{2,0\}$	$\emptyset$	1	1	$(n-1)$
$\bar{L} = \{1\}$	1	$\emptyset$	0	$(n-2)$
$\phi = 1$				

From Table III.1 and equation (32), the expression for b_0 is given by

$$\begin{aligned}
 b_0 = & n(n-1)\alpha_1\alpha_2/v_s^2 + \\
 & (n-1)(n-2)[(1-P_{m1}/P_{v1})\alpha_2 + (1-P_{m2}/P_{v2})\alpha_1]/v_s^2(1-P_{m0}) + \\
 & (n-2)(n-3)(1-P_{m1}/P_{v1})(1-P_{m2}/P_{v2})/[v_s(1-P_{m0})]^2 + \\
 & (n-1)[f_2(y_1) + f_2(y_2)]\alpha_1/(v_c P_{v2}) + \\
 & (n-2)[f_2(y_1) + f_2(y_2)](1-P_{m1}/P_{v1})/[v_c P_{v2}(1-P_{m0})]
 \end{aligned} \tag{III.1}$$

For $i=1$, the lists $\Lambda, \bar{\Lambda}$, parameters ϕ, δ and values for a_κ are listed in Table III.2

TABLE III.2 ORIGIN LISTS FOR EVENT PROBABILITY β_1

	Λ	$\bar{\Lambda}$	δ	a_κ
$L = \{1,0\}$ $\bar{L} = \{2\}$ $\phi = 1$	$\emptyset$	2	1	(n-1)
	2	$\emptyset$	0	(n-2)
$L = \{1,2\}$ $\bar{L} = \{0\}$ $\phi = 0$	$\emptyset$	$\emptyset$	0	1

As for b_0 , from Table III.1 and equation (32), the expression for b_1 is given by

$$\begin{aligned}
 b_1 = & (n-1)f_1(y_1)\alpha_2/(v_c P_{v1}) + \\
 & (n-2)f_1(y_1)(1-P_{m2}/P_{v2})/[v_c P_{v1}(1-P_{m0})] + \\
 & f_1(y_1)f_2(y_2)/(P_{v1}P_{v2})
 \end{aligned} \tag{III.2}$$

For $i=2$, similar results are obtained and the expression for b_2 is

$$\begin{aligned}
 b_2 = & (n-1)f_1(y_2)\alpha_2/(v_c P_{v1}) + \\
 & (n-2)f_1(y_2)(1-P_{m2}/P_{v2})/[v_c P_{v1}(1-P_{m0})] + \\
 & f_1(y_2)f_2(y_1)/(P_{v1}P_{v2})
 \end{aligned} \tag{III.3}$$

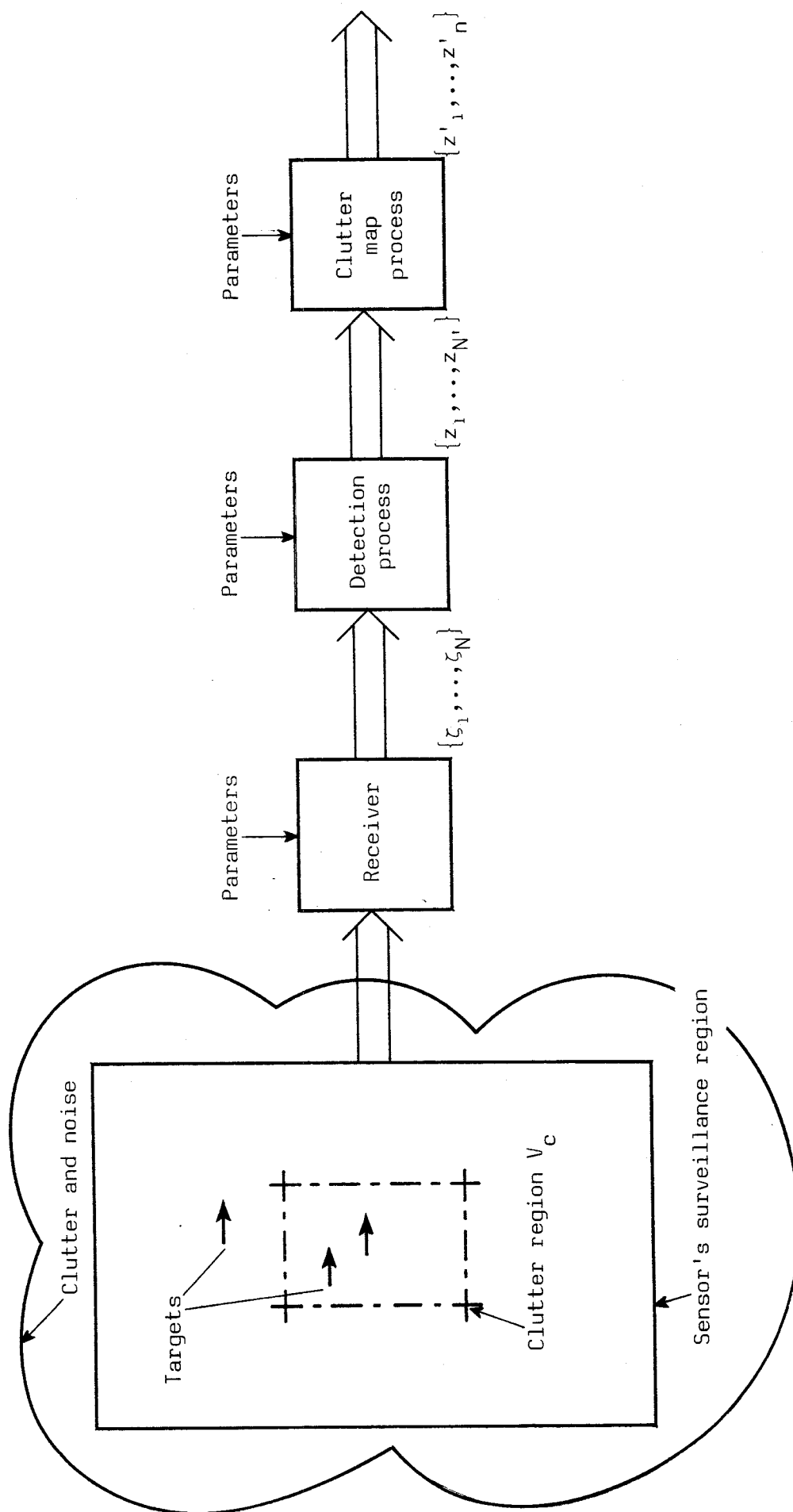


Figure 1. Sensor model to produce measurements in the region V_c

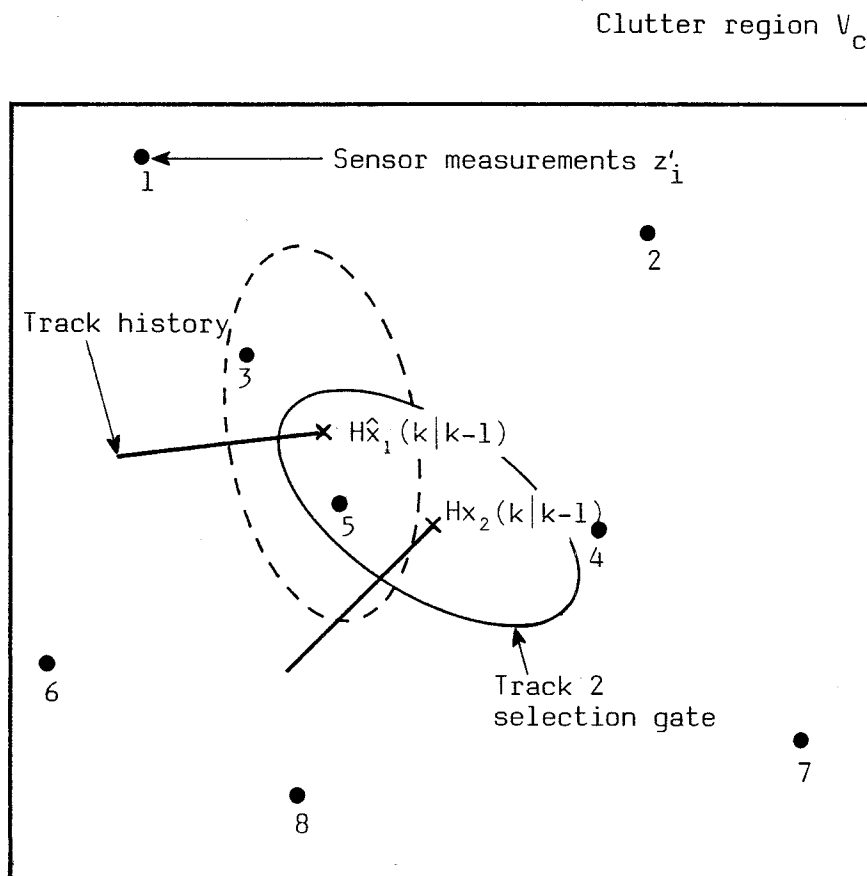


Figure 2. Example of two target tracks with sensor measurements

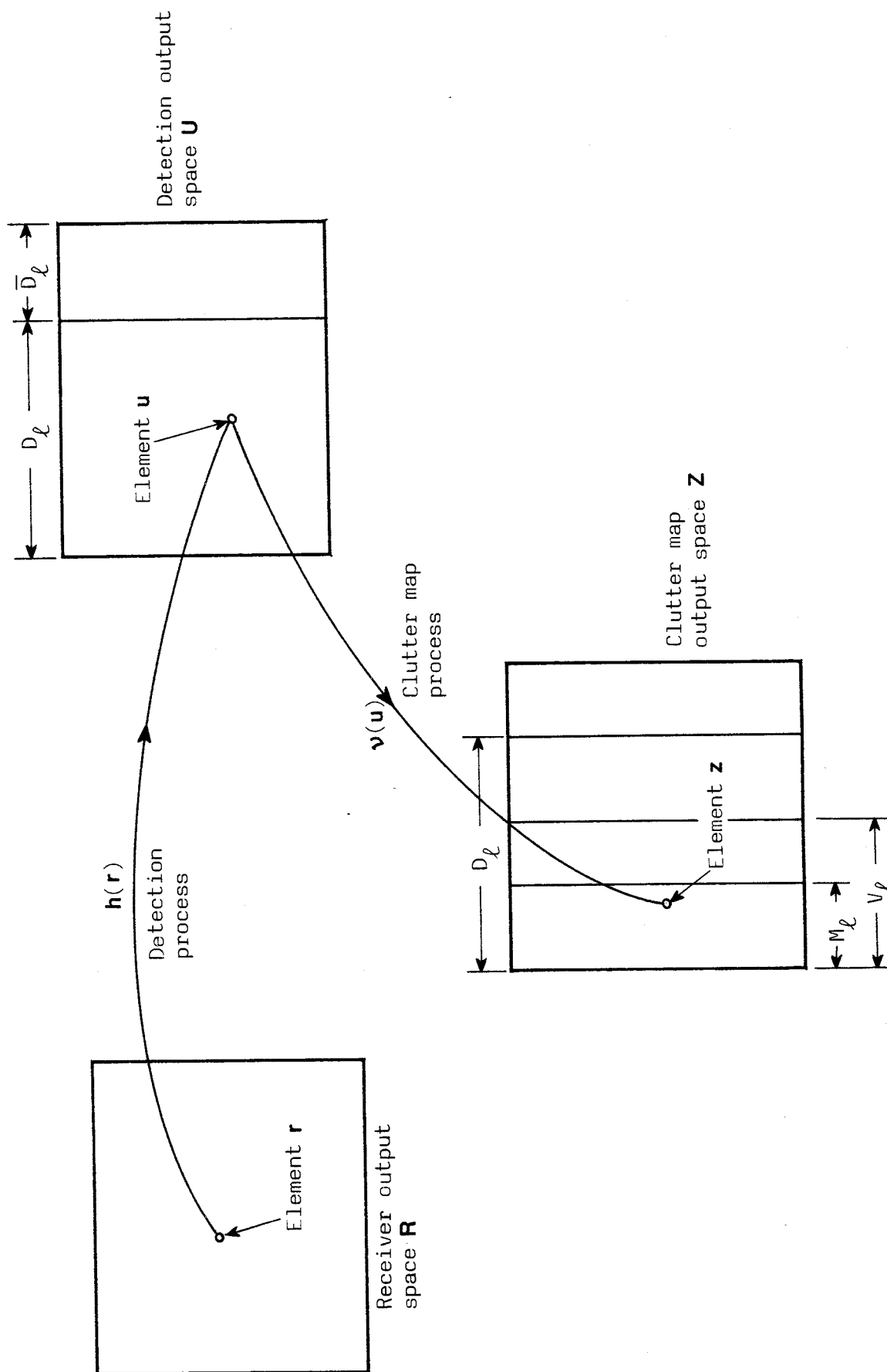


Figure 3. Sample spaces and mapping functions for the sensor model and one target

Sub-partitions W_{ijkh} formed from the permutation of clutter subscripts

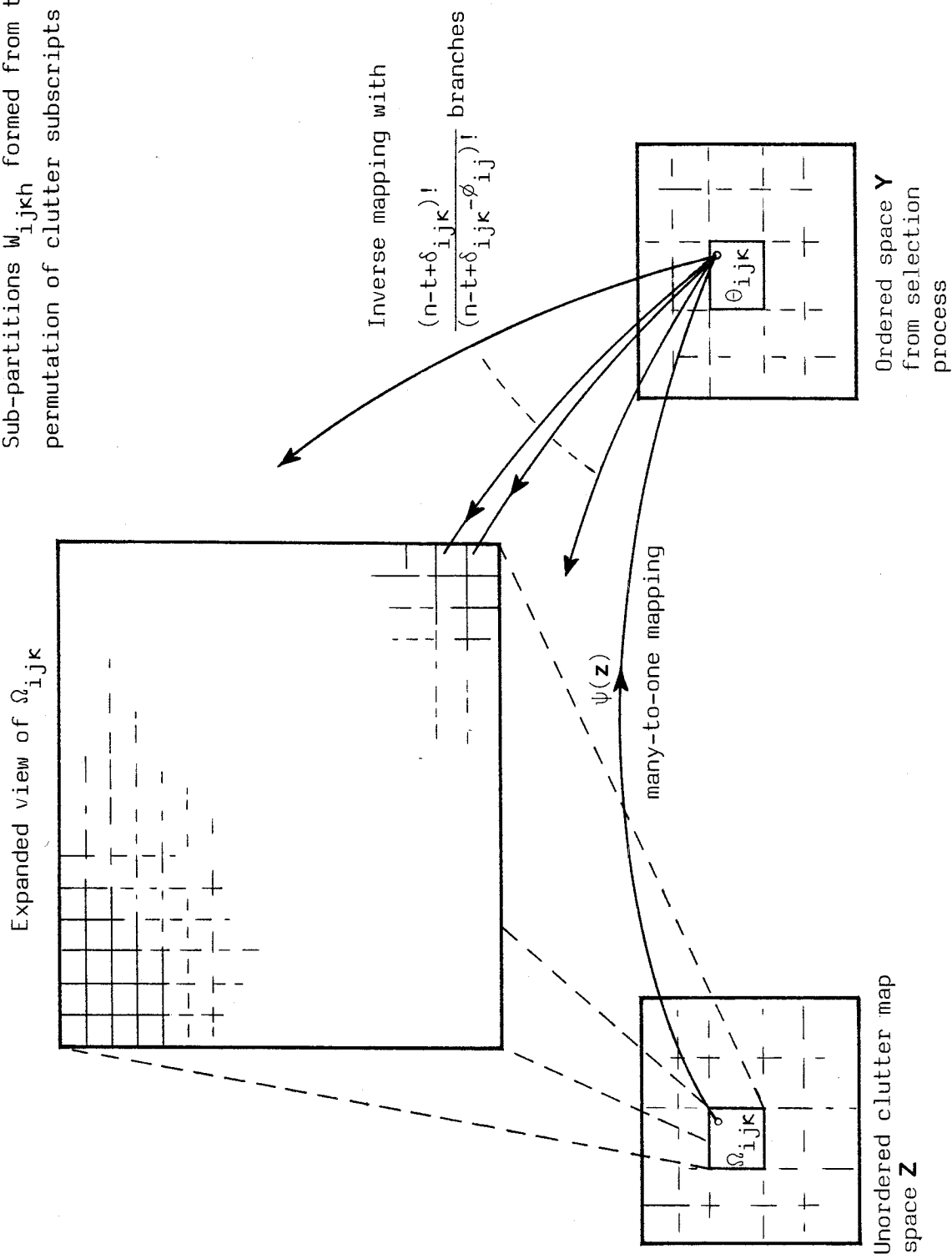


Figure 4. Illustration of mapping between the unordered and ordered spaces

DOCUMENT CONTROL DATA SHEET

Security classification of this page

UNCLASSIFIED

1	DOCUMENT NUMBERS	2	SECURITY CLASSIFICATION
AR Number: AR-004-230		a. Complete Document: Unclassified	
Series Number: ERL-0381-TR		b. Title in Isolation: Unclassified	
Other Numbers:		c. Summary in Isolation: Unclassified	
3	TITLE		
A NEW APPROACH TO MULTITARGET TRACKING USING PROBABILISTIC DATA ASSOCIATION			
4	PERSONAL AUTHOR(S):	5	DOCUMENT DATE:
S.B. Colegrove		September 1986	
6	6.1 TOTAL NUMBER OF PAGES 29		
6	6.2 NUMBER OF REFERENCES: 10		
7	7.1 CORPORATE AUTHOR(S):	8	REFERENCE NUMBERS
Electronics Research Laboratory		a. Task: DEF 77/036	
7.2 DOCUMENT SERIES AND NUMBER		b. Sponsoring Agency: Department of Defence	
Electronics Research Laboratory 0381-TR		9	COST CODE:
		125	
10	IMPRINT (Publishing organisation)	11	COMPUTER PROGRAM(S) (Title(s) and language(s))
Defence Research Centre Salisbury			
12	RELEASE LIMITATIONS (of the document):		
Approved for Public Release			

Security classification of this page:

UNCLASSIFIED

13 ANNOUNCEMENT LIMITATIONS (of the information on these pages):

No limitation

14 DESCRIPTORS:

a. EJC Thesaurus
TermsAutomatic tracking
Kalman filtering
Probability theoryb. Non-Thesaurus
TermsProbabilistic data association
Estimation

15 COSATI CODES:

17110

16 SUMMARY OR ABSTRACT:

(if this is security classified, the announcement of this report will be similarly classified)

This report develops the theory for a multitarget tracking algorithm based on Probabilistic Data Association with new selection rules for assigning sensor measurements to target tracks and for forming multitrack clusters. These new rules remove the requirement to form a gate about each target's predicted position for the selection of sensor measurements. The resultant algorithm is the same for all target tracks and clutter conditions. The algorithm adapts to the sensor measurements via probability terms which model the environment and sensor processing.