

TESIS DEFENDIDA POR
Darío Bonilla Hernández
Y APROBADA POR EL SIGUIENTE COMITÉ

Dr. David Hilario Covarrubias Rosales
Director del Comité

Dr. Roberto Conte Galván
Miembro del Comité

Dr. Hugo Homero Hidalgo Silva
Miembro del Comité

Dr. Luis Armando Villaseñor
González
Miembro del Comité

Dr. Arturo Velázquez Ventura
*Coordinador del programa de
posgrado en Electrónica y
Telecomunicaciones*

Dr. Raúl Ramón Castro Escamilla
Director de Estudios de Posgrado

22 de Agosto de 2005

**CENTRO DE INVESTIGACIÓN CIENTÍFICA Y DE EDUCACIÓN
SUPERIOR**

DE ENSENADA



**PROGRAMA DE POSGRADO EN CIENCIAS
EN ELECTRÓNICA Y TELECOMUNICACIONES**

**MODELADO Y SIMULACIÓN DE LA APROXIMACIÓN DE MÁXIMA
VEROSIMILITUD (ML) INCONDICIONAL EN LA DETERMINACIÓN DEL
DoA EN CAMPO CERCANO**

TESIS

que para cubrir parcialmente los requisitos necesarios para obtener el grado de
MAESTRO EN CIENCIAS

Presenta:

Darío Bonilla Hernández.

Ensenada, Baja California, México, Agosto del 2005.

RESUMEN de la tesis de **Darío Bonilla Hernández**, presentada como requisito parcial para la obtención del grado de MAESTRO EN CIENCIAS en ELECTRÓNICA Y TELECOMUNICACIONES. Ensenada, Baja California. Agosto del 2005.

MODELADO Y SIMULACIÓN DE LA APROXIMACIÓN DE MÁXIMA VEROSIMILITUD (ML) INCONDICIONAL EN LA DETERMINACIÓN DEL DoA EN CAMPO CERCANO

Resumen aprobado por:

Dr. David H. Covarrubias Rosales
Director de Tesis

La detección de la Dirección de Arribo (DoA) de múltiples terminales móviles de banda estrecha, es uno de los temas de mayor investigación en la actualidad dentro de la tecnología de Antenas Inteligentes. Para este propósito se han estudiado distintos métodos, empleando técnicas sub-óptimas y óptimas. Los más utilizados han sido los métodos de subespacio, dadas sus ventajas de rapidez y bajo costo computacional. Pensando en algoritmos que proporcionen una mejor distinción de fuentes cercanas, menor error de estimación y buena eficiencia en ambientes ruidosos, es necesario utilizar nuevos métodos que resuelvan el problema de la localización de fuentes de manera eficiente. Dadas las cualidades de los métodos de inferencia estadística, tales como la consistencia, sesgo y de varianza mínima, es interesante considerar los métodos de Máxima Verosimilitud (ML) para aplicarse al problema de localización de fuentes en sistemas de comunicaciones celulares, ya que pueden ser aplicables a sistemas con señales de banda ancha, con mayor número de usuarios, e incluso bajo ambientes muy agresivos donde se tenga una baja Relación Señal a Ruido (SNR). En este trabajo se realiza el estudio de las propiedades de modelado y simulación de la estimación de la posición de las fuentes de interés bajo la condición de campo cercano por medio de algoritmos de Máxima Verosimilitud Incondicional basados en el método de Maximización de la Esperanza. Así como, también, la evaluación de las bondades del estimador de Máxima Verosimilitud, su modelado y simulación bajo la variación de los parámetros de: Nivel de ruido, número de muestras tomadas, número de fuentes a localizar y la separabilidad entre fuentes cercanas.

A partir de los resultados obtenidos se definirán las ventajas y desventajas del estimador ML, comparado con el algoritmo MUSIC (Multiple Signal Classification) basado en subespacios. Se establecerán las limitaciones principales de ambos métodos y los escenarios bajo los que responden adecuadamente, encontrándose que el estimador de máxima verosimilitud es más robusto, y resuelve el problema de la localización de fuentes de manera más eficiente bajo las consideraciones establecidas.

Palabras clave: Antenas inteligentes, Localización de fuentes, Máxima verosimilitud, Campo cercano

ABSTRACT of the thesis presented by **Darío Bonilla Hernández** as a partial requirement to obtain the MASTER OF SCIENCE degree in ELECTRONICS AND TELECOMMUNICATIONS. Ensenada, Baja California, Mexico. August 2005.

MODELING AND SIMULATION OF UNCONDITIONAL MAXIMUM LIKELIHOOD APPROACH FOR DIRECTION OF ARRIVAL OF NEAR FIELD SOURCES

The Direction of Arrival (DoA) is applied to detection of multiple sources, and is one of the most studied problem in smart antennas technology. For this purpose, there had been studied many methods, based on optimal and sub-optimal techniques. Sub-space based methods had been most popular given their advantages such as speed and easy implementation. Thinking in improve performance algorithms which had better resolution on localization of near field sources, less estimation error, and a good efficiency under high noise scenarios is necessary the use of methods that solves the source localization problem in a more efficient way. Due to many attractive characteristics of the statistical inference methods such as consistence, minimum variance and asymptotic unbiasedness, is interesting to considerate the Maximum Likelihood Estimation methods for the application on the source localization problem on mobile communications systems. Also, they could be applied to wideband signals systems, for multiple sources and for aggressive scenarios with a low Signal to Noise Ratio (SNR).

This thesis studies the main features, model and simulation of the source location estimation under near field condition by Unconditional Maximum Likelihood method, based on Expectation Maximization algorithm. Also, this thesis evaluates the performance, modeling and simulation of this estimator for the variance of the parameters: Noise level, number of snapshots, number of sources, and closely spaced sources.

According to the results obtained we found the advantages and disadvantages of the ML Estimator compared with the MUSIC (Multiple Signal Classification) algorithm, which is a based in sub-space method. For both methods were established the principal limitations and the scenarios for which they work properly. Also we found that the ML estimator is more robust and solves the source localization problem in a more efficient way for the considerations above established.

Key words: Smart antennas, source localization, maximum likelihood estimator, near field.

DEDICATORIA

A mi familia,

**Dario Bonilla Olivares,
Maria del Carmen Hernández Rivera,
Marcos y Sacnicté.**

Por el amor, soporte, enseñanzas y confianza incondicionales que me han dado durante toda mi vida.

Para **Ady**

Por su amor incondicional, y por creer en mi capacidad para lograr mis metas. Te amo.

Para **Hiram , Mary, Aaron y Karen**

Quienes me ayudaron y estuvieron presentes en los momentos alegres y difíciles.

AGRADECIMIENTOS

Quiero agradecer a mi *familia* por el amor que me dieron, por saber que siempre estarían ahí cuando los necesitara, y porque ellos me dieron las bases y el empuje inicial para lograr esta meta.

A mis amigos, aquellos que me apoyaron durante mi estancia y me recibieron en sus corazones, gracias *Mary* e *Hiram*, por adoptarme y hacerme sentir querido, gracias *Aaron* y *Karen* por ser mis amigos en tierras tan lejanas. Gracias *Ady* por reñirme, quererme, hacerme dar más de mí mismo y creer en mí, y porque al final esto logro es para ti.

A todos mis *compañeros de generación*, ya que con su compañía y experiencias compartidas durante la estancia en CICESE fue más fácil el lograr esta meta.

A mi director de tesis y amigo *Dr. David H. Covarrubias Rosales*, por brindarme todo su apoyo, confianza y sugerencias durante el desarrollo de este trabajo y que bajo su dirección fue posible este trabajo.

A los miembros del comité de tesis *Dr. Roberto Conte Galván*, *Dr. Luis A. Villaseñor González* y *Dr. Hugo Hidalgo Silva*, gracias por sus consejos y aportaciones hechas durante el desarrollo de esta trabajo de investigación.

Al *Grupo de Comunicaciones Inalámbricas (GCI)*, en especial a José, por el gran trabajo en equipo que desarrollamos durante este periodo de investigación y que sin su ayuda no hubiera sido posible.

Al *CICSE*, *investigadores y estudiantes que laboran en él* por brindarme la oportunidad de llevar a cabo esta etapa de superación.

Al Consejo Nacional de Ciencia y Tecnología (CONACYT) por la beca otorgada.

Se agradece al CONACYT por el apoyo otorgado al proyecto “*Modelado y Simulación de Algoritmos de la Dirección de Llegada (DOA) y Conformación Digital de Haz (DBF) Aplicados a Comunicaciones Móviles Celulares con Antenas Inteligentes*” con clave U39514-Y, sobre el cual se enmarcó este trabajo de tesis.

Contenido

	Página
Introducción	1
Capítulo I Antecedentes y Marco de referencia de la investigación.....	6
1.1 Introducción	6
1.2 Antenas Inteligentes.....	7
1.2.1 Agrupamiento de antenas.....	10
1.2.2 Modelos espaciales de canal radio.....	12
1.3 El problema de la localización de fuentes	14
1.4 Métodos de estimación de DoA.....	16
1.4.1 Métodos Convencionales.....	16
1.4.2 Métodos de sub-espacio.....	17
1.4.3 Métodos de Máxima Verosimilitud	18
1.4.4 Métodos Integrados.....	19
1.5 Metodología	20
1.6 Conclusiones	21
Capítulo II Método de estimación basado en Máxima Verosimilitud	23
2.1 Introducción	23
2.2 Inferencia estadística.....	23
2.2.1 Principios de inferencia	24
2.2.2 Estimador puntual y propiedades.....	24
2.2.3 Estimadores fundamentales	28
2.3 Método de estimación del DoA basado en Máxima Verosimilitud.....	30
2.3.1 Características de los estimadores de máxima verosimilitud	32
2.3.2 Función de Verosimilitud	33
2.3.3 Maximización de la función de verosimilitud	36
2.3.4 Cota de Cramér Rao.....	37
2.4 Algoritmo de Maximización de la Esperanza.....	39
2.4.1 Inicialización.....	42
2.4.2 Calculo de la Esperanza.....	43
2.4.3 Maximización de la Esperanza	44
2.5 Conclusiones	44
Capítulo III Modelo del sistema	46
3.1 Introducción	46
3.2 Modelo de señal	47
3.2.1 Campo cercano	47
3.2.2 Campo lejano	50
3.3 Modelo matemático estimador.....	52
3.3.1 Función de verosimilitud	52
3.3.2 Algoritmo EM.....	54
3.3.3 Modelo matemático de la Cota de Cramér Rao	58
3.4 Conclusiones	59
Capítulo IV Simulaciones y resultados.....	61
4.1 Introducción	61
4.2 Consideraciones de simulación.....	62
4.3 Estadísticas obtenidas	63

Contenido (continuación)

	página
4.3.1 Simulaciones comparativas con MUSIC	66
4.3.2 Simulaciones del estimador de Máxima verosimilitud.....	70
4.4 Resultados.....	81
4.5 Conclusiones.....	84
Capítulo V Conclusiones	85
5.1 En cuanto al Canal Radio y el Modelo de campo cercano.	85
5.2 En cuanto a los métodos de localización de fuentes.....	87
5.3 En cuanto a la evaluación en las prestaciones del estimador de Máxima Verosimilitud.	89
5.5 Trabajos Futuros.	95
Apéndice A. Método de estimación, algoritmo MUSIC	97
Referencias.....	101

Lista de Figuras

	Página
Figura 1. Antenas inteligentes aplicadas a comunicaciones móviles celulares	9
Figura 2. Filtraje espacial de fuentes de interés.....	9
Figura 3. Configuraciones geométricas de agrupamientos de antenas.	11
Figura 4. Metodología de investigación y desarrollo de la tesis.....	21
Figura 5. Grafica correspondiente a la función de verosimilitud	31
Figura 6. Diagrama de flujo del algoritmo de Maximización de la Esperanza.	42
Figura 7. Modelo de datos, distribución de elementos del agrupamiento de antenas, considerando un frente de onda esférico.....	48
Figura 8. Modelo de datos, distribución de elementos del agrupamiento.	50
Figura 9. Diagrama de flujo del algoritmo EM	56
Figura 10. Comparación de estimación, Algoritmo MUSIC y Estimador ML.....	67
Figura 11. Comparación de comportamiento Estimador ML vs MUSIC.....	69
Figura 12. Convergencia del estimador de máxima verosimilitud bajo ambientes adversos de SNR y distinto número de fuentes	71
Figura 13. Evaluación de número de muestras tomadas para la estimación.....	73
Figura 14. Evaluación del número de fuentes de interés estimadas	75
Figura 15. Análisis de Separación de fuentes	78
Figura 16. Evaluación de separabilidad entre Fuentes	80
Figura 17. Modelo de datos, distribución de elementos del agrupamiento.	98

Lista de Tablas

	Página
Tabla I. Consideraciones numéricas de simulación.....	63
Tabla II. Comparación entre algoritmo MUSIC y estimador ML.	67
Tabla III. Error de estimación en grados para la comparación de ML contra MUSIC ..	69
Tabla IV. Evaluación de Número de Muestras.....	74
Tabla V. Evaluación del número de fuentes de interés estimadas.....	76
Tabla VI. Análisis de Separación de fuentes	78

Introducción

La motivación de este trabajo es el empleo de la tecnología de antenas inteligentes en comunicaciones móviles, y dentro de esta tecnología particularmente aquello que tiene que ver con la localización de las fuentes de datos o usuarios de interés, con lo cual se comprobará los beneficios potenciales del empleo de arreglo de antenas, medido en términos de mayor alcance, mayor diversidad, mayor rechazo a interferentes y transmisión espacial selectiva, en resumen, mejora de la capacidad del sistema [Liberti y Rappaport, 1999; Rappaport, 2002].

El incremento en la capacidad de los sistemas 3G, se puede abordar en base a dos vertientes [Liberti y Rappaport, 1999]:

- Mediante técnicas de acceso al medio (MAC), sea por división de códigos (CDMA), por división de tiempo (TDMA), por división espacial (SDMA), o por híbridos entre estos.
- Mediante las mejoras en las prestaciones en la antena de la estación base.

La detección de la Dirección de Arribo (DoA) de múltiples terminales móviles de banda estrecha, es uno de los temas de mayor investigación en la actualidad dentro de la

tecnología de Antenas Inteligentes. Para este propósito se han estudiado distintos métodos, empleando técnicas sub-óptimas y óptimas.

Para la localización de fuentes en sistemas de comunicaciones móviles, durante un tiempo las más utilizadas fueron las técnicas sub-óptimas basadas en los métodos de subespacio, dadas sus ventajas de rapidez y bajo costo computacional. Pero pensando en algoritmos que proporcionen una mejor distinción de fuentes cercanas, menor error de estimación y buena eficiencia en ambientes ruidosos, es necesario utilizar nuevos métodos que resuelvan el problema de la localización de fuentes de manera eficiente. Dadas las cualidades de los métodos de inferencia estadística, tales como la consistencia, sesgo y de varianza mínima, es interesante considerar los métodos de Máxima Verosimilitud (ML) para aplicarse al problema de localización de fuentes en sistemas de comunicaciones celulares [Dempster et al, 1977]. En particular, presentan ventajas estadísticas y en prestaciones para los nuevos sistemas sobre los métodos anteriores (subespacio) [Miller y Fuhrmann, 1990]. Esto, debido a que pueden ser aplicables a sistemas con señales de banda ancha, con mayor número de usuarios, e incluso bajo ambientes muy agresivos donde se tenga una baja Relación Señal a Ruido (SNR), todas estas variables afectan la eficiencia de estimación de los métodos de subespacio [Van Trees, 2002; Stoica y Nehorai, 1989].

Así, a lo largo de esta tesis, se propone analizar las propiedades de modelado y simulación de la estimación de la posición de las fuentes de interés bajo la condición de campo cercano por medio de algoritmos de ML basados en el método de Maximización de la Esperanza (EM) incondicional. Este método presenta amplias ventajas en comparación con

otros métodos de estimación [Feder y Weinstein, 1988]. Se escogió el campo cercano ya que actualmente la mayoría de los trabajos reportados realizan un análisis en campo lejano y este parámetro ha sido muy poco estudiado, dejando de lado las aplicaciones para las que podría aplicarse al hacerse un estudio de esta naturaleza.

El objetivo de esta tesis es, considerando un entorno celular en comunicaciones móviles:

Modelar y simular la localización de fuentes en campo cercano mediante el método de estimación de máxima verosimilitud incondicional (UML) para la determinación del DoA y parámetros angulares de fuentes cercanas al arreglo de antenas, como alternativa a las limitaciones de métodos de subespacio, bajo ambientes poco estudiados.

Para lograr este objetivo, se evalúan las prestaciones de los métodos de estimación de la dirección de arribo (DoA) de las señales de interés, en un entorno de comunicaciones celulares, bajo condiciones poco estudiadas (campo cercano, ambientes de ruido). Se utilizan arreglos lineales de antenas para los métodos de Inferencia Estadística (Máxima Verosimilitud, ML) y el algoritmo MUSIC (Multiple Signal Classification) [Schmid, 1986], basado en técnicas de subespacio. Para eliminar los problemas de maximización no lineales, se utiliza el método iterativo Maximización de la Esperanza (EM) para el cálculo de la estimación de máxima verosimilitud. Dados estos métodos de estimación del DoA, se evalúa el error de estimación y se obtiene el error cuadrático medio (RMSE) contra la relación señal a ruido (SNR) bajo ambientes adversos como niveles de SNR a partir de -5 dB, entre otros.

Como resultado de esta evaluación se espera demostrar las ventajas del estimador de Máxima Verosimilitud sobre el algoritmo MUSIC en cuanto a resolución en la estimación de fuentes y robustez del algoritmo bajo escenarios, no solo favorables, si no también cuando estos se tornan agresivos y completamente dañinos. Se espera encontrar los límites para ambos métodos de estimación, detectando los ambientes y condiciones bajo los que se desempeñan favorablemente y así establecer las limitaciones, ventajas y desventajas de cada uno.

El utilizar técnicas óptimas supone la posibilidad de encontrar mejoras en la estimación de fuentes. Y el hecho de realizar el estudio bajo condiciones poco estudiadas, aporta un análisis poco realizado y documentado hasta el momento, y abre la puerta a nuevos estudios y la posibilidad de nuevas aplicaciones a partir de estas consideraciones. Se espera aportar nuevas estadísticas, las cuales, por un lado, refuercen las conclusiones a las que se llega a lo largo de todo el trabajo presentado, pero también que logren demostrar las ventajas y desventajas de cada método de estimación y a partir de ahí determinar la viabilidad de ser utilizados en un sistema real.

Para evaluar estos estimadores, es necesario tener conocimientos básicos sobre el tema, durante el primer capítulo se establecerá el marco de referencia necesario para poder entender el problema de la localización de fuentes. Una vez logrado esto tenemos que comenzar a estudiar el método que nos interesa, es decir, a lo largo del segundo capítulo revisaremos la teoría y bases para el estimador de Máxima verosimilitud. Dando paso al

tercer capítulo donde plantearemos el modelo matemático necesario para evaluar las condiciones descritas anteriormente, por último, en orden para terminar este trabajo, en el capítulo cuarto, se presentan las simulaciones y análisis de datos resultantes. Y durante el quinto, las conclusiones y análisis del estimador de Máxima Verosimilitud, basado en el algoritmo EM, aplicado a la localización de fuentes de interés.

Capítulo I Antecedentes y Marco de referencia de la investigación

1.1 Introducción

En este capítulo se planteará el problema al que se avoca esta tesis, estableciendo los conceptos necesarios para la comprensión del estimador que se describe a lo largo de los capítulos siguientes. Para esto, es necesario conocer conceptos, que nos adentren en la tecnología de antenas inteligentes, los factores que determinan su funcionamiento, ventajas y desventajas, y otros fenómenos relacionados a su funcionamiento

Analizaremos, así mismo, el problema de la localización de fuentes de manera breve, principios, limitaciones, alcances y la forma en que se ha intentado resolver este problema. Los métodos utilizados y sus alcances y limitaciones de cada uno de ellos. Todo esto como preámbulo al problema a que se avoca esta tesis.

Por último, pero no menos importante, se establecerá la metodología que se observó para la realización de esta tesis, sobre la cual se ha desarrollado el proceso de investigación, modelado y simulación en que se sustentan nuestras conclusiones.

1.2 Antenas Inteligentes

Antes de entrar en materia de lo que son los estimadores, es importante establecer lo que es un sistema de Antenas Inteligentes, ya que esta tecnología es la que engloba el problema al que se avoca esta tesis y, a que implica consideraciones importantes en el trabajo de investigación y desarrollo de la misma.

Una ANTENA INTELIGENTE es aquella que en vez de disponer de un diagrama de radiación fijo, es capaz de generar o seleccionar haces muy directivos enfocados hacia el usuario deseado [Liberti y Rappaport, 1999]. El uso de este tipo de antenas en la estación base sirve para mejorar tanto la capacidad del sistema, como también la calidad de la señal, incrementa el alcance e introduce la posibilidad de ofrecer nuevos servicios, debido precisamente a esa directividad del haz de radiación, tales como [Buon, 2002]:

- ✓ Controles de potencia más eficientes.
- ✓ Traspasos inteligentes.
- ✓ Mayor área de cobertura.
- ✓ Aumento en la capacidad del sistema.
- ✓ Reducción del nivel de interferencia (BER, QoS).
- ✓ Mejora en la seguridad.

Sin embargo esta tecnología también tiene sus desventajas, las cuales se centran principalmente en el procesamiento realizado para el tratamiento de la señal:

- ✗ Mayor complejidad del sistema.
- ✗ Mayores costos.
- ✗ Métodos más sofisticados de filtraje espacial/temporal.
- ✗ Mejores métodos de procesamiento de la señal.
- ✗ Tecnología en proceso de madurez.
- ✗ Localización de la fuentes

El objetivo de esta tesis se centra principalmente en el último de los problemas: el problema de localización de fuentes o detección de la dirección de llegada (DoA) de múltiples usuarios móviles mediante un arreglo pasivo de antenas. Este problema es de gran importancia, ya que a partir de esta información es como se puede generar una respuesta del sistema adecuada para atender a cada usuario de manera individual y así aumentar la capacidad del sistema [Covarrubias Rosales, 2004], ver Figura 1.

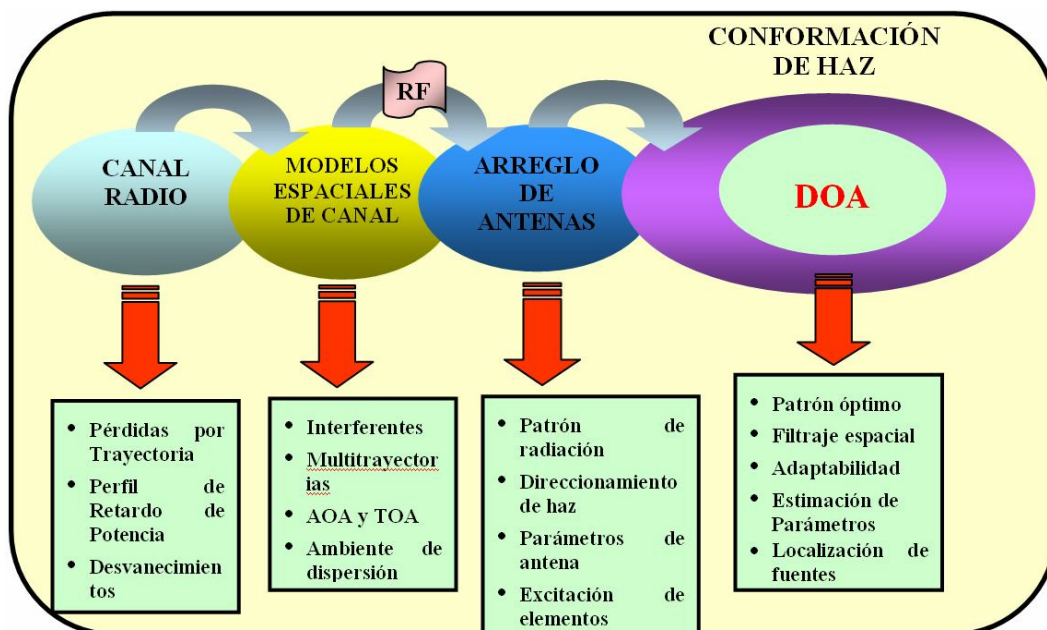


Figura 1. Antenas inteligentes aplicadas a comunicaciones móviles celulares

Un sistema de antenas inteligentes se encuentra constituido por un agrupamiento de antenas y el procesamiento ligado a éste para identificar la ubicación en dirección de cada usuario, y así establecer el filtraje espacial entre usuarios, al dirigir patrones de radiación/recepción a cada uno de ellos en la misma frecuencia y en la misma ranura de tiempo, Figura 2.

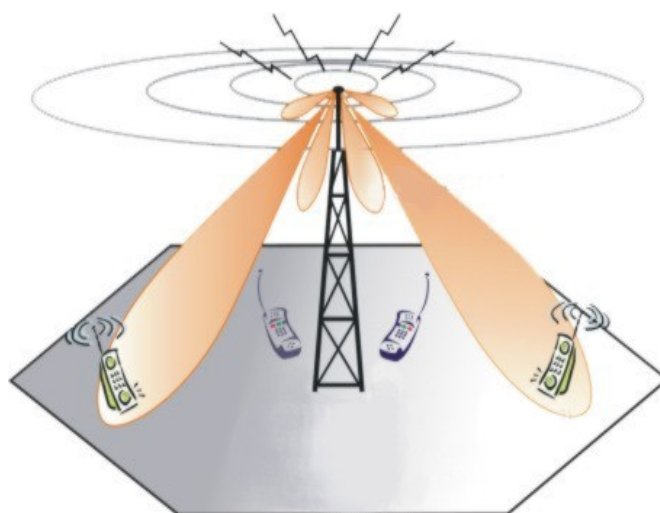


Figura 2 Filtraje espacial de fuentes de interés.

Para comprender esto, primero analizaremos lo que es un agrupamiento de antenas, como se encuentra constituido y sus características principales; los ambientes dispersivos y modelos de canal con los cuales se modela y simula el comportamiento de un sistema de comunicaciones móviles.

1.2.1 Agrupamiento de antenas

Un agrupamiento de antenas se define como un grupo de antenas espacialmente distribuidas que, en vez de disponer de un diagrama de radiación fijo, es capaz de generar o seleccionar haces muy directivos enfocados hacia un usuario deseado. Además, con múltiples antenas en la estación base se obtiene una ganancia en diversidad lo cual presenta varias ventajas importantes en la localización de fuentes [Covarrubias Rosales, 2004].

Los elementos de antena que conforman este agrupamiento se encuentran distribuidos de acuerdo a un cierto patrón geométrico y espaciados entre sí una cierta distancia, medida en términos de longitudes de onda. Y al controlar la amplitud de las corrientes y fase de cada elemento, se puede modificar la forma del diagrama de radiación generando, así, la posibilidad de dirigir los haces de radiación en una dirección determinada. [Panduro y Covarrubias Rosales 2004]

Un agrupamiento de antenas se puede ordenar en diferentes geometrías, dependiendo de ésta, se obtienen características específicas tanto en cobertura angular, resolución,

directividad, etc. Las configuraciones geométricas más comunes de agrupamientos de antenas son: Linear, circular, elíptico y planar [Covarrubias Rosales, 2004] como se muestran en la Figura 3.

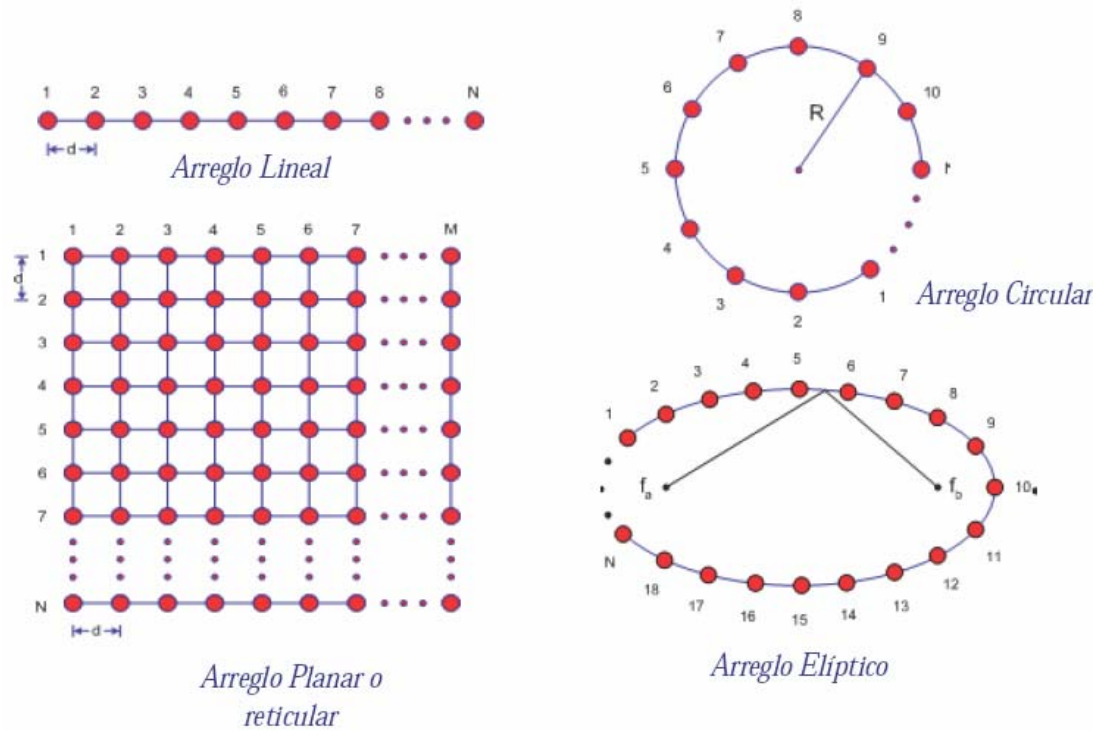


Figura 3. Configuraciones geométricas de agrupamientos de antenas.

En el caso de esta tesis, se trabajó con arreglos lineales, por ser los mas simples de analizar, además de que son con los que se trabaja más comúnmente. El uso de este tipo de arreglos introduce varias ventajas, una de ellas es la diversidad espacial, la cual es posible usar a nuestro favor gracias, precisamente, a tener un número finito de elementos de antena [Covarrubias Rosales, 2004].

La diversidad espacial se obtiene al emplear más de un elemento de antena de recepción, separados una distancia física medida en términos de longitud de onda. Esta diversidad se obtiene cuando se muestrea y combina una señal simultáneamente a partir de diferentes puntos de observación en el espacio. Y que mejor escenario que un agrupamiento de antenas donde se encuentra un número de elementos de antena distribuidos en el espacio y que muestrean y combinan la misma señal provenientes de las fuentes.

Pero al generar el muestreo de estas señales también se están recibiendo réplicas de la señal (multitrayectorias), generadas al rebotar las señales de interés en múltiples dispersores distribuidos dentro del área celular, es por esto que también es importante analizar el canal radio sobre el cual se lleva a cabo la transmisión.

1.2.2 Modelos espaciales de canal radio

Un papel muy importante para la estimación de la posición espacial del móvil dentro de un escenario celular, si se quiere hablar lo más apegado a la realidad posible, lo juega la caracterización del canal radio, ya que el tipo de entorno varía en función tanto de la posición de las fuentes como de los elementos dispersores de señal. Esto requiere generar información de la distribución espacial de los móviles dentro de la célula, lo cual implica generar modelos más precisos relacionados con la distribución de la presencia de posibles elementos dispersivos alrededor del terminal móvil y de la estación base, es decir: Modelos Espaciales de Canal. Para esto se pueden utilizar los modelos espaciales modernos de canal

radio en un entorno dispersivo. Logrando generar información sobre el desfaseamiento de las señales y las multitrayectorias generadas, desvanecimientos, etc.

Estos modelos de canal consideran una distribución estadística de dispersores, basados en una geometría considerando un solo salto. Los modelos pueden variar su geometría dependiendo del entorno (macrocelular, microcelular), adaptando la distribución de los dispersores a las consideraciones de cada entorno [Covarrubias Rosales, 2004]. Dentro de estos modelos, los más importantes son:

- Modelos Geométricos:
 - o Circular
 - o Elíptico.
- Modelos estadísticos:
 - o Gaussiano.

Éste último considera dispersores con una distribución normal [Andrade y Covarrubias Rosales, 2003].

Una vez revisados estos conceptos y de haber asentado el escenario sobre el que se está trabajando, ya que se han establecido las ventajas y desventajas de la tecnología de antenas inteligentes, la utilización de agrupamientos de antenas y lo que esto implica, se puede pasar ahora al problema de interés, es decir, la localización de fuentes.

1.3 El problema de la localización de fuentes

Dentro del procesamiento de señal de un agrupamiento de antenas, el problema de la localización de fuentes, consiste en determinar la dirección de arribo de una señal de interés en presencia de ruido y señales interferentes. Este problema se ha venido estudiando como uno de los problemas más acuciantes dentro de la tecnología de Antenas Inteligentes.

Para resolver este problema, el agrupamiento explota la separación espacial de la localización de las fuentes, es decir la naturaleza de las señales incidentes en el agrupamiento y la diversidad espacial que ofrece el espaciamiento entre elementos [Covarrubias Rosales, 2004]. Para esto se ha venido trabajando con diversos tipos de métodos, entre ellos, los más importantes son los basados en subespacios, los cuales explotan las características de eigenvectores y eigenvalores de los datos muestreados [Schmid, 1986]. Estos algoritmos son rápidos y de fácil cálculo, pero se encuentran limitados por variables como: el número de elementos en el arreglo, el nivel de ruido (SNR), el número de fuentes y su separación, entre las más importantes. Por otro lado se cuenta con los métodos basados en inferencia estadística, los cuales explotan la naturaleza estadística de las señales, y que se han venido utilizando en otro tipo de aplicaciones.

Dadas las limitaciones de los métodos de estimación de DoA basados en subespacio para resolver el problema de localización de fuentes, surge la necesidad de evaluar nuevos métodos de estimación aplicados al problema de localización de fuentes. Así, en esta tesis se propone analizar las propiedades de modelado y simulación de la estimación de la

posición de las fuentes de interés bajo la condición de campo cercano por medio de algoritmos de ML incondicional basados en el método EM bajo condiciones poco estudiadas.

El énfasis que se hace sobre el análisis en campo cercano viene, también motivado, a que la mayor parte de los estudios realizados sobre este tema se han hecho bajo la condición de campo lejano, solo hasta fechas recientes se ha abordado esta condición, además, bajo esta suposición se pueden aplicar a otros tipos de agrupamiento de sensores (micrófonos, sensores sísmicos, etc.) y problemas. Además, la estimación de la posición de fuentes en campo cercano introduce, además de la estimación del ángulo de azimut, la posibilidad de estimar la separación entre estación base y las fuentes de interés (alcance).

Se puede adelantar que el método basado en máxima verosimilitud ofrece ventajas no solo de eficiencia, también estadísticas en cuanto a consistencia, sesgo y mínima varianza y, a las condiciones del entorno celular y del canal radio; Pero esto será analizado durante el siguiente capítulo, y sustentado en el capítulo 4. Primeramente se debe introducir en materia de los métodos de estimación.

1.4 Métodos de estimación de DoA

Los métodos de estimación del DoA se pueden clasificar en 4 grandes grupos [Liberti y Rappaport, 1999]:

- Métodos Convencionales
- Métodos Basados en Subespacios (eigen estructuras)
- Métodos de Máxima Verosimilitud
- Métodos Integrados

Cada uno de estos métodos tiene sus ventajas y desventajas por lo que fueron desarrollados y utilizados. En los siguientes apartados se dará una breve reseña de los principales tipos, a fin de que se conozca las limitaciones y ventajas de cada método de cara al objetivo de la tesis.

1.4.1 Métodos Convencionales

Estos métodos dirigen lóbulos de radiación en todas las direcciones posibles, tratando de identificar máximos de potencia que pudieran representar fuentes activas.

Este tipo de métodos se basan en los conceptos de conformación de haz y no explotan la naturaleza del modelo estadístico de señales. Además requieren de arreglos de antenas con

un gran número de sensores para obtener una buena resolución. Algunos ejemplos son el método de Bartlett y el de Capon [Covarrubias Rosales, 2004; Liberti y Rappaport, 1999].

Estos métodos son muy simples, pero poseen muy baja resolución, además de que son altamente dependientes del número de elementos del arreglo y de la relación SNR (nivel de ruido). También presentan una baja eficiencia en ambientes muy ruidosos o con multitrayectorias, y poca versatilidad.

1.4.2 Métodos de sub-espacio

Los métodos de eigen-estructuras, clasificación a la cual pertenecen MUSIC, ROOT MUSIC, ESPRIT y derivados, siguen las siguientes propiedades de la matriz de correlación del muestreo:

- El espacio abarcado por sus eigen-vectores puede ser particionado en dos subespacios, llamados el subespacio de señal y el subespacio de ruido.
- Los vectores de dirección correspondientes a las fuentes direccionales son ortogonales al subespacio de ruido.

Estos métodos requieren de agrupamientos de antenas con el fin de crear una matriz de correlación empleando las señales recibidas por el arreglo, lo cual implica el conocimiento del vector de dirección [Schmidt, 1986].

Estos métodos han sido los más populares, por sus ventajas de rapidez y exactitud bajo ambientes favorables. Dentro del grupo de investigación que sustenta esta tesis, se han realizado ya trabajos sobre estos métodos. Pero al igual que los métodos anteriores, presentan desventajas de baja eficiencia en ambientes muy ruidosos o con multitrayectorias y son dependientes del número de elementos del arreglo, aunque, cabe aclarar, en menor medida. Como complemento a esta Tesis y al problema en que se avoca, en el apéndice A se tiene una explicación detallada sobre el algoritmo MUSIC de estimación de DoA.

1.4.3 Métodos de Máxima Verosimilitud

Estos métodos se encuentran basados en la teoría de la inferencia estadística, la cual, por medio del comportamiento estadístico de las señales, hace una estimación de la posición angular de las fuentes.

Estos métodos se consideran óptimos, proporcionan una alta resolución, incluso en ambientes ruidosos o con un número de muestras reducido, pero también requieren una gran cantidad de cómputo [Van Trees, 2002].

A diferencia de los métodos de subespacio, los métodos ML presentan una buena eficiencia, aún en casos en que las señales de distintos usuarios pudieran estar correlacionadas, o bajo una relación de SNR baja, lo cual los hace aplicables bajo

ambientes adversos y a señales de banda ancha, y por tanto a los sistemas de comunicación de 3G actuales.

1.4.4 Métodos Integrados

Por último, los métodos híbridos incorporan características de los métodos de subespacio junto con métodos de restauración. Éstos métodos determinan la dirección de arribo a partir de las características espaciales de las señales incidentes en el arreglo de antenas. Estos métodos no solo requieren una gran cantidad de recursos computacionales, también incrementan la complejidad de los sistemas, por tanto su costo [Covarrubias Rosales, 2004; Liberti y Rappaport, 1999].

Una vez revisadas estas alternativas se puede entender el porque se considera a los métodos de subespacio y de ML, como opciones viables para resolver el problema de la localización de fuentes. Dado que los convencionales tienen una eficiencia pobre y los métodos integrados un alto costo tanto de cómputo como del sistema, se descartan estas técnicas.

Ahora, si se compara las características de los métodos de subespacio con las bondades de los métodos ML, se puede esperar que, para ambientes con señales correladas y con bajos niveles de SNR, los métodos que presentarán mejores prestaciones en la estimación del DoA serían los métodos ML.

A continuación, podemos ya establecer la metodología seguida a lo largo de este trabajo y por la cual se ha llegado a los resultados y conclusiones reportados.

1.5 Metodología

La metodología de investigación que se siguió durante el desarrollo de esta tesis se encuentra planteada en la Figura 4.

Como primer paso se hizo un análisis de la problemática de la localización de fuentes, a fin de establecer los alcances y limitaciones del trabajo y de generar las primeras hipótesis a seguir durante este trabajo.

En este trabajo de investigación, primero fue necesario investigar sobre los modelos de canal y el modelo matemático de la señal, es decir los entornos dispersivos, y como se caracterizaría una señal bajo la condición de campo cercano. A continuación se procedió a la investigación y modelado del estimador de máxima verosimilitud.

Ya con el modelo matemático se procedió a la validación del estimador ML. Y por ultimo, a la simulación y obtención de estadísticas que sustentaran este trabajo, y que arrojarán conclusiones concisas.

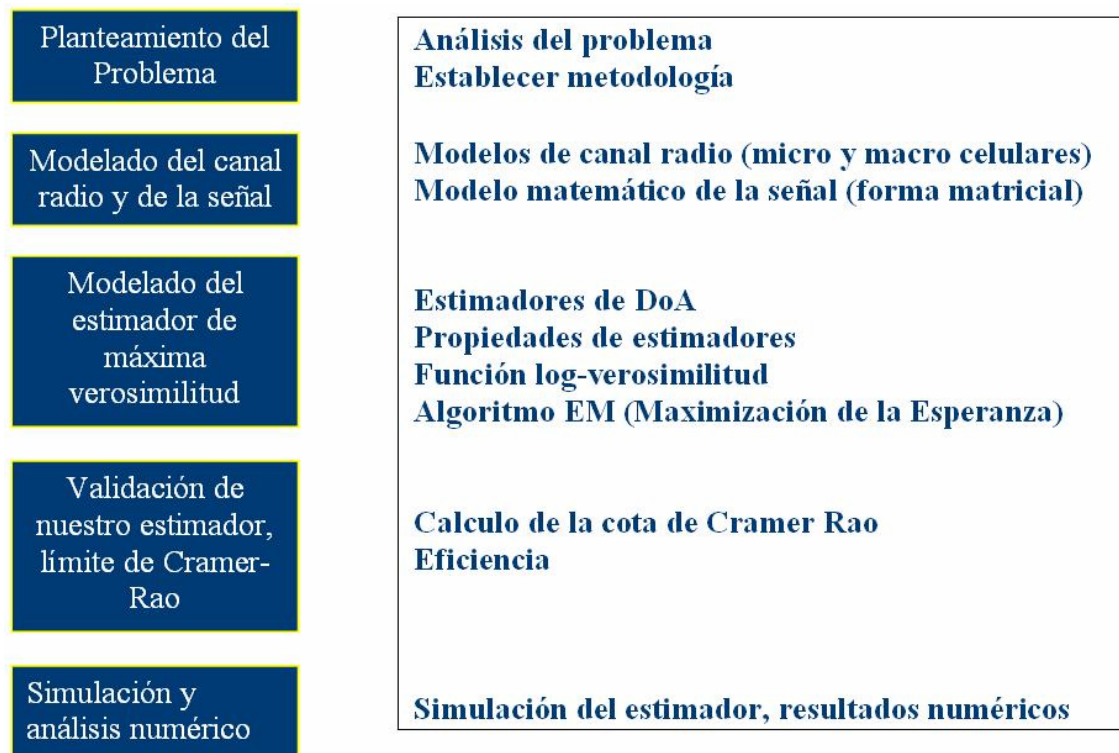


Figura 4. Metodología de investigación y desarrollo de la tesis

Ya establecido cada uno de estos pasos, se pudieron generar los resultados y conclusiones de esta tesis, las cuales se reportaran a lo largo del último capítulo.

1.6 Conclusiones

Una vez hecho el planteamiento del problema, establecidos los principios sobre los cuales se basa el problema de la estimación de fuentes, los conceptos básicos que aseguren un mejor entendimiento de la tecnología de antenas inteligentes, y de los estimadores mencionados en esta tesis, es momento de dar paso a la investigación que se ha realizado. Lo que a continuación se plantea, es la teoría necesaria para entender el estimador de

Máxima Verosimilitud, principal herramienta de esta investigación. Así pues, durante el siguiente capítulo se explicarán, de manera detallada, las características, principios y funcionamiento de este estimador, para posteriormente dar paso al modelo matemático y a las simulaciones y análisis de datos generados.

Capítulo II Método de estimación basado en Máxima Verosimilitud

2.1 Introducción

El objetivo de este capítulo es presentar la teoría de la inferencia estadística incluyendo lo que es un estimador, características y algunos de los métodos más comunes, para posteriormente analizar en detalle la teoría del estimador ML. Plantear los pasos que lo componen y su funcionamiento. Además, por ser una de las partes más importantes de este estimador, también se dedicará un apartado para estudiar el algoritmo de maximización de la esperanza.

2.2 Inferencia estadística

Antes de entrar en materia del estimador ML, es necesario establecer la naturaleza a la que éste obedece, la inferencia estadística. La cual se puede definir como sigue: La inferencia estadística persigue la obtención de conclusiones sobre un gran número de datos, basándose en la observación de una muestra obtenida de ellos; también intenta medir su significancia, es decir la confianza que nos merecen [Van Trees, 2002].

2.2.1 Principios de inferencia

Dentro del marco de la identificación de sistemas y la estimación de parámetros, se trabaja con el problema de extraer información a partir de observaciones que pueden ser no fiables ya que pueden estar corrompidas por factores externos. Cuando estas perturbaciones se modelan como procesos aleatorios, las observaciones pueden describirse como realizaciones de variables aleatorias [Van Trees, 2002]:

$$x^N = (x(1), x(2), \dots, x(N)), \quad (1)$$

las cuales toman valores en R^N , suponiendo, también, que x^N tiene una función de densidad de probabilidad (pdf) asociada. Dentro de este contexto, se considera el estadístico o estimador como una variable aleatoria con una determinada distribución, la cual será la pieza clave en las dos amplias categorías de la inferencia estadística: la estimación y el contraste de hipótesis.

Pero antes de entrar más en detalle, los primeros términos, obligados, a los que se debe hacer referencia, serán los de estadístico y estimador.

2.2.2 Estimador puntual y propiedades

El concepto de estimador, como herramienta fundamental, se caracteriza mediante una serie de propiedades que nos sirven para elegir el “mejor” para un determinado parámetro

de una población, así como algunos métodos para la obtención de ellos, tanto en la estimación puntual como por intervalos [Kay, 1993; Van Trees, 2002].

Desde el punto de vista de la estadística clásica, se considera a los parámetros de una población como cantidades fijas, las cuáles es posible *estimar* mediante estadísticos que son función de los datos obtenidos en una o más muestras y que no contienen cantidades desconocidas. A estos estadísticos se les conoce como estimadores del parámetro de interés [Kay, 1993; Van Trees, 2002]. Por ejemplo, la media muestral es un estimador de la media poblacional, la proporción observada en la muestra es un estimador de la proporción en la población. Para nosotros es importante definir lo que es el estimador puntual, ya que lo que se busca, en esta tesis, es un solo valor estimado del parámetro.

Estimador Puntual es cualquier estadístico que se utiliza para inferir un parámetro desconocido de la población y que proporciona un único valor como estimación de dicho parámetro [Kay, 1993; Van Trees, 2002]. En general se representará con θ a un parámetro cualquiera en el que estamos interesados, a su estimador con $\hat{\theta}$ y al estimado o estimación con θ' (la notación puede variar en algunos casos).

Los criterios más usuales para evaluar un estimador son:

- Varianza
- Sesgo
- Eficiencia
- Consistencia
- Suficiencia

De manera más específica, los términos mencionados en la lista anterior consisten en:

Varianza: es el segundo momento con respecto a la media del estimador.

Sesgo: El valor medio que se obtiene de la estimación para diferentes muestras debe ser el valor del parámetro.

Se dice que un estimador $\hat{\theta}$ de un parámetro θ es **insesgado** si:

$$E[\hat{\theta}] = \theta . \quad (2)$$

La carencia de sesgo puede interpretarse del siguiente modo: Suponiendo que se tiene un número indefinido de muestras de una población, todas ellas del mismo tamaño n . Sobre cada muestra el estimador ofrece una estimación concreta del parámetro que buscamos. Pues bien, el estimador es insesgado, si sobre dicha cantidad indefinida de estimaciones, el valor medio obtenido en las estimaciones es θ (el valor que se desea conocer).

Eficiencia: Al estimador, al ser una variable aleatoria, no puede exigírsele que para una muestra cualquiera se obtenga como estimación el valor exacto del parámetro. Sin embargo, podemos esperar que su dispersión con respecto al valor central (varianza) sea tan pequeña como sea posible.

Dados dos estimadores $\hat{\theta}_1$ y $\hat{\theta}_2$ de un mismo parámetro θ , diremos que $\hat{\theta}_1$ es más **eficiente** que $\hat{\theta}_2$ si :

$$Var[\hat{\theta}_1] < Var[\hat{\theta}_2] . \quad (3)$$

Consistencia: Cuando el tamaño de la muestra crece arbitrariamente, el valor estimado se aproxima al parámetro desconocido.

Decimos que $\hat{\theta}$ es un **estimador consistente** con el parámetro θ si:

$$\forall \varepsilon > 0, \lim_{n \rightarrow \infty} P[|\hat{\theta} - \theta| > \varepsilon] = 0, \quad (4)$$

o lo que es equivalente:

$$\forall \varepsilon > 0, \lim_{n \rightarrow \infty} P[|\hat{\theta} - \theta| < \varepsilon] = 1. \quad (5)$$

Este tipo de propiedades definidas cuando el número de observaciones n , tiende a infinito, es lo que se denominan *propiedades asintóticas* [Kay, 1993; Van Trees, 2002].

Suficiencia: Esta propiedad indica que el estimador debería aprovechar toda la información existente en la muestra. Diremos, entonces, que $\hat{\theta} \equiv \hat{\theta}(X_1, \dots, X_n)$ es un *estimador suficiente* del parámetro θ si la probabilidad P

$$P[X_1 = x_1, X_2 = x_2, \dots, X_n = x_n | \theta = a] \quad (6)$$

donde: a = cualquier valor posible de θ

no depende de θ para todo posible valor de θ [Kay, 1993; Van Trees, 2002].

Esta definición, así enunciada, tal vez resulte un poco oscura, pero lo que expresa es que un estimador es suficiente si agota toda la información existente en la muestra que sirva para estimar el parámetro.

Tengamos en cuenta que el estimador **no es un valor concreto** sino una variable aleatoria, ya que aunque depende unívocamente de los valores de la muestra observados ($X_i=x_i$), la elección de la muestra es un proceso aleatorio. Una vez que la muestra ha sido elegida, se denomina **estimación** el valor numérico que toma el estimador sobre esa muestra x .

Existen técnicas para encontrar estimadores de algún parámetro, las más utilizadas son el método de máxima verosimilitud y el método de los momentos, los cuales se describirán brevemente en la siguiente sección.

2.2.3 Estimadores fundamentales

Ya que se han establecido las propiedades deseadas en un estimador, ahora se revisarán las características de ciertos estimadores que por su importancia en las aplicaciones resulta fundamental mencionar: estimadores de la esperanza matemática y varianza de una distribución de probabilidad, y el estimador insesgado de varianza mínima.

Estimador de la esperanza matemática

Consideremos las muestras de tamaño n , $x_1, x_2, \dots, x_n$, de un carácter sobre una población que viene expresado a través de una v.a. X que posee momentos de primer y segundo orden, es decir, existen $E[X]$ y $Var[X]$:

$$x_1, x_2, \dots, x_n, \begin{cases} E[X] = \mu \\ Var[X] = \sigma^2 \end{cases} \quad (7)$$

El estimador *media muestral* que denotaremos normalmente como $\bar{X}$ es:

$$\bar{X} = \frac{1}{n}(x_1 + x_2 + \dots + x_n). \quad (8)$$

Como podemos ver, este es el más simple de los estimadores.

Estimador insesgado de varianza mínima uniforme

Si existe un estimador insesgado $\hat{\theta}$ para el parámetro θ , tal que la varianza de $\hat{\theta}$ sea la menor de entre todos los estimadores insesgados de θ , para cualquier valor de θ , entonces $\hat{\theta}$ se denomina el estimador de varianza mínima uniforme.

El Estimador de Máxima Verosimilitud (MLE: Maximum Likelihood Estimator), además de ser el más versátil, ya que se puede aplicar en gran cantidad de situaciones, y por ello uno de los más empleado, es también una alternativa al estimador insesgado de mínima varianza (MVUE: minimum variance unbiased estimator). Ya que para muchos problemas de estimación el estimador MVUE no existe, y cuando exista no está asegurado un procedimiento sistemático para encontrarlo. En el siguiente apartado nos ocuparemos del MLE de manera más detallada.

2.3 Método de estimación del DoA basado en Máxima Verosimilitud

El método de ML se basa en el principio de verosimilitud, el cual establece que toda la información que se tiene acerca de un parámetro desconocido θ está contenida en la función de verosimilitud.

Principio de verosimilitud: Sea X un conjunto de muestras aleatorias, la información traída por una observación x acerca de θ , está completamente contenida en la función de verosimilitud $l(X)$. Sin embargo, si x_1 y x_2 son dos observaciones dependientes del mismo parámetro θ , tal que exista una constante c que satisfaga $l(x_1) = c l(x_2)$ para cada θ , entonces traerán la misma información acerca de θ y conducirán a estimadores idénticos [Van Trees, 2002; Liberti y Rappaport, 1999].

Para explicar mejor este principio y también el estimador, consideremos al conjunto X de muestras aleatorias, las cuales tienen una función de densidad de probabilidad $f(x; \theta)$. Entonces, las muestras aleatorias simples, de tamaño n , $x_1, x_2, \dots, x_n$, tienen una función de distribución de probabilidad conjunta f_c de la siguiente manera:

$$f_c(x_1, x_2, \dots, x_n; \theta) = f(x_1; \theta) f(x_2; \theta) \dots f(x_n; \theta). \quad (9)$$

Esta función se puede considerar de dos maneras:

- Fijando θ , es una función de las n cantidades x_i . A ésta se le llama función de densidad de probabilidad.

- Fijando los x_i , como consecuencia de los resultados de elegir una muestra mediante un experimento aleatorio, es únicamente función de θ . A esta función de θ la denominamos **función de verosimilitud** $l(\theta)$.

La función de verosimilitud se obtiene a partir de la función de densidad, intercambiando los papeles entre parámetro y estimador. En una función de verosimilitud consideramos que las observaciones $x_1, \dots, x_n$, están fijas, y se representa en la gráfica la función de densidad de probabilidad para todos los posibles valores del parámetro θ , Figura 5. El estimador de máxima verosimilitud del parámetro buscado, $\hat{\theta}_{MV}$, es aquel que maximiza su función de verosimilitud, $l(\theta)$.

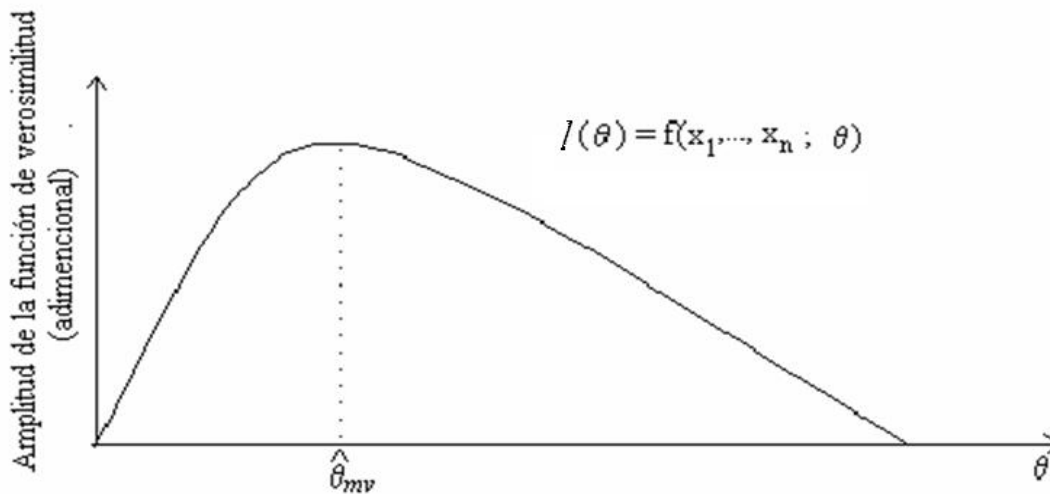


Figura 5. Gráfica correspondiente a la función de verosimilitud

De modo más preciso, se define el **estimador máximo verosímil** como la v.a.

$$\hat{\theta}_{MV} = \max_{\theta \in \mathcal{R}} f(x_1, x_2, \dots, x_n; \theta). \quad (10)$$

Una vez establecido el principio de verosimilitud, y estimador de máxima verosimilitud, se puede proseguir con el análisis de las características de este método, así como a las partes que lo conforman.

2.3.1 Características de los estimadores de máxima verosimilitud

Los estimadores de máxima verosimilitud tienen, al igual que otros estimadores puntuales, propiedades que son deseadas y que proporcionan el grado de utilidad de éstos. A continuación se enuncian las más importantes [Kay, 1993; Van Trees, 2002].

1. Son consistentes
2. Son invariantes frente a transformaciones biunívocas, es decir, si $\hat{\theta}_{MV}$ es el estimador máximo verosímil de θ y $g(\theta)$ es una función biunívoca de θ , entonces $g(\hat{\theta}_{MV})$ es el estimador máximo verosímil de $g(\theta)$.
3. Si $\hat{\theta}$ es un estimador suficiente de θ , su estimador máximo verosímil, $\hat{\theta}_{MV}$ es función de la muestra a través de $\hat{\theta}$.
4. Son asintóticamente normales.
5. Son asintóticamente eficientes, es decir, entre todos los estimadores consistentes de un parámetro θ , los de máxima verosimilitud son los de varianza mínima.
6. Son insesgados si su valor esperado es igual al parámetro que pretende estimar.

Una forma de determinar la varianza del estimador insesgado de varianza mínima la proporciona un resultado que asegura que la varianza de cualquier estimador insesgado siempre será mayor o a lo sumo igual que cierto límite conocido como *cota inferior de Cramér-Rao* (CCR).

Una vez que hemos comprobado que nuestro estimador cumplirá con las características deseadas, podemos dividirlo en tres fases, las cuales son de suma importancia, tanto para el cálculo, la obtención y la verificación de nuestro estimador:

1. Obtención de la función de verosimilitud
2. Maximización de dicha función mediante técnicas iterativas
3. Cálculo de la cota de Cramér-Rao

En los siguientes apartados analizaremos cada una de estas fases por separado, pero sin olvidar que juntas nos llevarán a obtener el estimador deseado.

2.3.2 Función de Verosimilitud

Como ya se había mencionado, el estimador lleva asociada una función de verosimilitud. Esta función de verosimilitud se puede definir como sigue: Sea $\mathbf{X}^n = (x_1,$

$x_2, \dots, x_n$) una muestra aleatoria de la variable aleatoria X , entonces la función de densidad de probabilidad conjunta para el vector aleatorio a ser observado está dada por:

$$f(\theta; x_1, x_2, \dots, x_N) = f_X(\theta; x^n). \quad (11)$$

Esta función es la llamada Función de Verosimilitud y refleja la “probabilidad de que el evento observado ocurra”. Sería razonable entonces, seleccionar un estimador de θ de forma tal que los eventos observados resulten “tan probables como sea posible”. Esto es, busquemos el valor que:

$$\hat{\theta}_{ML}(x^N) = \arg \max_{\theta} f_X(\theta; x^n). \quad (12)$$

En una muestra aleatoria cada observación x_i tiene la misma distribución de X y además son independientes entre si mismas (i. i. d. independientes e idénticamente distribuidas). Dado este principio, la expresión (11) queda como:

$$l(\theta) = l(\theta|X) = f(x_1|\theta) \bullet \dots \bullet f(x_n|\theta) = \prod_{i=1}^n f(x_i|\theta). \quad (13)$$

Se puede observar que esta función no es nada nuevo: simplemente es una forma de reescribir la pdf conjunta de X . Pero el tener un multiplicatorio puede complicar las cosas al momento de obtener el estimador de (12).

Para resolver esto generalmente se utiliza el logaritmo de la función, ya que es equivalente a maximizar una función a maximizar su logaritmo (al ser esta una función estrictamente creciente), entonces se trabaja con:

$$L(\theta|X) = \log(l(\theta|x)) = \log\left(\prod_{i=1}^n f(x_i|\theta)\right) = \sum_{i=1}^n \log(f(x_i|\theta)). \quad (14)$$

Es más sencillo trabajar con expresiones de este tipo, facilitando el cálculo del estimador. Lo que resta ahora es obtener el valor que maximiza esta función, el cual puede calcularse derivando con respecto a θ la función de verosimilitud (en este caso su logaritmo) y tomando como estimador de máxima verosimilitud al que haga la derivada nula:

$$\frac{\partial}{\partial \theta_i} [L(\theta)]_{\theta=\hat{\theta}_{ML}} = 0 \quad i = 1, 2, \dots, D. \quad (15)$$

En ciertos casos, tales como pdfs de forma exponencial, maximizar esta expresión puede ser fácilmente resuelto empleando cálculo y álgebra lineal convencional, pero generalmente este proceso es complicado, resultando en problemas de optimización no lineales, sobre todo cuando D crece o cuando no existe la expresión final o es más compleja de evaluar, para resolver este problema se utilizan técnicas eficientes para resolver el problema de la maximización. Los métodos utilizados son algoritmos iterativos, los cuales aunque representan un mayor costo computacional, aseguran la convergencia hacia un valor. En el siguiente apartado revisaremos algunos de ellos dando pie a la sección 2.4 donde explicaremos más a fondo el método que se escogió para aplicarse en este trabajo.

2.3.3 Maximización de la función de verosimilitud

Aunque se pueden utilizar cálculo y álgebra convencionales en casos donde la función de verosimilitud es simple, para nuestro caso no es tan viable dada la naturaleza de las funciones utilizadas. Entonces se pueden emplear técnicas numéricas avanzadas de maximización tales como:

- Iteración Newton-Raphson
- Iteración por métodos cuadrados
- Proyección Alternativa
- Algoritmo Maximización de la Esperanza

Dada la naturaleza de los datos que se reciben en el agrupamiento de antenas, donde existen pérdidas, interferencia y ruido, lo cual origina pérdida de datos a lo largo de nuestro muestreo, además del tiempo limitado para realizar esta medición; se hace necesario elegir un método apropiado para la maximización. Los métodos de Newton-Raphson e Iteración por métodos cuadrados, aunque son de los más sencillos, no siempre llegan a expresiones que se puedan utilizar, además que no aseguran, bajo las condiciones anteriormente mencionadas, un buen funcionamiento.

Los métodos de proyección alternativa (AP) y maximización de la esperanza fueron diseñados para trabajar en situaciones donde existen datos incompletos [Van Trees, 2002].

Ambos han sido aplicados a este tipo de problemas, pero el que genera una mejor respuesta es el algoritmo de EM [Ziskind y Wax, 1998].

Dada su importancia y complejidad, en la sección 2.4 explicaremos más a detalle éste algoritmo, los pasos que lo conforman y la base de su funcionamiento.

2.3.4 Cota de Cramér Rao

Una vez encontrado el estimador, es necesario verificar que es eficiente; ya que no podemos estar comparándolo contra otros estimadores, verificamos su eficiencia a través de la cota de Cramér-Rao (CCR). El estimador de máxima verosimilitud es insesgado. Un estimador insesgado se dice que es eficiente desde el punto de vista estadístico si la covarianza de la estima alcanza la CCR [Van Trees, 2002].

Al calcular la CCR, puede conocerse al menos la varianza del estimador insesgado de varianza mínima, y así estar en posibilidades de encontrar ese estimador; en lugar de buscar cualquier estimador, se buscará únicamente aquel cuya varianza satisface la CCR, con la seguridad de que no habrá otro estimador insesgado mejor. Por otro lado, si ya se cuenta con un estimador insesgado, al comparar su varianza con la cota inferior de Cramér-Rao, puede saberse si éste es el mejor estimador insesgado o es posible que exista otro mejor.

Frecuentemente, también, resulta difícil el cálculo exacto de las propiedades estadísticas de un estimador, como por ejemplo la matriz de covarianza del error de estimación. En lugar de esto, se suele realizar un análisis asintótico cuando el número de muestras N tiende a infinito. Para analizar el comportamiento asintótico de la eficiencia del estimador también se utiliza la CCR como un indicador de su desempeño estadístico.

La cota Cramer-Rao para la matriz de varianza de un estimador insesgado de θ está dada por:

$$\text{Var}(\hat{\theta}) \geq \mathbf{J}^{-1}(\theta). \quad (16)$$

donde $\mathbf{J}$ es la *Matriz de Información de Fisher* (FIM) dada por:

$$J(\theta) = E \left(\left[\frac{\partial}{\partial \theta} \log f(X; \theta) \right]^2 \right) = -E \left[\frac{\partial^2}{\partial \theta^2} \log f(X; \theta) \right]. \quad (17)$$

La información de Fisher (IF) describe la cantidad de datos de información provista acerca de un parámetro desconocido. La principal aplicación de la IF es para determinar la varianza y el comportamiento asintótico de un estimador. Además, la IF es un concepto clave en la teoría de inferencia estadística y se define de la siguiente manera: Sea $\mathbf{X}=(x_1, \dots, x_n)$ una muestra aleatoria, y sea $f(\mathbf{X}|\theta)$ la función de densidad para algún modelo de datos, el cual tiene como vector parámetros $\theta=(\theta_1, \dots, \theta_k)$. Entonces la matriz de IF $\mathbf{I}_n(\theta)$ de una muestra de tamaño n viene dada por la matriz simétrica $k \times k$ cuyo ij -ésimo elemento

viene dado por la covarianza entre la primera derivada parcial de la función log verosimilitud:

$$I_n(\theta)_{i,j} = \text{Cov} \left[\frac{\partial \ln f(X|\theta)}{\partial \theta_i}, \frac{\partial \ln f(X|\theta)}{\partial \theta_j} \right]. \quad (18)$$

La cota inferior de Cramér Rao aplicada al problema de interés produce una matriz J donde los elementos se dan de la siguiente manera:

$$J_{ij} = -E \left[\frac{\partial^2}{\partial \theta_i \partial \theta_j} [L_x(\theta)] \right] = K \text{tr} \left[K_x^{-1}(\theta) \frac{\partial K_x(\theta)}{\partial (\theta_i)} K_x^{-1}(\theta) \frac{\partial K_x(\theta)}{\partial (\theta_j)} \right] + 2 \text{Re} \left[\frac{\partial m^H(\theta)}{\partial (\theta_i)} K_x^{-1}(\theta) \frac{\partial m(\theta)}{\partial (\theta_j)} \right], \quad (19)$$

como la media es cero

$$J_{ij} = K \text{tr} \left[K_x^{-1}(\theta) \frac{\partial K_x(\theta)}{\partial (\theta_i)} K_x^{-1}(\theta) \frac{\partial K_x(\theta)}{\partial (\theta_j)} \right].$$

donde: $m(\theta)$ es la media del conjunto de muestras X , K denota el número de muestras tomadas, tr denota la traza de la matriz A y Re la parte real de A [Van Trees, 2002].

En la mayoría de los casos el parámetro θ contiene, no sólo los parámetros de interés, sino que puede contener parámetros indeseados o el caso en que los parámetros deseados son de distinta índole. Un ejemplo de esto es cuando hay información sobre la varianza del ruido o la potencia de la señal.

2.4 Algoritmo de Maximización de la Esperanza

El método de Maximización de la Esperanza (EM) como ya se dijo, es un método iterativo diseñado para casos donde existen datos incompletos. El significado más general

de este concepto implica la existencia de dos espacios muestrales y una función inyectiva de X a Y : los datos observados, los cuales se encuentran deformados por los efectos del canal (ruido, multitrayectorias, pérdida de datos, etc.) y se consideran como incompletos. Y el segundo espacio de datos completos desconocidos, es decir los datos originales, que no sufren alteraciones pero que no los podemos conocer.

- $X_k = \{x_1, \dots, x_k\}$ *datos observados incompletos*
- $Y_k = \{y_1, \dots, y_k\}$ *datos completos desconocidos*

De los datos completos Y desconocidos que han generado a X , solo se sabe que pertenecen a la preimagen de X por la aplicación $Y(x)$.

En un escenario, como en el que nos situamos, con D fuentes de interés, se tienen D problemas de maximización. Aquí nosotros consideramos la señal proveniente de cada una de las fuentes $Y_i(t)$, que incide en el arreglo de forma similar a los datos observados, incompletos, considerándolos como datos sin interferencia ni pérdidas, pero desconocidos.

El modelo correspondiente a los datos completos es: $f(x|\theta)$. Como x no es observable, pero sí lo es una función $y = y(x)$, entonces el modelo correspondiente a los datos observados es:

$$f(y|\theta) = \int X(y)f(x|\theta)dx. \quad (20)$$

Y podemos establecer que existe una relación entre los datos incompletos $X(k)$ y los datos completos $Y_l(k)$ a través de una transformación lineal, definida en (21).

$$X(k) = \sum_{l=1}^D Y_l(k), \quad k = 1, \dots, K \quad (21)$$

Podemos observar que como (21) es una transformación lineal, entonces $X(k)$ y $Y_l(k)$ observan un mismo comportamiento aleatorio y por tanto una pdf similar. Así ahora se debe maximizar la función $f_y(\theta|x)$.

El algoritmo de Maximización de la Esperanza consta de tres pasos principales:

- Partir de un valor inicial de los parámetros
- Estimar la esperanza de los valores ausentes, dados los parámetros iniciales y el resto de las observaciones (prever los valores ausentes)
- Obtener un nuevo valor de los parámetros, que maximice la función de verosimilitud. Si se cumple que $|\hat{\theta}^{n+1} - \hat{\theta}^n| < \varepsilon$, se detiene, de lo contrario regresa al segundo paso.

En la Figura 6 se muestra el diagrama de flujo correspondiente a este algoritmo.

Posteriormente en los siguientes apartados se explicarán más a detalle en que consiste cada uno de estos pasos y su importancia.

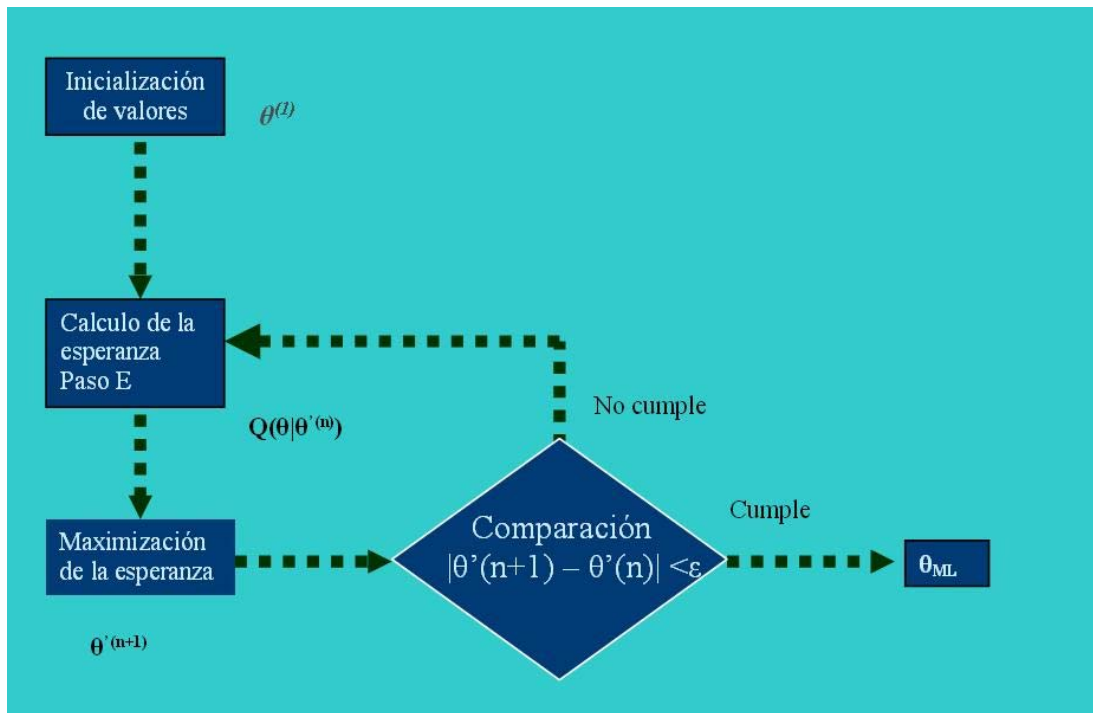


Figura 6. Diagrama de flujo del algoritmo de Maximización de la Esperanza.

2.4.1 Inicialización

Aunque este algoritmo nos llevará a los parámetros estimados, necesita una inicialización adecuada a fin que pueda converger rápidamente hacia éstos.

Un método presentado en [Zizkind y Wax, 1998] y en [Van Trees, 2002] es resolver el problema para una sola fuente de manera que:

$$\hat{\theta}_1(0) = \arg \max_{\theta_1} \left\{ \text{tr} \left[P_{v(\theta_1)} C_x \right] \right\}, \quad (22)$$

donde C_x es la matriz de covarianza y P_v es la matriz de proyección sobre el rango del vector de dirección.

Entonces se resuelve para la segunda fuente suponiendo que la primera fuente es $\hat{\theta}_1(0)$:

$$\hat{\theta}_2(0) = \arg \max_{\theta_2} \left\{ \text{tr} \left[P_{[v(\theta_1), v(\theta_2)]} C_x \right] \right\}. \quad (23)$$

Otra forma de inicialización es utilizar algún método de estimación de DoA, menos complicado y más rápido, i.e. MUSIC, Root MUSIC, ESPRIT, etc. y utilizar los valores arrojados por cualquiera de estos algoritmos [Çirpan y Çekli, 2002].

2.4.2 Cálculo de la Esperanza

Una vez que se han obtenido los valores de inicialización, para cada iteración ahora se calcula la esperanza de la función de máxima verosimilitud condicionada a los datos observados X y al valor del parámetro $\theta^{(n)}$, esta función se define como:

$$Q(\theta | \hat{\theta}^{(n)}) = E[\ln f(X : \theta | Y, \hat{\theta}^{(n)})], \quad (24)$$

donde: $\hat{\theta}^{(n)}$ es la estimación actual de θ para la n -ésima iteración; como la esperanza del logaritmo de la función de verosimilitud correspondiente al valor del parámetro θ y a los datos completos x . La esperanza se calcula respecto de la distribución condicionada a los datos observados y y al valor del parámetro θ .

Se usa el valor esperado, porque la variable que representa los datos completos (Y) es una variable aleatoria.

2.4.3 Maximización de la Esperanza

El siguiente paso es maximizar $Q(\theta|\hat{\theta}^{(n)})$ con respecto a θ para obtener $\hat{\theta}^{(n+1)}$ el cual utilizaremos en la siguiente iteración.

$$\hat{\theta}^{(n+1)} = \arg \max_{\theta} Q(\theta|\hat{\theta}^{(n)}), \quad (25)$$

$\hat{\theta}^{(n+1)}$ será el nuevo valor estimado

Estos dos pasos se repetirán hasta que:

$$|\hat{\theta}^{(n+1)} - \hat{\theta}^{(n)}| < \varepsilon. \quad (26)$$

Al término de lo cual se asegura que el valor final de $\hat{\theta}$ maximiza a nuestra función de verosimilitud.

2.5 Conclusiones

Durante este capítulo se ha desarrollado la teoría necesaria para entender el estimador de Máxima Verosimilitud, se han presentado sus orígenes, características y propiedades. Es importante considerar que la cota de Cramer- Rao se utiliza sólo para verificar la eficiencia de nuestro estimador, aunque puede utilizarse para otros fines, incluso como un medio para la obtención del estimador en sí mismo. Por sus propiedades matemáticas, este

estimador puede ser aplicable a sistemas con señales de banda ancha, con mayor número de usuarios, e incluso bajo ambientes muy agresivos donde se tenga una baja SNR.

Ya que se ha revisado la teoría necesaria podemos ahora realizar el modelado matemático del sistema, el cual se llevará a cabo durante el siguiente capítulo, para posteriormente llegar a las simulaciones y análisis de resultados.

Capítulo III Modelo del sistema

3.1 Introducción

En este capítulo desarrollaremos el modelo matemático para el estimador de Máxima Verosimilitud basado en el algoritmo de Maximización de la Esperanza, explicados conceptualmente en el capítulo anterior. En este capítulo se desarrollarán las expresiones específicas, de acuerdo a la aplicación que se da en esta tesis, que son necesarias para poder aplicar este estimador a los entornos considerados en el proceso de simulación , el cual será abordado en el siguiente capítulo.

Una vez establecida la teoría para el estimador ML, en este capítulo se desarrollarán los modelos de las fuentes de interés, su incidencia e influencia en el agrupamiento de antenas, considerando los escenarios de campo cercano y de campo lejano. Se tratará también el modelo del estimador en estos escenarios, considerando cada una de las etapas de las cuales está compuesto.

3.2 Modelo de señal

Es importante definir, en el escenario que nos proponemos, el modelo matemático del sistema, ya que el estar trabajando en campo cercano introduce nuevos factores a considerar, por un lado el frente de onda ya no se considera plano, dada la proximidad de las fuentes ahora el desfaseamiento se considera de forma esférica, tal y como se muestra en la Figura 7. Con la finalidad de abarcar los ambientes posibles, se presenta el modelo matemático para la señal tanto en campo cercano como en campo lejano.

3.2.1 Campo cercano

El modelo utilizado para construir la *matriz de datos* $\mathbf{X}(\mathbf{t})$, se encuentra formado por las K muestras de las señales provenientes de D fuentes que inciden en el agrupamiento, donde D es el número de fuentes de interés. El agrupamiento utilizado está compuesto por M elementos de antena, distribuidos sobre el eje x , según se indica en la Figura 7. Cada elemento tiene un peso I_m , con una separación entre elementos d en términos de longitud de onda. Por simplicidad, se asignan ganancias I_m unitarias para todos los elementos.

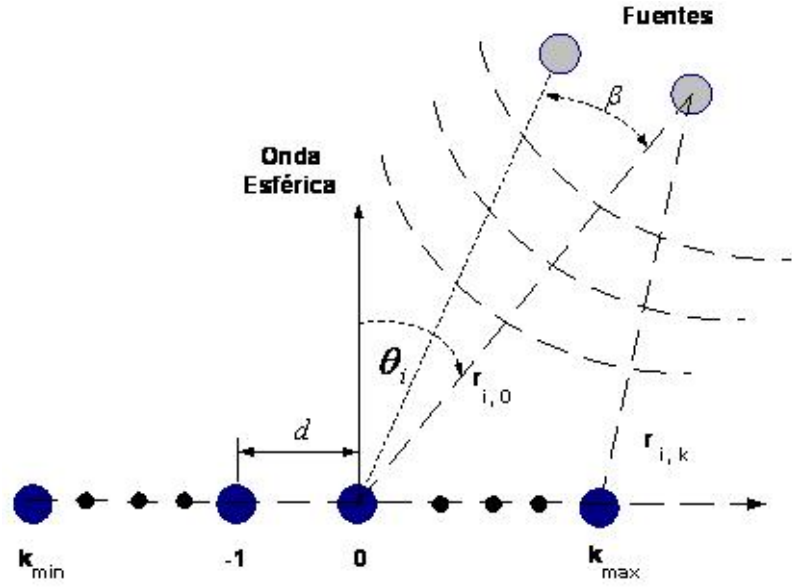


Figura 7. Modelo de datos, distribución de elementos del agrupamiento de antenas, considerando un frente de onda esférico.

La salida de cada elemento del agrupamiento se representa por la ecuación (27), donde el término exponencial conforma el vector de dirección ν del agrupamiento, por la ecuación (28), el cual establece la diferencia de fase entre elementos, q es el índice de fuentes. El cambio de fase se representa por dos componentes, μ_q que representa la componente de fase de campo lejano y ζ_q es el término para aproximar el campo cercano dado el frente de onda esférico [Çekli y Çirpan, 2003].

$$x_m(k) = \sum_{q=1}^D e^{j(\mu_q m_q + \zeta_q m_q^2)} s_q(k) + n(k), \quad (27)$$

$$\nu(\mu_q, \zeta_q) = \left[e^{j(\mu_q m_{\min} + \zeta_q m_{\min}^2)}, \dots, 1, \dots, e^{j(\mu_q m_{\max} + \zeta_q m_{\max}^2)} \right]^T. \quad (28)$$

Donde:

$$\mu_q = -\frac{2\pi d}{\lambda} \sin(\theta_q), \quad \zeta_q = \frac{\pi d^2}{\lambda r_q} \cos^2 \theta_q,$$

El primer componente de fase de (27) corresponde a la fase en campo lejano $\mu_i = -(2\pi d / \lambda) \sin \theta_q$. El término que aproxima el efecto de la fuente en campo cercano, incluido en la fase por medio de un polinomio de segundo orden, es una función no lineal de la distancia r_i y el azimut $\zeta_i = -(2\pi d^2 / \lambda r_q) \cos^2 \theta_q$. La forma matricial se representa en la ecuación (29) donde la matriz V está constituida por los vectores dirección v para cada una de las q fuentes, S es la señal y N el vector de ruido

$$X(k) = V(\mu, \zeta)S(k) + N(k), \quad k = 1, 2, \dots, K \quad (29)$$

Las D fuentes radiantes son omni-direccionales, representadas en magnitud y fase por

$$s_q(t) = A(t)\alpha(t)e^{-j2\pi ft}, \quad (30)$$

donde: $A(t)$ es la amplitud de la señal,

$\alpha(t) = 1/r^\varepsilon$ son las pérdidas por propagación,

ε es el exponente de pérdidas y

f la frecuencia de interés.

En $t=r/c+kT$ se consideran el tiempo de propagación y el muestreo, en el instante k con periodo T , [Liberti y Rappaport, 1999]. La amplitud $A(t)$ es un proceso aleatorio gaussiano con media cero que también depende del tiempo pero con una frecuencia menor a la frecuencia portadora de la señal [Van Trees, 2002].

3.2.2 Campo lejano

En este modelo la matriz de datos $\mathbf{X}(t)$, contiene las N muestras de D fuentes o señales que arriban al agrupamiento de antena. El agrupamiento está compuesto por M elementos, distribuidos sobre el eje x . Cada elemento tiene un peso o ganancia I_m , con una separación entre elementos d . En este caso, también, las ganancias I_m serán unitarias, según se indica en la Figura 8.

El agrupamiento de antena estará formado por un conjunto de elementos omnidireccionales, es decir, dipolos colocados en dirección del eje z . De tal forma que la salida $\mathbf{X}(t)$ del agrupamiento es producida por la incidencia de ondas planas, de D fuentes decorreladas de banda estrecha a una distancia r . Estas inciden en el agrupamiento en dirección $\theta_q = [\theta_1, \dots, \theta_D]$.

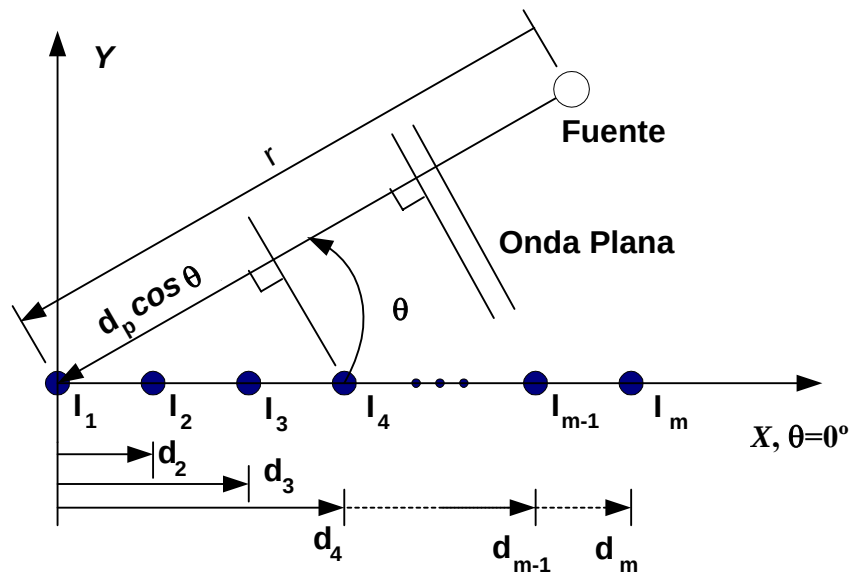


Figura 8. Modelo de datos, distribución de elementos del agrupamiento.

La salida de cada elemento del agrupamiento se representa por (31), donde el término exponencial de cada elemento conforma el vector de dirección $\mathbf{v}$ del agrupamiento (32), el cual establece la diferencia de fase entre elementos, q es el índice de fuentes. Como se puede observar el componente de fase posee, ahora, una forma más simple.

$$x(t) = \sum_{q=1}^D \mathbf{v}(\theta_q) s_q(t) + n(t), \quad (31)$$

$$\mathbf{v}(\theta_q) = [e^{-j2\pi d \cos(\theta_q)} \quad \dots \quad e^{-j2\pi d \cos(\theta_q)}]^\top. \quad (32)$$

La forma matricial se representa en la ecuación (33) donde la matriz V está constituida por los vectores columna de dirección $\mathbf{v}$ para cada una de las D fuentes, S es la señal y N el vector de ruido

$$X(k) = V(\theta)S(k) + N(k), \quad k = 1, 2, \dots, K. \quad (33)$$

Las consideraciones en cuanto a las fuentes radiantes y al comportamiento del vector de ruido y por tanto de las propiedades de los datos en X son las mismas. Fuentes omnidireccionales, representadas en magnitud y fase por (30), con pérdidas por propagación, frecuencia portadora y periodo de muestreo. Podemos ver ahora que la matriz de dirección solo depende de la dirección de arribo de las señales.

3.3 Modelo matemático estimador

Se supone que la señal $s(t)$ y el ruido $n(t)$ están decorrelados para cualquier *tiempo*, y que ambos poseen una función de densidad de probabilidad (pdf) compleja Gaussiana con media igual a cero. Esto nos lleva a que los datos observados X a la salida del agrupamiento, son un vector complejo Gaussiano con media cero y con matriz de covarianza $K_x(\theta)$. La señal de ruido Gaussiano tiene una desviación σ^2 conocida, [Van Trees, 2002; Liberti y Rappaport, 1999].

3.3.1 Función de verosimilitud

Con estas consideraciones ahora podemos empezar a establecer el modelo matemático para el estimador de máxima verosimilitud, explicado en el capítulo anterior. Primero es necesario establecer la función de máxima verosimilitud, considerando que la información traída por una observación x , acerca del parámetro θ , está completamente contenida en la función de densidad de probabilidad $p(x)$. Donde $p(x)$ para un agrupamiento de M elementos de antena es de la forma siguiente, dada la naturaleza establecida para el comportamiento de las señales [Van Trees, 2002; Çekli,y Çirpan,, 2003]:

$$p_{x|\theta}(x) = \frac{1}{\det[\pi K_x(\theta)]} \exp[-(x^H - m_x^H(\theta))K_x^{-1}(\theta)(x - m_x(\theta))]. \quad (34)$$

Dado que la media es igual a cero:

$$p_{x|\theta}(x) = \frac{1}{\det[\pi K_x(\theta)]} \exp[-x^H K_x^{-1}(\theta)x], \quad (35)$$

donde el parámetro θ indica la dirección de arribo de la señal de la fuente de interés.

La pdf conjunta para varias observaciones sucesivas, dado que cada observación x es independiente e idénticamente distribuida (i.i.d.), queda entonces:

$$p_{x_1, x_2, \dots, x_k|\theta}(x) = \prod_{k=1}^K \frac{1}{\det[\pi^N K_x(\theta)]} \exp[-x_k^H K_x^{-1}(\theta)x_k] . \quad (36)$$

Como se puede ver, esta expresión es algo difícil de manejar, por lo cual se utiliza la función logarítmica (log-verosimilitud). La función log-verosimilitud se define como el logaritmo de la pdf conjunta de las observaciones de la siguiente manera:

$$L_x(\theta) = \log p_{x_1, x_2, \dots, x_k|\theta}(x) . \quad (37)$$

Como los términos que no se encuentran dentro de la exponencial no dependen de k , se puede eliminar el multiplicatorio y solo aplicarlo a la exponencial, despreciando los términos innecesarios[Çekli y Çirpan., 2003], con lo que (37) se reduce a la forma:

$$L_x(\theta) = -\left[\log \det[K_x(\theta)] + \frac{1}{K} \sum_{k=1}^K X_k^H K_x^{-1}(\theta) X_k \right], \quad (38)$$

La ecuación (38) se puede simplificar por las propiedades de la matriz de covarianza y utilizando la operación de traza, así se puede reducir a (39):

$$L_x(\theta) = -\left[\log \det[K_x(\theta)] + \text{tr}[K_x^{-1}(\theta) \bullet C_x] \right]. \quad (39)$$

Donde C_x es la *matriz de correlación* del muestreo dada por (40).

$$C_x = \frac{1}{K} \sum_{k=1}^K x(k)x^H(k), \quad (40)$$

$$\text{y } K_x \text{ por: } K_x = E[X(k)X^H(k)] = V(\theta)K_sV^H(\theta) + \sigma_w^2 I. \quad (41)$$

$$\text{Donde: } K_s \text{ está dado por } K_s = E[s(k)s^H(k)] . \quad (42)$$

Dada la función de verosimilitud logarítmica $L_x(\theta)$, ahora hay que encontrar el estimador $\hat{\theta}_{ML}$ el cual maximiza dicha función.

$$\frac{\partial}{\partial \theta_i} [L(\theta)]_{\theta=\hat{\theta}_{ML}} = 0 \quad i = 1, 2, \dots, D \quad (43)$$

O lo que es lo mismo, que cumpla con la condición:

$$\nabla_{\theta} [L(\theta)]_{\theta=\hat{\theta}_{ML}} = 0 \quad (44)$$

Maximizar esta expresión se puede complicar bastante por problemas no lineales, principalmente cuando D crece, por lo cual se deben utilizar técnicas eficientes para resolver el problema. El método utilizado es el algoritmo EM diseñado para problemas de maximización donde se tiene pérdida de datos y alta complejidad matemática.

3.3.2 Algoritmo EM

El algoritmo EM, el cual ya se explicó en la sección 2.4, está diseñado para calcular estimadores de máxima verosimilitud en aquellas situaciones donde el cálculo directo es complicado.

El algoritmo EM utilizado para la maximización de la función de verosimilitud considera dos espacios muestrales: los datos observados, los cuales se encuentran deformados por los efectos del canal (ruido, multitrayectorias, pérdida de datos, etc) y que consideraremos como incompletos. El segundo espacio de datos completos desconocidos, es decir los datos originales, no sufren alteraciones.

- $X_k = \{x_1, \dots, x_k\}$ *datos observados incompletos*
- $Y_k = \{y_1, \dots, y_k\}$ *datos completos desconocidos*

En un escenario con D fuentes de interés, se tienen D problemas de maximización. Se considerará la señal proveniente de cada una de las fuentes $Y_l(t)$, que incide en el arreglo de forma similar a los datos observados (45), suponiendo los datos sin interferencia ni pérdidas, pero desconocidos.

$$Y_l(k) = v(\theta_l) s_l(k) + n_l(k), \quad l = 1, 2, \dots, D \quad (45)$$

Y podemos establecer que la relación entre los datos incompletos $X(k)$ y los datos completos $Y_l(k)$ está dada por la transformación definida por:

$$X(k) = \sum_{l=1}^D Y_l(k), \quad k = 1, \dots, K. \quad (46)$$

Podemos observar que (46) es una transformación lineal, por tanto $X(k)$ y $Y_l(k)$ son conjuntamente Gaussianas y se tiene que la función log-verosimilitud para $Y_l(k)$ es de la misma forma que para $X(k)$, [Van Trees, 2002]:

$$L_y(Y : \theta, K_s) = -[\ln \det K_{yl} + \frac{1}{K} \text{tr}[K_{yl}^{-1} \sum_{k=1}^K Y_l(k) Y_l^H(k)]], \quad (47)$$

donde K_{yl} es la matriz de covarianza de los datos muestreados.

Una vez definida (47), lo que resta es maximizarla con respecto a (θ, Ks) , para esto utilizamos el algoritmo EM el cual consta de los siguientes pasos, [Van Trees, 2002; Miller y Fuhrmann, 2002]:

- Inicialización de los parámetros
- Estimar la esperanza de los valores ausentes (datos incompletos), dados los parámetros iniciales y el resto de las observaciones (estimar los valores ausentes)
- Obtener un nuevo valor de los parámetros, que maximice la función de verosimilitud.
- Si se cumple que $|\theta^{n+1} - \theta^n| < \varepsilon$, se detiene, de lo contrario regresa al segundo paso.

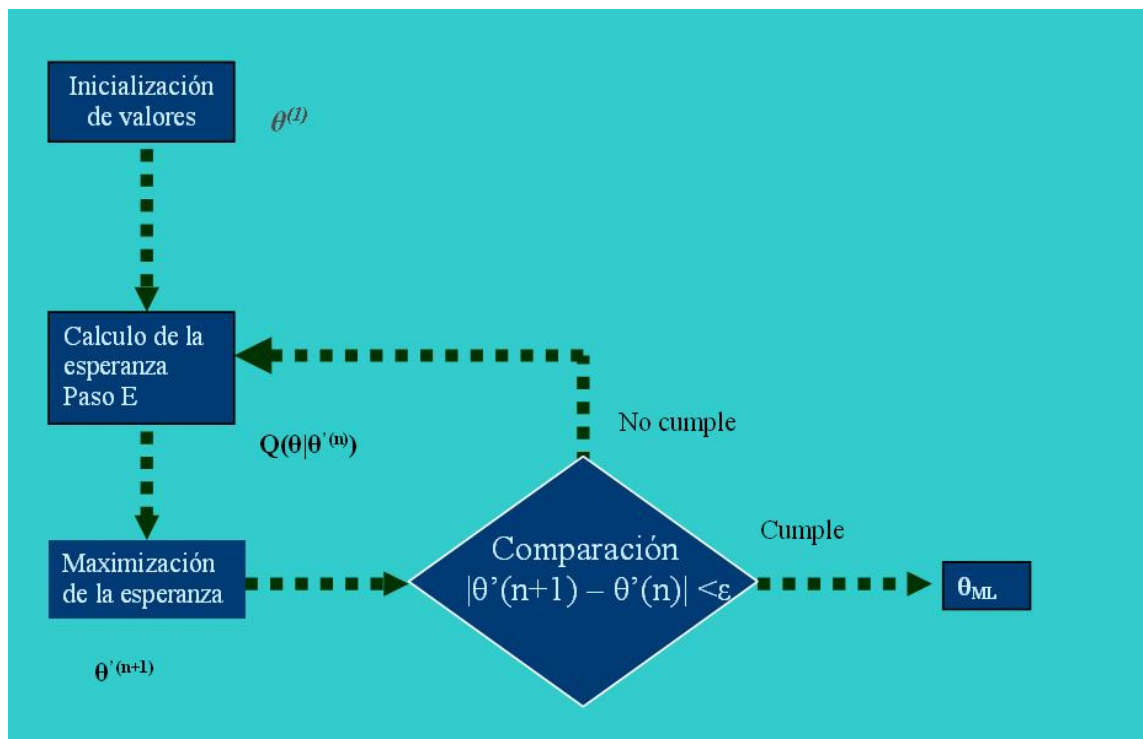


Figura 9 Diagrama de flujo del algoritmo EM

El primer paso, la inicialización, puede darse utilizando una aproximación por algún método como el presentado en [Ziskind y Wax, 1998], o incluso utilizando los valores que arroje algún otro algoritmo de menor definición, como el algoritmo MUSIC.

En el segundo paso, en forma más detallada, para cada iteración se calcula la esperanza condicionada, de la función de máxima verosimilitud, a los datos observados X y al valor del parámetro $\theta^{(n)}$:

$$Q(\theta|\theta^{(n)}) = E[L_Y(Y:\theta|X, \theta^{(n)})]. \quad (48)$$

Donde: $\theta^{(n)}$ es la estimación actual de θ para la n -ésima iteración.

En este caso es suficiente que se calcule la esperanza condicionada de la matriz de covarianza de los datos completos C_Y , ya que es *estadísticamente suficiente* para inferir los datos [Çekli y Çirpan, 2003].

$$C_Y^{n+1} = E[C_Y^n | K_Y^n, K_X^n, C_X]. \quad (49)$$

El siguiente paso es maximizar (48) con respecto a θ , para obtener $\theta'^{(n+1)}$, el cual se utilizará en la siguiente iteración:

$$\theta'^{(n+1)} = \arg \max_{\theta} Q(\theta|\theta'^{(n)}). \quad (50)$$

Así se llega a sustituir la esperanza condicional de C_Y , calculado en el paso anterior, en la función de verosimilitud para los datos completos (38). A continuación se obtiene el

determinante de K_{Y_l} y se desarrolla $K_{Y_l}^{-1}$ para volver a sustituirse en la función de verosimilitud, a fin de simplificar (38), se sustituye K_{Y_l} y $K_{Y_l}^{-1}$ en (38) quedando:

$$\{\theta_q^{p+1}, r_q^{p+1}\} = \arg \max_{\{\mu_q, \zeta_q\}} \frac{v^H(\mu_q, \zeta_q) \hat{C}_{y_q}^{p+1} v(\mu_q, \zeta_q)}{|v(\mu_q, \zeta_q)|^2} \quad (51)$$

Estos dos pasos se repetirán hasta que:

$$|\theta'(n+1) - \theta'(n)| < \varepsilon . \quad (52)$$

Al término de lo cual se asegura que el valor final de θ' maximiza la función de verosimilitud y será el estimador propuesto [Çekli y Çirpan, 2003].

3.3.3 Modelo matemático de la Cota de Cramér Rao

Una vez escogido nuestro estimador, es necesario verificar que es eficiente, esto se hace a través de la Cota de Cramer-Rao. Con el objetivo de medir la bondad de la estimación, se compara con la CCR, la cual representa la cota inferior de la varianza de cualquier estimador insesgado.

Un estimador insesgado se dice que es eficiente desde el punto de vista estadístico si la covarianza de la estima alcanza la CCR, es decir, cuando nuestro estimador alcanza la CCR podemos decir que no habrá un estimador “mejor” a este.

$$\theta' \text{ es un estimador “eficiente” de } \theta \text{ si: } \text{var}(\hat{\theta}) = \frac{1}{J(\theta)} \quad (53)$$

Explícitamente, la cota CCR está determinada por el inverso del Hessiano de la función de verosimilitud:

$$Var(\hat{\theta}) = CCR(\hat{\theta}) \geq \mathbf{J}^{-1}(\theta) = \frac{1}{E\left[\left[\frac{\partial}{\partial \theta} \log L(X; \theta)\right]^2\right)}. \quad (54)$$

O también como el inverso de la matriz de Fisher ($\mathbf{J}$), (53).

Esta cota inferior aplicada al problema de interés nos da una matriz $\mathbf{J}$ donde los elementos se dan de la siguiente manera:

$$CCR^{-1}(\tau) = \frac{2N}{\sigma^2} \text{Re}\left\{\mathbf{A}^H \mathbf{\Pi}^c \mathbf{A} \bullet (\mathbf{1} \otimes \mathbf{K}_s \mathbf{A}^H \mathbf{K}_x \mathbf{A} \mathbf{K}_s)^T\right\} \quad (55)$$

$$\text{Donde:} \quad \mathbf{\Pi} = \mathbf{A}[\mathbf{A}^H \mathbf{A}]^{-1} \mathbf{A}^H \quad (56)$$

$$\text{y } \mathbf{\Pi}^c = \mathbf{I} - \mathbf{\Pi} \quad (57)$$

Al simular esta función de acuerdo a los estimados obtenidos con las simulaciones presentadas en el siguiente capítulo, podremos asegurar la eficiencia de nuestro estimador.

3.4 Conclusiones

Las consideraciones tomadas dentro de la generación del modelo, permiten una mejor comprensión del modelo y del sistema en general. Para empezar se consideraron ambientes tanto en campo cercano como en campo lejano, así podemos soportar un mayor número de escenarios ya sea micro o macro células o los llamados “Hot Spots”. Para la inicialización

del algoritmo EM se pueden utilizar dos métodos, uno puede ser la utilización de algún otro método de estimación de DOA, que sea más rápido aunque impreciso, o utilizar el método utilizado en [Zizkind y Wax, 1999]. Es necesaria una buena inicialización a fin de asegurar una convergencia óptima hacia nuestro estimador, y asegurar su rapidez. La cota de Cramer- Rao se utiliza solo para verificar la eficiencia de nuestro estimador, aunque puede utilizarse para otros fines, incluso como un medio para la obtención del estimador en si mismo.

Una vez que se ha establecido el modelo matemático requerido para este trabajo, ahora se presentarán las simulaciones realizadas. En el siguiente capítulo se presentaran las simulaciones que se consideran pertinentes para evaluar este estimador, considerando los escenarios más representativos y que sustentaran de forma más firme los resultados del presente trabajo.

Capitulo IV Simulaciones y resultados

4.1 Introducción

Este capítulo tiene como principal objetivo presentar simulaciones para la localización de fuentes de interés, basadas en el método de estimación de máxima verosimilitud y en el algoritmo Maximización de la Esperanza, hechas a partir del modelo matemático establecido durante el capítulo anterior. A través de las estadísticas presentadas se puede establecer una comparación con otros métodos en cuanto a su comportamiento bajo condiciones poco estudiadas, (evaluando condiciones de separabilidad de fuentes, campo cercano, ambientes de ruido, número de muestras, número de fuentes).

Los valores que se reportan están basados, en cada simulación y estadística, al error de estimación, es decir el error entre el valor estimado y la posición real de la fuente de interés. Las estadísticas finales constan de valores promedio (simulaciones Montecarlo), es decir se ha repetido el experimento para cada uno de los casos un número determinado de veces y se han obtenido valores promedio del error que arroja cada caso.

4.2 Consideraciones de simulación

Para realizar las estadísticas presentadas a lo largo de este capítulo se establecieron consideraciones y condiciones que fueron necesarias observar, tanto para establecer entornos que fueran congruentes con los objetivos de esta tesis y que se aproximasen, lo más posible, a las condiciones reales de un entorno real. Estas consideraciones de simulación también se agregan a las consideraciones del modelado, las cuales también debieron ser tomadas en cuenta.

Estas consideraciones se listan a continuación:

- Se utilizó un agrupamiento de antenas constituido por 7 elementos, separados $1/4 \lambda$ entre sí.
- Las amplitudes de la señal se simularon como procesos aleatorios complejos con distribución Gaussiana y media cero.
- Se consideró a la matriz de correlación de la señal, $\mathbf{K}_s$, como conocida.
- Se consideró la desviación del ruido gaussiano, σ^2 , como conocida.
- Se realizaron simulaciones estilo Montecarlo con repetición del experimento, $k=50$ repeticiones.

En la siguiente tabla se enumeran las consideraciones correspondientes al canal radio y a las características de cómo fueron simuladas las fuentes de interés.

Tabla I. Consideraciones numéricas de simulación

Parámetros de las fuentes de interés	
Separación entre base y fuente de interés (campo cercano)	1 a 16λ
Posición de fuentes	$-\pi/3$ a $\pi/3$
Separación angular	1° a 25°
SNR	-5 a 20
Número de fuentes	2 a 7
Parámetros de canal	
Frecuencia de portadora	$f=1900$ MHz
No. De muestras	30, 50, 100, 200, 300, 500, 1000, 1800
Periodo de muestreo	$1.365 \text{ e-}010 \text{ seg} \Rightarrow 4f$
Ruido	Aleatorio Gaussiano con media cero

Una vez enlistadas las consideraciones, ahora plantearemos los aspectos y entornos que se evaluaron durante las simulaciones, y que estadísticas se obtuvieron a partir de estas.

4.3 Estadísticas obtenidas

Las estadísticas obtenidas y durante este capítulo son reflejo, primeramente, de las consideraciones presentadas anteriormente pero también de los aspectos que se quisieron analizar durante las simulaciones. Para empezar debemos comentar que el origen de varios de estos entornos obedece a estadísticas y resultados consultados en otras publicaciones [Çekli, y Çirpan, 2003; Feder y Weinstein, 1988; Miller y Fuhrmann, 1990]. Tal es el caso del análisis para distintos niveles de ruido, pero también, son fruto de inquietudes y dudas

que surgieron al analizar estas fuentes de información y a las limitaciones que se han observado en los anteriores métodos de localización de fuentes. Además a través de dichas limitaciones se puede comprobar la robustez del método de estimación ML para la localización de fuentes calculando el RMSE de la estimación.

Los entornos que se decidieron simular se enfocan principalmente a condiciones poco estudiadas. Específicamente bajo la consideración de campo cercano y en la variación de los siguientes parámetros:

- Nivel de SNR.
- Número de muestras.
- Número de fuentes.
- Separabilidad de fuentes:

Todas las simulaciones tienen como variable común tanto el error RMS de estimación y también, al nivel de ruido (SNR). Se disminuyó esta relación aumentando el nivel de ruido a fin de evaluar el comportamiento del estimador bajo ambientes adversos de ruido. A partir de esta variable (SNR) se realizaron las estadísticas para las tres variables restantes.

Las simulaciones más importantes realizadas y que se reportan a continuación son:

- ❑ **Comparación con MUSIC:** para evaluar las prestaciones con respecto a los métodos de subespacio, se plantean estas simulaciones las cuales servirán como referencia del comportamiento de ambos estimadores.

- ❑ **Convergencia del algoritmo:** con esta simulación se puede evaluar el tiempo que le lleva al estimador ML para obtener los resultados, ésta como se verá, es una de las principales limitaciones de este estimador

- ❑ **Evaluación de número de muestras tomadas para la estimación:** con esta simulación se evalúa la robustez del estimador ML, ya que es una de las principales limitaciones de los métodos de subespacio [Liberti y Rappaport, 1999], y con esto se muestra la dependencia del estimador ML a esta variable.

- ❑ **Evaluación de número de fuentes de interés a detectar:** Esto está limitado al número de elementos del agrupamiento de antenas para los métodos de subespacio, es decir, MUSIC sólo puede detectar $n-1$ fuentes, donde n es el número de elementos del agrupamiento. Para el estimador ML no existe una limitación tan directa a este número. Esta simulación no se ha encontrado reportada en otros trabajos, por lo que podría ser considerada una aportación. Salvo esta simulación, todas las demás se llevaron a cabo considerando solo 2 fuentes de interés

- ❑ **Evaluación de separabilidad de fuentes:** con esta simulación podemos ver la definición y el aislamiento de fuentes que puede encontrar el estimador cuando las fuentes de interés se encuentran muy cercanas. Al igual que la simulación anterior, no se han encontrado estadísticas semejantes reportadas en otros trabajos, por lo que también se pueden catalogar como una aportación de este trabajo.

Ya que hemos establecido las estadísticas que reportaremos, pasemos ahora al análisis y discusión de las mismas.

4.3.1 Simulaciones comparativas con MUSIC

Estas simulaciones comparativas con MUSIC se hacen necesarias, como una referencia y complemento para las siguientes estadísticas. A través de éstas podemos evaluar el comportamiento de ambos métodos de estimación bajo las mismas condiciones de simulación.

La Figura 10 es una comparación gráfica de los resultados obtenidos por el algoritmo MUSIC y el estimador ML. En este caso sólo se realizó un experimento para la localización de dos fuentes en campo cercano, dada la característica de MUSIC de poseer un espectro se puede representar de esta manera. El eje horizontal está dado en grados, emulando la ventana que observa el agrupamiento de antenas de -60 a 60 grados, y el eje vertical es la potencia del espectro de MUSIC en dB. En la Figura 10 podemos observar las líneas verticales continuas, las cuales indican la posición de los valores estimados que arroja el estimador ML, que al ser valores puntuales, se agregan a fin de tener una mejor perspectiva de la comparación. Las líneas verticales punteadas indican la posición real de las fuentes de interés. Las fuentes se localizaron en -5 y 30 grados de posición angular, con una SNR de 10 dB y para 1000 muestras tomadas. En la tabla II se encuentran los datos de la simulación, se muestran los valores estimados por cada método y el error correspondiente.

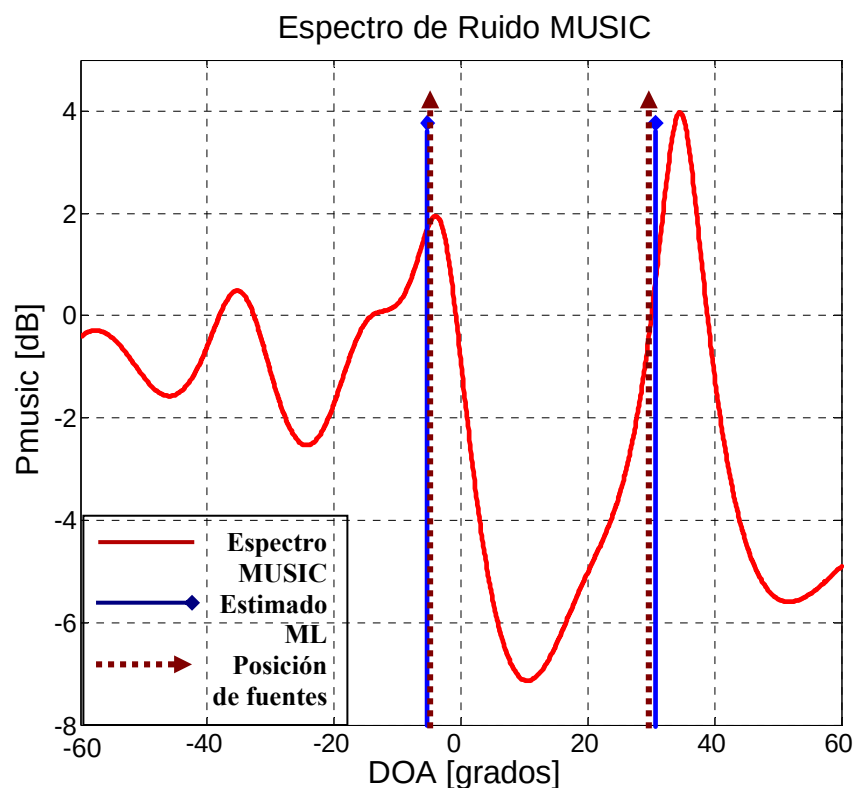


Figura 10. Comparación de estimación, Algoritmo MUSIC y Estimador ML

Tabla II. Comparación entre algoritmo MUSIC y estimador ML.

2 Fuentes MUSIC		
Para una SNR de 10dB	Fuente I	Fuente II
Posición fuentes (°)	-5	30
θ Estimado (°)	-3.9611	34.5657
Error (°)	1.0389	4.5657
2 Fuentes ML		
Para una SNR de 10dB	Fuente I	Fuente II
Posición fuentes (°)	-5	30
θ Estimado (°)	-5.1823	30.339
Error (°)	0.1823	0.339

Los números en negritas corresponden al error de estimación de cada uno de los métodos para cada una de las dos fuentes que se consideraron, podemos ver claramente que el menor error corresponde a los resultados arrojados por el estimador ML, obviamente

para un solo caso en donde se ha corrido una sola vez pues no se puede considerar como un indicativo de cómo se comportaría el estimador bajo otras condiciones o en otros casos, donde, debido a la aleatoriedad del ruido, cambie el error de estimación. Por esto se decidió recurrir a las simulaciones tipo Montecarlo.

Otra estadística importante, y que ha sido reportada en otros escritos, es un análisis para diferentes niveles de ruido. En la Figura 11 observemos una simulación donde ahora el eje horizontal representa la variación del SNR, y el eje vertical el error de estimación en grados. Este tipo de análisis es útil para evaluar como se comportan ambos estimadores bajo condiciones adversas y favorables. Además también podemos ver el comportamiento del estimador ML con respecto a la cota de Cramér-Rao y que verifica la eficiencia de nuestro estimador al ser aproximadamente igual pero menor a los valores estimados. Esto cumple con la condición explicada en las secciones 2.3.3 y 3.3.3.

En la Figura 11 podemos ver como para un mayor SNR el error de estimación disminuye y al disminuir el SNR, el error de estimación aumenta considerablemente para el algoritmo MUSIC, mientras que el estimador ML no se ve afectado en tanta magnitud. Esta simulación se hizo con 2 fuentes de interés, un número de muestras de 1000 y con una separación angular de fuentes de 6 grados. Se escogió esta separación ya que en este punto es donde el algoritmo MUSIC comienza a dar errores demasiado grandes si continúan aproximándose las fuentes o se disminuye la SNR. Notemos que para niveles de ruido óptimo, el error arrojado por ambos métodos tiende a parecerse.

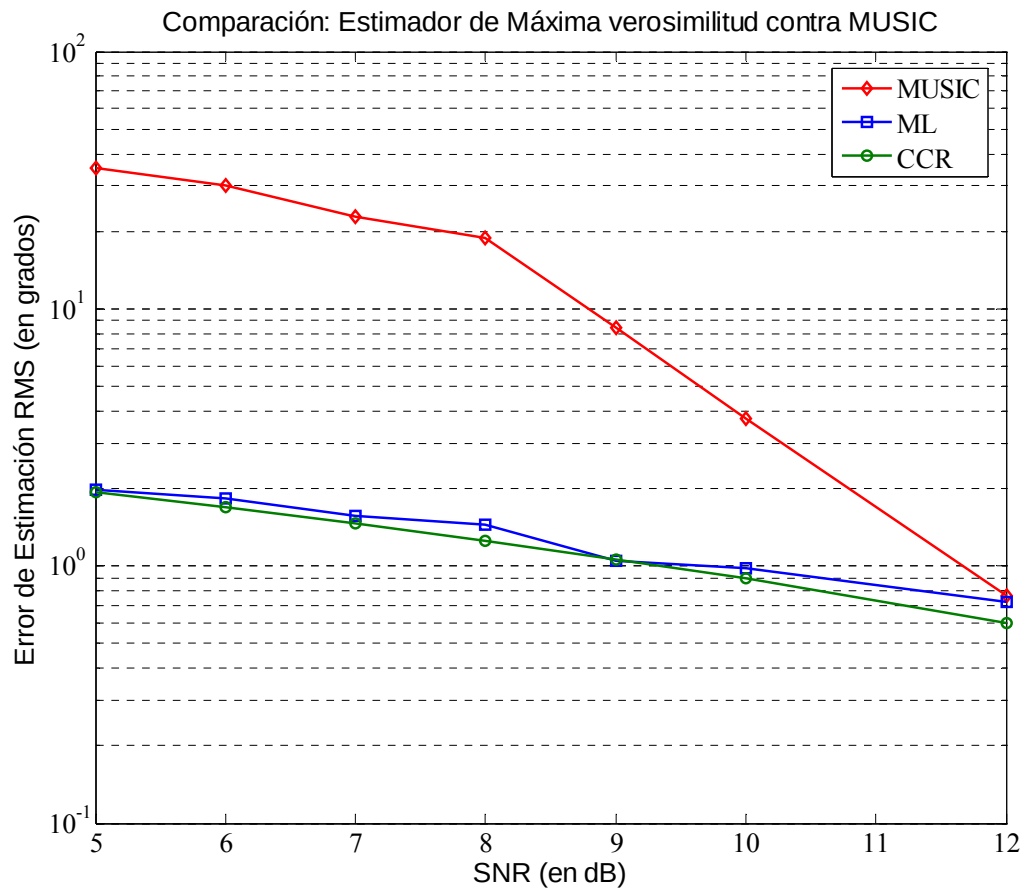


Figura 11. Comparación de comportamiento Estimador ML vs MUSIC.

Tabla III. Error de estimación en grados para la comparación de ML contra MUSIC

SNR (en dB)	Máxima Verosimilitud	MUSIC
5	1.9783	35.148
6	1.8388	30.104
7	1.5569	22.802
8	1.4523	18.751
9	1.051	8.4388
10	0.9812	3.7592
12	0.7256	0.76437

En estas dos figuras se hacen patentes las limitaciones de los algoritmos basados en subespacios, en particular MUSIC, podemos observar como para niveles de ruido relativamente alto este algoritmo deja de arrojar resultados útiles para un sistema real.

Durante los siguientes apartados se mostrarán las estadísticas que se obtuvieron para el estimador ML, evaluando las variables que se mencionaron en la sección 4.3.

4.3.2 Simulaciones del estimador de máxima verosimilitud

Estas simulaciones, como se mencionó anteriormente, tienen como objetivo mostrar la robustez del estimador de ML bajo condiciones y escenarios poco estudiados y reportados. También se evalúa el desempeño de este estimador con respecto al algoritmo MUSIC bajo las mismas condiciones, que si bien no es el objetivo de esta tesis, sirve como referencia de comparación.

4.3.2.1 Convergencia del algoritmo

La primera estadística que se consideró importante es la convergencia del algoritmo, ya que como se ha mencionado, el costo computacional y por tanto el tiempo de procesamiento son la principal desventaja de este estimador. De la Figura 12 podemos observar el número de iteraciones necesarias que necesita el algoritmo EM para converger a un valor, en este caso el eje vertical representa el número de iteraciones realizadas por el algoritmo para llegar a un valor estimado, el eje horizontal indica el nivel de SNR y la familia de curvas varía para diferentes número de fuentes.

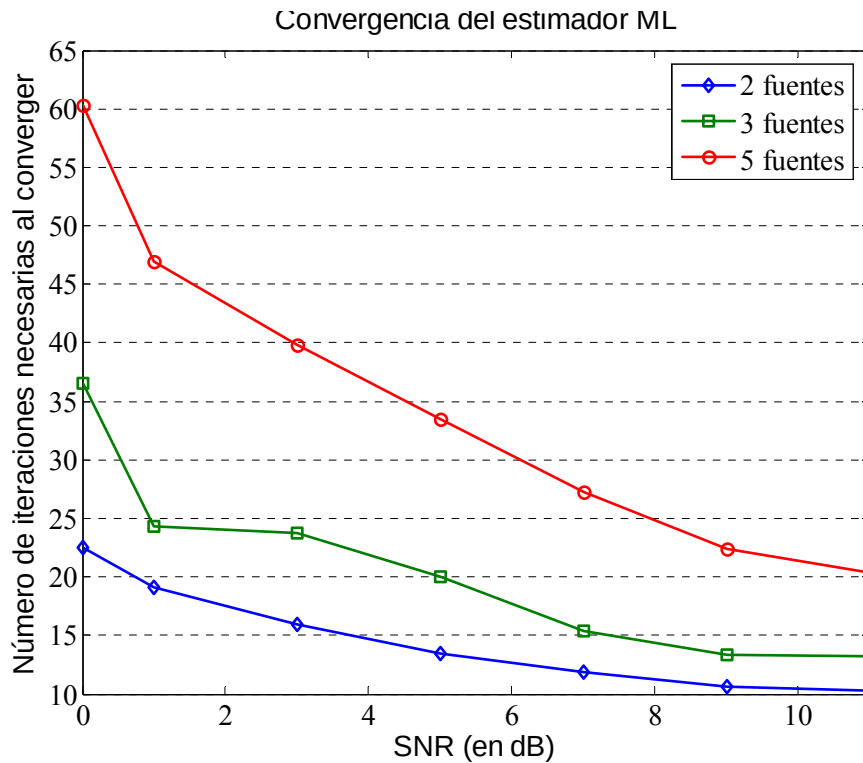


Figura 12 Convergencia del estimador de máxima verosimilitud bajo ambientes adversos de SNR y distinto número de fuentes

Podemos observar, en la Figura 12, como, para poder llegar a un valor estimado se hacen necesarias mas iteraciones conforme el nivel de SNR disminuye. También al momento de aumentar el número de fuentes a localizar el algoritmo incrementa el número de iteraciones necesarias; esto es lógico, ya que a mayor ruido hay mayor pérdida de datos y con más fuentes existe mayor interferencia, y se esperaba una mayor dificultad en la inferencia de los datos.

Debido a estos resultados, al aumentar el número de fuentes de interés a detectar el tiempo necesitado por el algoritmo EM, para generar los valores estimados, aumentaría. Por esto las demás simulaciones, exceptuando la evaluación del número de fuentes a detectar,

se realizaron considerando solo 2 fuentes de interés. Ya que con dos fuentes se pueden observar los efectos de pérdida de datos y de interferencia entre fuentes.

4.3.2.2 Evaluación de número de muestras

Se realizaron simulaciones para distintos números de muestras tomadas para realizar la estimación de la posición de fuentes de interés, a fin de evaluar la dependencia del algoritmo a esta variable. En particular, esta es una limitación para los métodos estudiados anteriormente (MUSIC, Root MUSIC), ya que el funcionamiento depende mucho del muestreo y del número de muestras tomadas. Como ya se mencionó, la frecuencia de muestreo es 4 veces la frecuencia de la portadora y la separación entre fuentes y estación base se estableció en 10 y 12 longitudes de onda para las fuentes 1 y 2 respectivamente.

Fijando la separación angular entre fuentes, se varió el número de muestras del experimento en 30, 50, 100, 200, 300, 500, 1000 y 1800 muestras para hacer la estimación del DoA de dos fuentes. Se varía la SNR de 0 dB a 12 dB, considerando la posición de las fuentes aleatorias con distribución uniforme. Los resultados se muestran en la Figura 13 y numéricamente en la Tabla IV. Evaluación de Número de Muestras, e indica que al ir variando el número de muestras se pueden obtener estimaciones del DOA con errores menores conforme aumenta el número de muestras.

En condiciones desfavorables (30 muestras) el algoritmo logra un error RMSE $< 3^\circ$ para los casos en los que el número de muestras es 30 y SNR igual a 0 dB. Un error RMSE

$< 2^\circ$ para los casos en los que la SNR es mayor a 0 dB con un número de muestras igual a 100 o para una mejor SNR (2 dB). Un error RMSE menor a 1° para los casos en los que la SNR es mayor a 8 dB y con 30 muestras o más, o para SNR mayores a 4 dB y 1000 muestras o más.

En este caso podemos observar que la estimación arroja mayores errores conforme se reduce el número de muestras y la SNR, esto se esperaba ya que estas variables indican menor información recolectada y mayor pérdida de datos. Sin embargo, comparado con las prestaciones de MUSIC, este algoritmo no se ve tan afectado por estas variables, ya que aún se obtienen datos válidos que se pueden utilizar para dirigir el haz de respuesta, mientras que para MUSIC en las condiciones mas desfavorables el error es considerable.

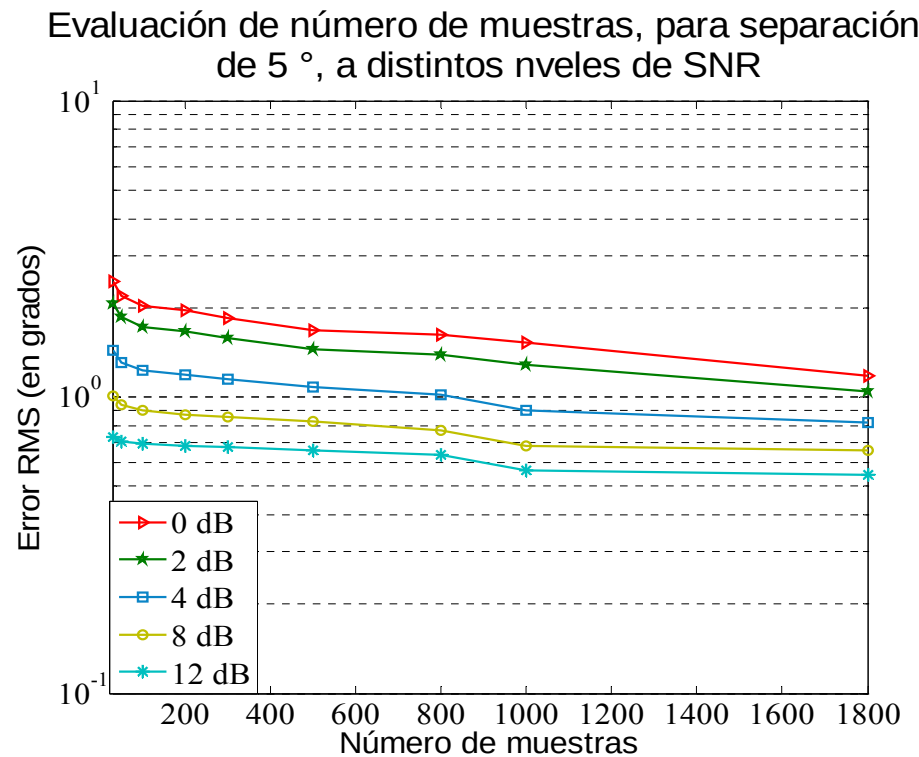


Figura 13. Evaluación de número de muestras tomadas para la estimación

Tabla IV. Evaluación de Número de Muestras

SNR/ Muestras	30	50	100	200	300	500	800	1000	1800
0	2.446	2.1985	2.0338	1.9626	1.8533	1.6791	1.6219	1.5236	1.1833
2	2.058	1.8587	1.727	1.6654	1.5843	1.4506	1.3901	1.2799	1.0456
4	1.720	1.5628	1.46	1.4074	1.3503	1.2511	1.1893	1.0717	0.9247
6	1.432	1.3106	1.2331	1.1885	1.1513	1.0806	1.0197	0.8991	0.8205
8	1.193	1.1022	1.0461	1.0089	0.9871	0.9391	0.8811	0.7619	0.7331
10	1.004	0.9375	0.8990	0.8684	0.8580	0.8266	0.7736	0.6850	0.6623
12	0.864	0.8167	0.7919	0.7671	0.7638	0.7430	0.7029	0.6603	0.6083
14	0.8028	0.7712	0.7634	0.7539	0.7365	0.6990	0.6973	0.6091	0.5711
16	0.774	0.7397	0.7249	0.7050	0.7045	0.6885	0.6548	0.5942	0.5595
18	0.7439	0.7169	0.7106	0.6984	0.6909	0.6665	0.6520	0.5686	0.5505
20	0.7346	0.7064	0.6977	0.6822	0.6802	0.6630	0.6379	0.5637	0.5467

4.3.2.3 Evaluación del número de fuentes de interés a detectar

También se realizó el experimento para un diferente número de fuentes de interés: 3, 5 y 7 fuentes. Esta variable, es una limitante muy importante para los métodos de subespacio, ya que solo pueden detectar un número de fuentes máximo igual al número de elementos del agrupamiento menos 1. Para el caso de este experimento solo 6 fuentes serían detectadas por el algoritmo MUSIC, pero en el caso del estimador ML se pueden detectar más. El aumentar el número de fuentes conlleva también a aumentar el tiempo de procesamiento necesario para llegar a los resultados del estimador ML, lo que hace patente su mayor desventaja.

Se tomaron **1000 muestras** para hacer la estimación del DoA de 3, 5 y 7 fuentes, fijando, también, la separación angular de las fuentes a 4° y modificando la relación SNR de -5 a 11 dB. Esto para comprobar las prestaciones y bondades del estimador, bajo ambientes todavía más adversos donde el algoritmo MUSIC simplemente no funcionaría. Los resultados obtenidos en la Figura 14 y numéricamente en la Tabla V, indican como el error de estimación disminuye conforme disminuye el nivel de ruido, y también al disminuir el número de fuentes.

Pero lo más importante, sin duda es el hecho de que, además de presentar errores por debajo de los 6° en comparación con MUSIC, se puede realizar la estimación de más de 6 fuentes. Lo que indica que el algoritmo no está sujeto por la misma limitación que los métodos basados en subespacios en cuanto al número de fuentes se refiere.

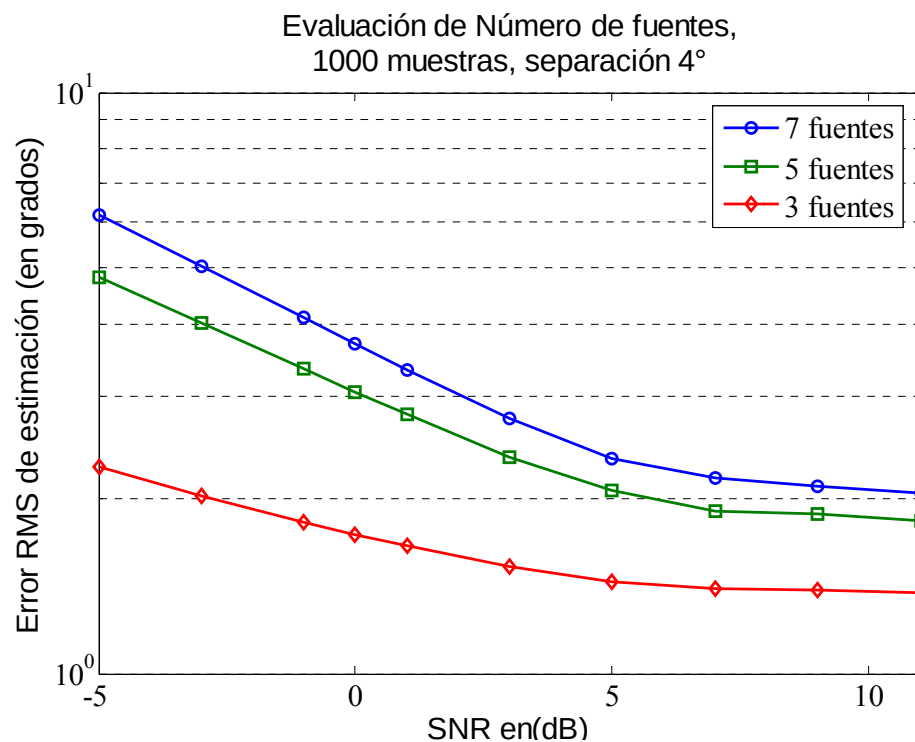


Figura 14. Evaluación del número de fuentes de interés estimadas

Tabla V. Evaluación del número de fuentes de interés estimadas

SNR/No. fuentes	7	5	3
-5	6.3498	5.0162	2.6678
-3	5.2362	4.2166	2.4263
-1	4.2986	3.5415	2.223
0	3.8958	3.2505	2.1356
1	3.5369	2.9907	2.0577
3	2.9511	2.5643	1.9305
5	2.5412	2.2623	1.8414
7	2.367	2.1026	1.8026
9	2.3072	2.0847	1.7904
11	2.2491	2.0315	1.7774

Es importante mencionar que este tipo de análisis no se había realizado anteriormente dentro del Grupo de Comunicaciones Inalámbricas, ni se ha encontrado reportado en publicaciones o libros, y que se consideró como un aspecto importante a resaltar y evaluar dentro del análisis del estimador ML

4.3.2.4 Evaluación de separabilidad de fuentes

Ahora pasemos a la que se ha considerado las más importante de todas las estadísticas generadas, ya que con esta se cierra la evaluación y se logra mostrar, en términos de la separabilidad de fuentes, la robustez del estimador ML, bajo escenarios adversos y que han sido poco estudiados.

Se generó una familia de curvas para diferentes separaciones angulares entre dos fuentes, con posiciones angulares aleatorias, las cuales van de 1° a 30° . Fijando la cantidad de muestras, se tomaron **1000 muestras** para hacer la estimación del DoA. Se varía la relación SNR de 0 a 20 dB.

Los resultados obtenidos en la Figura 15 y numéricamente en la Tabla VI indican que: con 1000 muestras se pueden obtener estimaciones de la dirección de arribo (DOA) con un error RMS menor a 2° para los casos en los que la SNR sea mayor a 10 dB y con separación de fuentes de 1° , o cuando la separación de fuentes es mayor a 4° para una SNR igual a 0 dB. También se puede observar un error RMS menor a 1° para los casos en los que la SNR es mayor a 0 dB y con separaciones angulares mayores a 12° . La separación entre fuentes y estación base estuvo en 10 y 12 longitudes de onda para las fuentes 1 y 2 respectivamente.

Se puede observar, como era de esperarse, que para mejor SNR el error de estimación disminuye bastante, igual que lo hace para mayores separaciones angulares entre fuentes de interés. Es importante analizar este aspecto del estimador ML, ya que dentro de los métodos de subespacio, la mejor definición para encontrar fuentes se ha encontrado alrededor de los 10 grados para niveles de SNR favorables, por encima de los 8 dB, que es el límite para poder distinguir entre ambas fuentes o sólo identificar una. A diferencia de esto, el estimador ML logra distinguir las fuentes aún cuando se trata de separaciones muy pequeñas, como en el caso de 1°

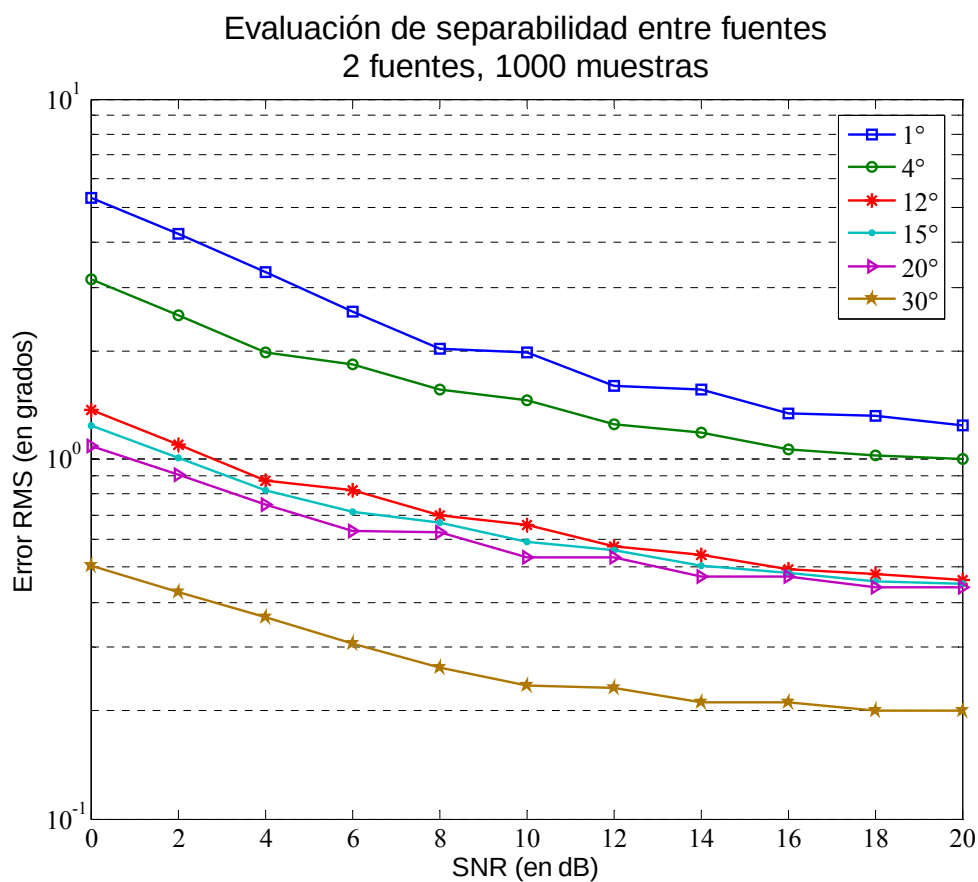


Figura 15. Análisis de separación de fuentes

Tabla VI. Análisis de separación de fuentes

SNR/ Separaciones	1°	4°	12°	15°	20°	30°
0	5.3215	3.1672	1.3656	1.2409	1.0898	0.50508
2	4.232	2.515	1.0951	1.0107	0.90371	0.42838
4	3.3115	1.9783	0.87258	0.81993	0.74875	0.36271
6	2.5598	1.8388	0.81752	0.715	0.62957	0.30807
8	2.0281	1.5569	0.69797	0.66866	0.62491	0.26446
10	1.9771	1.4523	0.65668	0.58997	0.53553	0.23445
12	1.5973	1.251	0.57128	0.55687	0.5322	0.23188
14	1.5633	1.1812	0.54376	0.50441	0.47261	0.21187
16	1.3354	1.0605	0.49252	0.48455	0.47062	0.21033
18	1.3184	1.0256	0.47876	0.45832	0.44082	0.20033
20	1.2425	0.9962	0.46168	0.4517	0.44016	0.19981

Para poder establecer el comportamiento a menores niveles de SNR se presenta la Figura 15, para la cual se consideran las mismas condiciones de la simulación reportada en la Figura 16.

Ahora la relación SNR es modificada de -5 a 11 dB. Los resultados obtenidos en la Figura 16, y numéricamente en la Tabla VI, nos muestran el error obtenido en las estimaciones de la *dirección de arribo* (DOA) para escenarios desfavorables donde la SNR es muy baja. El algoritmo logra un error RMS de estimación menor a 1° para los casos en los que la SNR es mayor o igual a -5 dB y para separaciones mayores o iguales a 20° . Un error RMSE menor a 2° para los casos en los que la SNR es mayor a -5 dB y para separaciones de 9° , o para los caso en los que la SNR es mayor a 1 dB y con separación de fuentes de 5° . Al igual que en la Figura anterior podemos ver que mientras disminuyen la SNR y la separación angular entre fuentes el error de estimación aumenta, pero siempre arrojando resultados menores a los 10 grados, en situaciones donde los métodos basados en subespacio simplemente no funcionan.

La separación entre fuentes y estación base estuvo en 10 y 12 longitudes de onda para las fuentes 1 y 2 respectivamente.

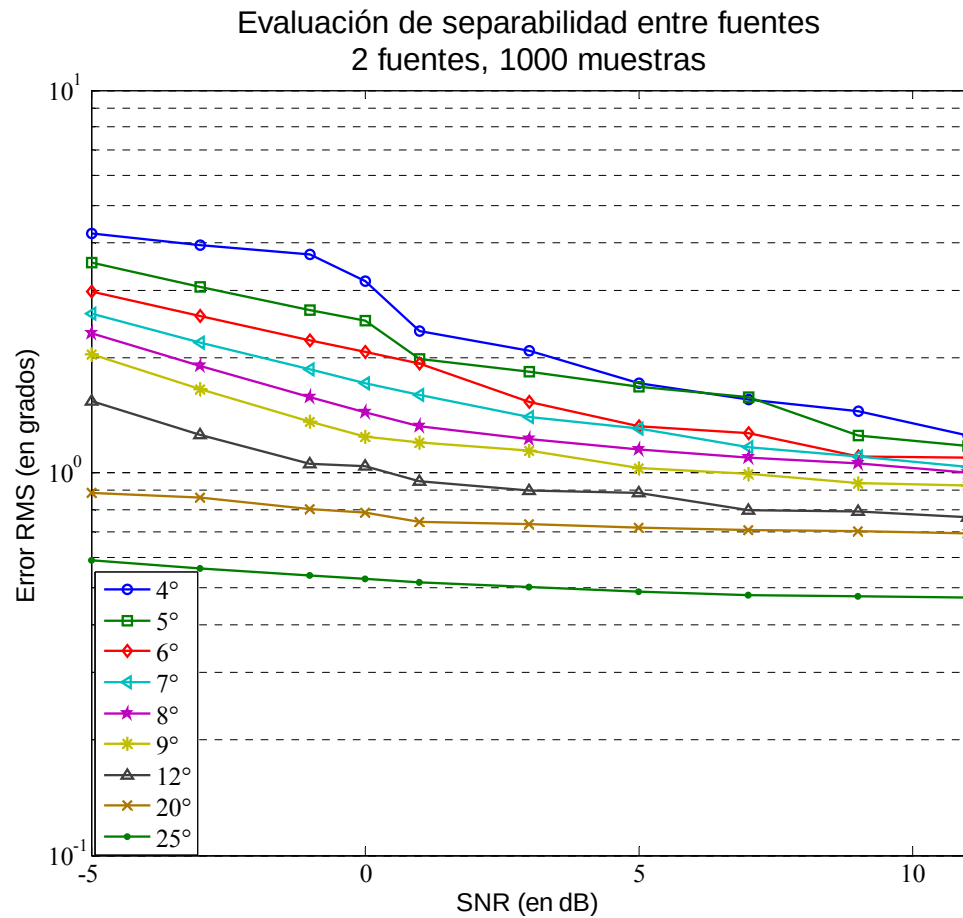


Figura 16. Evaluación de separabilidad entre fuentes

También es necesario remarcar que este análisis tampoco se ha realizado anteriormente, tanto dentro del Grupo de Comunicaciones Inalámbricas ni se ha encontrado reportado en publicaciones o libros. También se consideró como un aspecto importante a resaltar y evaluar dentro del análisis del estimador ML, como muestra concluyente de la robustez del estimador de Máxima Verosimilitud.

A partir de las simulaciones anteriores, con los datos obtenidos y del comportamiento observado, podemos establecer las principales diferencias entre los estimadores de Máxima Verosimilitud y el algoritmo MUSIC, las cuales se listan a continuación:

Máxima Verosimilitud

- Mayor definición bajo ambientes adversos:
- Separaciones menores a 5° y niveles SNR menores a 0 dB
- No tan limitado al nivel de SNR
- No dependiente del número de elementos del agrupamiento de antenas
- No limitado en número de fuentes
- Funciona con pocas muestras: 30
- Robusto bajo condiciones de bajos niveles de SNR (-5 dB) y separaciones entre fuentes pequeñas (1°)

MUSIC

- Rápido y preciso en ambientes controlados:
- Separaciones mayores a 8° y niveles SNR mayores a 10 dB
- Dependiente del nivel de SNR ≥ 10 dB
- Dependiente del número de elementos del agrupamiento de antenas.
- Limitado en número de fuente detectables ($n-1$)
- Dependiente del muestreo (frecuencia, muestras)

También, de las estadísticas reportadas, podemos establecer los límites para los cuales el algoritmo MUSIC todavía funciona adecuadamente y a partir de estos comparar y evaluar las bondades del estimador de Máxima Verosimilitud.

4.4 Resultados

Bajo las consideraciones de simulación establecidas encontramos que el estimador ML y el algoritmo MUSIC varían su comportamiento cuando el nivel de SNR, el número de muestras o de fuentes es cambiado, estos factores influyen en la estimación de la posición de las fuentes de interés, pero cada uno de los estimadores se ve afectado en diferente magnitud.

Para el primer caso, donde las simulaciones se hacen variando solamente el nivel de ruido (SNR), se puede observar en la Figura 11 como el estimador de Máxima verosimilitud mantiene un error de estimación relativamente bajo en comparación con el algoritmo MUSIC, el cual a partir de 9 dB, en este caso, arroja errores demasiado altos. Esto ya nos da información inicial, sobre los límites de aplicación del algoritmo MUSIC, límites que para el estimador ML se encuentran todavía bajo condiciones más adversas.

Cuando se reduce el número de muestras tomadas para hacer la estimación, la dependencia del estimador ML a este parámetro no es muy importante, como se muestra en la Figura 13, incluso con 30 muestras, el estimador ML arroja resultados con un error de estimación pequeño. Cosa muy contraria al algoritmo MUSIC, el cual es muy dependiente a este parámetro, por lo que se debe utilizar un número alto, alrededor de 1000 o más, esto aumentaría el tiempo necesario para realizar la estimación.

Para el caso de variar el número de fuentes a estimar, el estimador ML aumenta el error de estimación conforme se incrementa éste número, en la Figura 14 podemos ver como de 3 a 7 fuentes el error aumenta en casi 4 grados, pero por otra parte el algoritmo MUSIC no podría realizar estimaciones más que para 6 fuentes. Aun cuando para el estimador ML no existe limitación matemática para el número de fuentes a estimar, como es el caso del algoritmo MUSIC, eventualmente el error de estimación, al ir aumentando, acotará el máximo número posible de fuentes.

Por último, cuando nosotros variamos el parámetro de separabilidad, es decir cuando alejamos o acercamos las fuentes de interés en posición angular, el error de estimación también varía, ver Figura 15. El estimador ML, aún cuando las fuentes se encuentren a un grado de separación, provee una distinción de ambas fuentes, el algoritmo MUSIC en ocasiones no logra distinguir ambas fuentes cuando la separación disminuye a menos de 8 grados o más, dependiendo del nivel de SNR y del número de fuentes.

Con todo lo dicho anteriormente, podríamos ya establecer las fronteras en donde el algoritmo MUSIC arroja resultados útiles. A partir de niveles de SNR aproximados a 8 dB, con separaciones angulares entre fuentes mayores a 8 grados y para un número de muestras mínimo de 1000, y con un máximo de 6 fuentes a localizar. En el caso del estimador ML estos límites se encuentran más allá, con niveles de SNR incluso de hasta -5 dB, separaciones de fuentes de 1 grado, o incluso menores, factible de utilizarse incluso con solo 30 muestras y con la posibilidad de más de 6 fuentes de interés para ser localizadas.

Por otro lado, aún con todas las ventajas claras que presenta el estimador ML, éste se encuentra limitado por la maximización de la función de verosimilitud, que se refleja en el alto costo computacional. Ya que, a medida que aumentamos el número de fuentes se incrementa el número de maximizaciones.

El estimador de máxima verosimilitud es, en resumen, más robusto, y resuelve el problema de la localización de fuentes de manera más eficiente bajo las consideraciones

de: Campo cercano, niveles bajos de SNR, y el aislamiento de fuentes de interés que se encuentren muy cercanas.

4.5 Conclusiones

Hasta este punto concluiríamos con el capítulo de simulaciones, a través de éstas, se han evaluado las bondades, características, ventajas y desventajas del estimador de Máxima Verosimilitud, para comprobar su robustez y viabilidad bajo condiciones poco estudiadas y escenarios adversos de comunicaciones inalámbricas. Tanto las consideraciones, como los escenarios evaluados se escogieron intentando dar una aproximación, lo mas fiel posible, a escenarios reales, donde pudieran aplicarse este tipo de estimadores, y para resolver casos reales de localización de fuentes.

Podríamos adelantar conclusiones sobre las bondades del estimador ML, basados en los resultados y en la comparación con el algoritmo MUSIC, pero esto corresponde al siguiente capítulo, donde se presentarán el análisis y las conclusiones finales con respecto al estimador y al trabajo que se ha realizado.

Capítulo V Conclusiones

De acuerdo a la metodología de investigación establecida, y tomando en cuenta el planteamiento del problema y objetivo de esta tesis se pueden concluir y mencionar simultáneamente las principales aportaciones y resultados de este trabajo de investigación.

5.1 En cuanto al Canal Radio y el Modelo de campo cercano.

Se modeló y simuló la señal generada por las fuentes de interés, propagada por el canal radio y que es recibida en el agrupamiento de antenas, para esto se tomaron en cuenta las consideraciones hechas para el modelo del estimador bajo los escenarios reportados en el capítulo 4. Los aspectos más importantes se listan a continuación:

- Se generó un modelo de canal radio considerando las características de un ambiente de comunicaciones de banda estrecha, con ruido Gaussiano con media cero. Para esto se diseñó un programa realizado sobre la plataforma de simulación MATLAB.

- Para las simulaciones se consideraron sólo dos fuentes de interés ya que el tiempo de procesamiento debido a la cantidad de datos analizados y al algoritmo EM limitaba este factor, a excepción del estudio donde se evalúa este parámetro. Aún así los programas requirieron de un gran recurso computacional.
- Se realizó el modelado y simulación de la incidencia de señales para un escenario de campo cercano, lo cual es de estudio reciente, introduciendo nuevos términos que modelaran esta condición, por tanto también las ventajas y desventajas concernientes.
- No se consideró movilidad en los usuarios, debido a la desventaja del alto costo computacional del estimador de máxima verosimilitud.
- Uno de los principales problemas aquí fue la generación de la atenuación de la señal, ya que la mayoría de los estudios y modelos se han realizado para la condición de campo lejano y cuando se aplicaban a campo cercano no funcionaban de la manera esperada, para resolver esto se tuvo que generar un modelo matemático adecuado a esta condición.
- Otra problemática fue la generación del ruido ya que generar un ruido para cada una de las fuentes de interés significaba mayor tiempo de procesamiento, para resolver esto se optó por generar un solo vector de ruido a partir de la potencia RMS de todas las señales y aplicado a la salida del agrupamiento de antenas.

Las principales aportaciones que se pueden mencionar son:

- El estudio y análisis de la condición del campo cercano, el cual nos permite no solo la aplicación sobre otro tipo de sensores (como micrófonos o sensores telúricos). También introduce la ventaja de poderse estimar varios parámetros a partir de un conjunto de muestras y abre la posibilidad de trabajos futuros a partir de este escenario.
- Además, las consideraciones de simulación se realizaron para que fueran lo más apegadas a las condiciones reales.

5.2 En cuanto a los métodos de localización de fuentes.

Una vez obtenidas las señales incidentes en el agrupamiento de antenas, se modelaron y simularon el estimador de máxima verosimilitud y el algoritmo MUSIC, a fin de resolver el problema de la localización de fuentes.

- Se modeló y simuló el estimador de **máxima verosimilitud incondicional** basado en el algoritmo de **maximización de la esperanza**, para la condición de campo cercano.
- Se utilizó el algoritmo de maximización de la esperanza, para obtener los valores que maximizaran a la función de verosimilitud. Se escogió este algoritmo debido

a sus ventajas al momento de trabajar con datos incompletos y corrompidos por factores externos.

- Se modeló y simuló el algoritmo MUSIC bajo la condición de campo cercano. A fin de poder utilizarlo como referencia para comparar las prestaciones del estimador ML bajo las mismas condiciones.
- Los principales problemas tienen que ver con la maximización de la función de verosimilitud, ya que se tuvo que trabajar con expresiones analíticas con matrices, además de que era un tema completamente nuevo y a que no existen muchas fuentes sobre el tema aplicado al problema de interés. El primer problema fue la convergencia del algoritmo, esto se resolvió a través de múltiples pruebas y de la revisión constante de la teoría. El segundo problema fue el alto costo computacional requerido por el algoritmo EM, ya que la maximización se debe realizar para cada fuente de interés considerada. Para esto se decidió simular solo dos fuentes de interés, además se hizo mano de recursos de cómputo más potentes, en nuestro caso se utilizó tiempo de la supercomputadora Calafia de CICESE.
- Otro de los problemas fue la congruencia de valores en las matrices, esto también afectaba a la convergencia del estimador, ya que algunas matrices se generaban con valores muy pequeños o demasiado altos para lo que se esperaba, este problema se resolvió al utilizar un modelo de pérdidas adecuado a la condición de campo cercano.

Las principales aportaciones en este punto fueron:

- El estudio y análisis del estimador de ML, el cual ha sido poco reportado en trabajos con el mismo problema de interés que esta tesis. Incluso dentro del grupo de investigación que sustenta este trabajo, es la primera vez que se realiza un estudio semejante. Y aún más cuando se incluye la condición de campo cercano.
- Dadas las propiedades de este estimador se genera la posibilidad de aplicar esta herramienta a otro tipo de sensores, áreas de estudio o para sistemas de señales de banda ancha.

5.3 En cuanto a la evaluación en las prestaciones del estimador de Máxima Verosimilitud.

Una vez que se ha analizado y presentado el modelado matemático del estimador ML y el algoritmo MUSIC y que se han realizado simulaciones para verificar su funcionamiento en el problema de la localización de fuentes, es necesario evaluar los aspectos que nos permitirán revisar las prestaciones y bondades del estimador ML comparado con los métodos basados en subespacios, generando así las estadísticas que apoyen este análisis. Respecto a este tema se resaltan los siguientes aspectos:

- Se realizaron simulaciones estilo Montecarlo con la finalidad de dar una representación más real de lo que sería el comportamiento del estimador ML y del algoritmo MUSIC en ambientes reales. Además con la naturaleza aleatoria en la que se basó el modelado matemático del estimador, de las señales de interés y del ruido, se hace lógico el pensar en este tipo de simulaciones. Para este fin, se realizaron 50 repeticiones de cada experimento. Y se calculó el error de estimación RMS a partir de los valores estimados por cada uno de los estimadores. El tiempo total de simulación para cada una de las estadísticas que se presentan a lo largo del capítulo 4 es de aproximadamente 60 horas con 20 minutos
- Se computó el tiempo de ejecución requerido del estimador ML a fin de encontrar las limitaciones del mismo. De acuerdo a las estadísticas presentadas, en el capítulo 4, observemos que ésta es una limitante importante para este estimador, debido principalmente al tiempo necesario para hacer que el algoritmo EM converja. Es en este punto donde se presenta la desventaja más importante del estimador al impedir su utilización para aplicaciones en tiempo real, y el requerimiento de altos recursos computacionales.
- Se evaluó el comportamiento del estimador ML y el algoritmo MUSIC para variaciones del nivel de ruido, generando un nivel de SNR distinto para cada experimento. El estimador ML arrojó un menor error de estimación en comparación con los errores presentados por el algoritmo MUSIC.

- Se evaluaron las prestaciones del estimador ML variando el número de muestras tomadas para realizar la estimación, con esto se evaluó la robustez del estimador ML ya que esta es una de las mayores limitaciones de los métodos de subespacios. Con esta estadística se pudo evaluar la dependencia del estimador ML a este parámetro. El error arrojado por el estimador ML no se incrementa de manera importante si se reduce el número de muestras incluso hasta 30, caso que para los métodos basados en subespacios deja de arrojar datos útiles.
- Se evaluó, también, el número de fuentes de interés a detectar, con esto se evalúa la limitación del estimador ML a este parámetro, el cual es una importante limitación para los métodos de subespacio, ya que sólo pueden detectar $n-1$ fuentes, donde n es el número de elementos del agrupamiento. En este caso, para el algoritmo MUSIC, el número máximo de fuentes posible a localizar se encuentra en 6, mientras que el estimador ML pudo detectar hasta 7 fuentes. No se llevaron a cabo simulaciones con más fuentes ya que, como hemos dicho, se tiene la limitante del tiempo. Esta simulación no se ha encontrado reportada en otros trabajos, por lo que se considera una aportación importante.
- También se evaluó la separabilidad entre fuentes, es decir, que tan cercanas, angular o espacialmente, podrían estar las fuentes de interés y con que definición es estimada su posición por el estimador ML. Mientras que para el algoritmo MUSIC la menor separación posible se encuentra alrededor de los 8 o 9 grados, el estimador ML puede detectar fuentes que se encuentren incluso a 1 grado de

separación angular. Al igual que la simulación anterior no se han encontrado estadísticas semejantes reportadas en otros trabajos, por lo que también se pueden catalogar como una aportación importante de este trabajo.

- Los límites donde el algoritmo MUSIC arroja resultados útiles se encuentra a partir de niveles de SNR aproximados a 8 dB, con separaciones angulares entre fuentes mayores a 8 grados y para un número de muestras mínimo de 1000, y con un máximo de 6 fuentes a localizar. En el caso del estimador ML estos límites se encuentran más allá, con niveles de SNR incluso de hasta -5 dB, separaciones de fuentes de 1 grado, o incluso menores, factible de utilizarse incluso con solo 30 muestras y con la posibilidad de más de 6 fuentes de interés para ser localizadas.
- Por otro lado, aún con todas las ventajas claras que presenta el estimador ML, éste se encuentra limitado a dos factores: La maximización de la función de verosimilitud, ya que esto se refleja en el alto costo computacional, y lo cual incrementa el tiempo de procesamiento. Y a que a medida que aumentamos el número de fuentes se incrementa el número de maximizaciones.
- El estimador de máxima verosimilitud es, en resumen, más robusto, y resuelve el problema de la localización de fuentes de manera más eficiente bajo las consideraciones de: Campo cercano, niveles bajos de SNR, con cantidades de muestras pequeñas y el aislamiento de fuentes de interés que se encuentren muy

cercanas. Pero esta limitado a aplicaciones que no necesiten respuestas en tiempo real.

Las aportaciones principales de este análisis y evaluación son:

- Localización de fuentes (DoA) por medio de técnicas óptimas: Método de inferencia estadística y de estimación.
- La introducción del análisis y evaluación de nuevas métricas: separabilidad de fuentes de interés, el cual no ha sido reportado en otros trabajos; y que nos permite evaluar la robustez del algoritmo bajo condiciones más cercanas a la realidad.
- La introducción del análisis y evaluación de nuevas métricas: número de fuentes de interés a detectar, el cual tampoco ha sido reportado anteriormente, y que nos permite evaluar una de las limitaciones que poseen los métodos basados en subespacios.
- Generación de nuevas estadísticas (punto anterior) que permiten evaluar las prestaciones en la determinación del DoA.
- Tomando en cuenta el estado del arte en que se encuentra el trabajo presentado, la evaluación de los parámetros y las estadísticas nuevas, mencionados anteriormente, presentan una buena referencia para estudiar las bondades del estimador de Máxima Verosimilitud aplicado al problema de la

localización de fuentes en campo cercano y bajo condiciones poco estudiadas.

5.4 Publicaciones resultado del trabajo de investigación

A partir del trabajo de investigación realizado, se generaron las siguientes publicaciones:

- Estimación del DOA con Arrays Optimizados de Separación No-Uniforme. Presentado en el XIX congreso de Instrumentación SOMI.
- Evaluación del Estimador de Máxima Verosimilitud Aplicado a la Localización de Fuentes en Comunicaciones Móviles Celulares. Presentado en el Encuentro de Investigación en Ingeniería Eléctrica, ENINVIE 2005.
- Mejora de la Separabilidad en campo cercano de fuentes, por medio del Estimador UML. Sometido a revisión a la revista IEEE Latinoamérica.

De acuerdo a las metas establecidas y a los resultados esperados que se mencionaron al inicio de este trabajo, y con los resultados obtenidos durante su proceso, se puede decir que se ha cumplido satisfactoriamente el objetivo planteado. Incluyendo aportaciones no contempladas al principio de la investigación.

5.5 Trabajos Futuros.

Como recomendaciones para líneas futuras de investigación, basadas en el problema que aborda este trabajo y como resultado de los problemas generados durante la investigación, se pueden mencionar las siguientes:

- ➔ Optimización de la maximización de la función de verosimilitud. Dado que el algoritmo EM requiere de una gran cantidad de procesamiento y por ende, de tiempo para converger, se requiere, ya sea, de mejoras a este algoritmo o de nuevos métodos que resuelvan el problema de la maximización de manera eficiente y más rápida.
- ➔ Desarrollo de un estimador para señales de banda ancha tales como cdma2000 o W-CDMA, el cual es aplicable a la tercera generación de comunicaciones móviles y para lo cual se necesitan nuevos modelos de canal radio, propagación de señal, función de verosimilitud, estimador, etc.
- ➔ El análisis bajo nuevas métricas y escenarios como:
 - Ambientes de multitrayectorias, lo cual indica trabajar con señales coherentes, y que representaría un escenario aún más real.

- Análisis de movilidad, para lo cual se establecerían nuevos parámetros de velocidad y desplazamiento, y también la velocidad de respuesta del algoritmo.

- Análisis para diferentes agrupamientos de antena. El probar este algoritmo para diferentes geometrías de agrupamientos de antenas con diferente número de elementos, y también con la variación de la separación entre elementos de antena.

Apéndice A. Método de estimación, algoritmo MUSIC

Para la estimación de las fuentes de interés (móviles activos), es utilizado el algoritmo MUSIC, propuesto por Schmidt en [Schmidt, 1986], el cual está basado en la determinación de los sub-espacios vectoriales, que explotan la eigenestructura de la matriz de datos, generada mediante el muestreo espacio-temporal. Para el caso general, el algoritmo considera sensores (eg. dipolos) con posición arbitraria y por lo tanto, aplicable a diferentes geometrías de agrupamiento. En el caso de este trabajo, el algoritmo se aplica a agrupamientos lineales con separación entre elementos de antena uniforme.

El modelo para construir la *matriz de datos* $\mathbf{X}$, contiene las K muestras de D fuentes o señales que arriban al agrupamiento de antena. El agrupamiento está compuesto por M elementos omnidireccionales (dipolos), distribuidos sobre el eje x , donde $\theta=0^\circ$. Cada elemento tiene un peso o ganancia I_p (unitaria), con una separación entre elementos dm_p . La distancia de separación es referida al elemento central denotado con “0”, $\Delta=\lambda/4$ es la distancia entre elementos, donde $k \in K = \{k_{\min}, \dots, k_{\max}\}$ es el índice de elementos según se indica en la Figura 17.

La salida $\mathbf{X}(t)$ del agrupamiento es producida por la incidencia de ondas esféricas, provenientes de D fuentes incorreladas de banda estrecha a una distancia r . Las ondas inciden en el agrupamiento en dirección $\theta_i=[\theta_1, \dots, \theta_D]$.

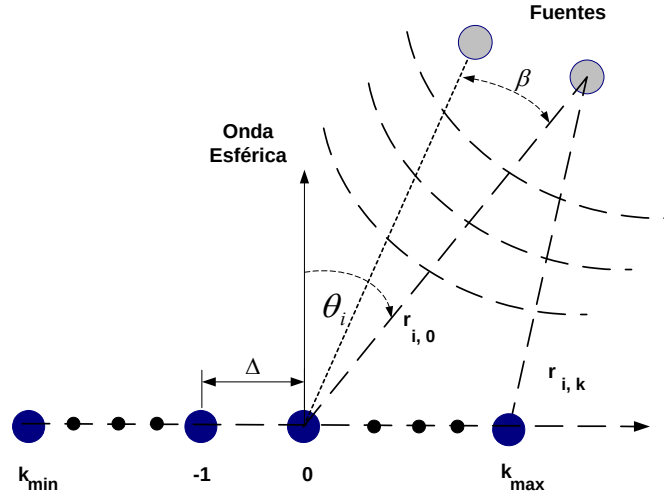


Figura 17. Modelo de datos, distribución de elementos del agrupamiento.

La salida del k -ésimo sensor es expresado como se indica en (58).

$$x_k(t_n) = \sum_{i=1}^d s_i(t_n) e^{j(\mu_i k + \zeta_i k^2)} + n_k(t_n), \quad 1 \leq t_n \leq N, \quad (58)$$

donde $n_k(t_n)$ es el ruido, $s_i(t_n)$ es la envolvente compleja de la i -ésima fuente. Como ya se explico en la sección 3.2, la ecuación (58) contiene el efecto de la curvatura de los frentes de onda en campo cercano.

La salida total del agrupamiento es representada por:

$$X(k) = V(\mu, \zeta) S(k) + N(k), \quad k = 1, 2, \dots, K \quad (59)$$

que depende del vector columna $v(\mu_q, \zeta_q)$ de la i -ésima columna de la matriz de dirección $V(\mu, \zeta)$, con la forma:

$$v(\mu_q, \zeta_q) = \left[e^{j(\mu_q m_{\min} + \zeta_q m_{\min}^2)}, \dots, 1, \dots, e^{j(\mu_q m_{\max} + \zeta_q m_{\max}^2)} \right]^T \quad (60)$$

el cual establece la diferencia de fase entre elementos cuando la onda es inducida en θ_i .

Por otra parte, las fuentes radiantes o terminales móviles son omni-direccionales, representadas en magnitud y fase por $s(t) = \alpha(t) \exp^{-j2\pi ft}$, donde $\alpha(t) = 1/(r^\varepsilon)$ son las pérdidas por propagación, ε es el exponente de pérdidas y f la frecuencia de interés. En $t=r/c+kT$ se considera el tiempo de propagación y el muestreo, en el instante k con periodo T y $\mathbf{N}(t)$ es el ruido, [Schmidt, 1986] [Çekli y Çirpan, 2003].

El algoritmo MUSIC, como ya se ha mencionado, trabaja con subespacios, dichos subespacios son obtenidos a partir del manejo de la matriz de covarianza de los datos muestreados ($\mathbf{R}_{xx}$). Para calcular la matriz de covarianza $\mathbf{R}_{xx}$ considerando un número de muestras K , utilizamos la expresión (61), donde H denota la matriz conjugada transpuesta. Descomponiendo la matriz de covarianza $\mathbf{R}_{xx}$ en sus eigen-valores y eigen-vectores se obtienen los valores y vectores característicos [Schmidt, 1986]. Ordenando los valores característicos de mayor a menor, $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{m-1} \geq \lambda_m$. Los $n-m$ vectores propios λ de los valores propios menores forman un espacio ortogonal $\mathbf{E}_n$, llamado *espacio de ruido*. El espacio de señal $\mathbf{E}_s$ se forma por los primeros n vectores característicos.

$$\mathbf{R}_{xx} = \frac{1}{K} \mathbf{x}(t) \mathbf{x}(t)^H \quad (61)$$

$$\mathbf{E}_s = [\lambda_1 \quad \dots \quad \lambda_n] \quad (62)$$

$$\mathbf{E}_n = [\lambda_{n+1} \quad \dots \quad \lambda_m]$$

El algoritmo MUSIC estima la dirección de arribo de cada usuario localizando los picos del espectro de ruido y señal mediante:

$$P_{MUSIC}(\mu, \varsigma)_n = \frac{1}{\mathbf{a}^H(\mu, \varsigma) \mathbf{E}_n \mathbf{E}_n^H \mathbf{a}(\mu, \varsigma)} \quad (63)$$

$$P_{MUSIC}(\mu, \varsigma)_s = \mathbf{a}^H(\mu, \varsigma) \mathbf{E}_s \mathbf{E}_s^H \mathbf{a}(\mu, \varsigma). \quad (64)$$

Los valores del espectro se calculan al sustituir cada valor posible, dentro de un rango) de μ y ς en el vector de dirección $\mathbf{a}(\mu, \varsigma)$, y se eligen los picos mayores como representantes de la localización de las fuentes de interés posteriormente solo se despeja el valor θ , el cual nos indica la posición angular de la fuente.

Referencias

- Andrade, A. y Covarrubias Rosales, D. 2003. *Radio Channel Spatial Propagation Model for Mobile 3G in Smart Antenna Systems*, IEICE Trans. Commun. E86-B. (1): 213-220.
- Arceo Olague, J. G., Bonilla Hernández, D., Covarrubias Rosales, D. 2004. *Estimación del DOA con Arrays Optimizados de Separación No-Uniforme*. XIX congreso de Instrumentación SOMI: memorias en CD
- Buon K., L. 2002. *Applications of Adaptive Antennas in Third-Generation Mobile Communications Systems*. PhD Dissertation, Curtin University of Technology-Australian Telecommunications Research Institute.
- Çekli, E. y Çirpan, H.A. 2003. *Unconditional Maximum likelihood Approach for localization of Near-Field*. AEÜ Int. J. of Electron. and Commun. 57(1):9–15.
- Çirpan, H.A. y Çekli, E. 2002. *Deterministic Maximum Likelihood Approach for Localization of Near-field Sources*. IEEE Trans. on Signal Processing. 46: 709-719.
- Covarrubias Rosales, D. 2004. No publicado. *Antenas Inteligentes para Comunicaciones Moviles Celulares*. Apuntes de la asignatura Antenas Inteligentes. Grupo de Comunicaciones Inalámbricas. Centro de Investigación Científica y de Educación Superior de Ensenada.
- Delmas, J.P. y Meurisse, J. 2003. *Robustness of narrowband DOA algorithms with respect to signal bandwidth*. Elsevier-Signal Processing, 83(3): 493-510.
- Dempster, A.P., Laird, N.M., y Rubin, D.B. 1977. *Maximum likelihood from incomplete data via the EM algorithm*. Journal of the Royal Statistical Society, Series B (Methodological). 39(1):1-38.

- Ertel R.B., y Reed, J.H. 1999. *Angle and time of arrival statistics for circular and elliptical scattering models*. IEEE J. Sel. Areas Commun., 17(11): 1829-1839.
- Feder, M. y Weinstein, E. 1988. *Parameter estimation of Superimposed Signals Using the EM algorithm*. IEEE Trans. On Acoustics, Speech, and Signal Processing. 36(4).
- Hochwald, B. y Nehorai, A. 1994. *Concentrated Cramer Rao Bound Expression*. IEEE Trans. on Information Theory. 40: 363–371.
- Kay, S. M. 1993. *Fundamentals of Statistical Signal Processing: Estimation Theory*. Prentice Hall, Inc. Primera edición, New Jersey. 595 pp.
- Liberti, J.C. Jr. y Rappaport, T.S. 1999. *Smart Antennas for Wireless Communications: IS-95 and Third Generation CDMA Applications*. Prentice Hall, Primera edición, New Jersey. 376 pp.
- Miller, M. I. y Fuhrmann, D. R. 1990. *Maximum-Likelihood Narrow-Band Direction Finding and the EM Algorithm*. IEEE Trans. Acoustics. Speech, and Signal Processing. 38: 1560–1577.
- Panduro, M.A., Covarrubias Rosales, D., Brizuela, C.A. y Marante, F.R. En proceso. *A Multi-Objective Approach in the Linear Antenna Array Design*. Aceptada para publicarse en International Journal of Electronics and Communications.
- Rappaport, T.S., 2002. *Wireless Communications: Principles and Practice*. Prentice Hall. Segunda edición. New Jersey. 707 pp.
- Schmidt, R.O. 1986. *Multiple Emitter Location and Signal Parameter Estimation*. IEEE Transactions on Antennas and Propagation. 34(3): 276-280.
- Stoica, P. y Nehorai, A. 1989. *MUSIC, Maximum Likelihood, and Cramer-Rao Bound*. IEEE Trans. On Acoustics, Speech, and Signal Processing. 37(5): 720-741.

- Van Trees, H.L. 2002. *Optimum array procesing. Part IV of detection, estimation, and Modulation Theory*. John Wiley & sons, Inc. Primera Edición. New York. 1443 pp.
- Ziskind I. y Wax M. 1998. *Maximum Likelihood Localization of Multiple sources by Alternating Projection*. IEEE Trans. on Acoustics, speech, and Signal processing. 36(10): 1553-1560.