



***FACULTAD
DE
CIENCIAS***

REPARTO DE SECRETOS

(Secret Sharing Schemes)

Trabajo de fin de Grado
para acceder al

GRADO EN MATEMÁTICAS

Autora: María Arce Muela

Director: Daniel Sadornil Renedo

Octubre-2016

Índice general

Resumen	1
Abstract	1
1. Introducción	3
2. Reparto de Secretos	7
2.1. Estructuras de acceso	7
2.2. Modelo general de un esquema de reparto de secretos	12
2.3. Medida de la información y entropía	13
3. Teoría de Códigos Correctores	17
3.1. Distancia mínima de un código y tasa de información	18
3.2. Matrices generadoras de códigos	21
3.3. Códigos MDS	24
4. Ejemplos de esquemas para repartir secretos	25
4.1. Esquema de Shamir	25
4.2. Construcción geométrica	29
4.3. Construcción de Ito, Saito y Nishizeki	33
4.4. Construcción vectorial	36
4.5. Construcción mediante grafos	42
4.6. Construcción basada en códigos	45
Bibliografía	53

Resumen:

Dado cualquier tipo de información que por uno u otro motivo deba de ser confidencial, los esquemas de reparto de secretos tienen como objetivo construir pedazos de información a partir de dicha información privada, a lo que denominaremos secreto, y repartirla entre un grupo de participantes. Además, mediante la tarea de reparto se legitimará a través de distintos tipos de algoritmos a las agrupaciones de participantes que serán capaces de recuperar el secreto. En este trabajo se muestra la idea general de reparto de secretos, así como distintos ejemplos basados en diferentes conceptos matemáticos.

Palabras clave: Reparto de secretos, Códigos correctores, Teoría de la Información, Autenticación, Protocolo.

Abstract:

If we have any kind of information which must be kept highly confidential, secret sharing schemes can distributing this private information, called secret, giving shares (pieces of the secret) to the players. Furthermore, only when special conditions be fulfilled, a group of players will be able to discover the value of the secret. In this paper, we're going to present a general idea of secret sharing schemes and different examples of them, which are based of different mathematical disciplines.

Key words: Secret Sharing, Error correcting codes, Information Theory, Authentication, Protocol.

Capítulo 1

Introducción

La distribución de los números de cuentas bancarias, los códigos de lanzamiento de un misil, o algunas claves para encriptar datos, que siempre pueden ser necesarias, por ejemplo, para guardar información confidencial de una gran empresa, están basados en esquemas para repartir secretos. Estos esquemas adquieren una gran importancia cuando la información debe de tratarse con especial cuidado, ya que si en cualquiera de los ejemplos anteriores se expone parte de dicha información, puede resultar crítico o desastroso para el ámbito donde se haya producido el error, pudiendo llegar a provocar graves consecuencias de ámbito global. Una vez podamos asegurar la privacidad de la información, deberá ser posible construir un protocolo de manera que si tenemos una agrupación adecuada, ésta sea capaz de recuperar la totalidad de la información de forma rápida y segura. Para ello, se utilizarán diferentes esquemas para repartir secretos, algunos de los cuales estudiaremos con detalle más adelante.

En resumen, los esquemas de reparto de secretos son una importante herramienta en seguridad y criptografía, con la cual no debe confundirse, ya que puede servirnos para guardar una clave con éxito, dividiendo su información en pedazos de forma que pueda ser recuperada sólo por quien deba tener acceso a ella, pero nunca será un protocolo criptográfico. Tan sólo es eso, una herramienta que puede estar al servicio de la criptografía con el objetivo de aportar mayor seguridad a la hora de guardar dicha clave.

Un ejemplo sencillo de estos esquemas que puede encontrarse en la vida real son las cajas de seguridad que los bancos ponen a disposición de sus clientes. Estas cajas se sitúan dentro de una cámara acorazada y su contenido es estrictamente confidencial, es decir, la persona que lo alquila no tiene que declarar su contenido a la entidad financiera y nadie más tiene porqué saber lo hay en ella. Por lo tanto, podría interpretarse como un esquema umbral, donde el único conocedor de la información es el cliente, y donde es necesaria la puesta en común de las llaves del cliente y el banco para poder acceder a dicha información.

En este trabajo, se va a realizar un estudio de distintos esquemas de reparto de secretos. La idea de realizar esquemas para repartir secretos no irrumpió en la matemática de forma casual. Se había producido una revolución en la ciencia de la información a lo

largo del siglo XX y era necesaria una respuesta contundente a nuevos tipos de problemas o situaciones. Como veremos más adelante, la ciencia de la información comenzó con los trabajos de Shannon [23], *A Mathematical Theory of Communication* y su creencia firme de que cualquier información podía ser medida, aportando el esquema general que describe cómo se transmite la información.

Los esquemas de reparto de secretos, fueron introducidos por primera vez en 1979 por los autores Adi Shamir [22] y George Blakley [7] de manera independiente.

Ambos autores dieron dos soluciones completamente diferentes para un problema común, que consistía en buscar un método para repartir un secreto, o cierta información específica desconocida entre un grupo de participantes. Se dotaba a cada uno de ellos de pedazos de dicha información, o participaciones de forma que sólo algunas agrupaciones de participantes fueran capaces de obtener la información completa poniendo en común sus participaciones.

Intuitivamente podemos decir que el secreto es un puzzle y que a cada persona se le da un conjunto de piezas (que pueden ser repetidas entre sí o varias personas tener piezas comunes) de forma que se pueda montar si un grupo de personas pone sus piezas en común. Como hay piezas repetidas y más de una persona puede tener las mismas piezas, puede haber participantes que queden fuera a la hora de componer el puzzle y a su vez, puede ser que alguno de ellos tenga una o más piezas que resulten determinantes y entonces sea necesario que dicho participante forme parte del grupo de personas que lo reconstruyan.

Los trabajos de Blakley y de Shamir, aún trantándose de dos propuestas totalmente diferentes, convergían en un punto común, su forma de definir las agrupaciones de participantes. Ambos describían un esquema donde, dado un grupo de n participantes, solamente las agrupaciones formadas por un determinado número t , fijado de antemano, de dichos participantes o más, podría reunir la información necesaria para recuperar el secreto. Además, cualquier otra agrupación formada por un conjunto de personas menor no podría obtener ningún tipo de información acerca del mismo. A este tipo de esquemas donde la recuperación del secreto viene determinada por el número de participantes de una agrupación se les denominaría a partir de ese momento como esquemas (t, n) -umbrales.

De estos primeros modelos se pudo obtener otra importante característica para esquemas posteriores. La idea de que las agrupaciones permitidas, que a partir de ahora denominaremos agrupaciones autorizadas, puedan recuperar el secreto en su totalidad, y que por el contrario cualquier agrupación de personas que no sea autorizada no obtenga ningún tipo de información uniendo sus participaciones, será determinante a la hora de realizar distintos tipos de construcciones. Tal es su importancia que dió lugar a lo que se conoce como esquemas perfectos, y son todos aquellos que cumplen con dicha condición.

Posteriormente a Shamir y Blakley, Ito, Saito y Nishizeki [15] describieron un esquema perfecto, mostrando que cualquier estructura de acceso (conjunto definido por todas las agrupaciones autorizadas posibles) puede realizar un esquema de reparto de

secretos. Al dar una solución global para realizar esquemas, muchas de estos esquemas son muy poco eficientes, ya que para poder realizarlos a veces es necesario dar demasiada cantidad de información a cada participante. Para poder entender estos esquemas en su totalidad tendremos que esperar al capítulo 4, cuando se hayan introducido los conceptos necesarios y donde se dará una definición con más detalle de los mismos.

Al hablar de estructuras de acceso poco eficientes, podemos sospechar que necesitaremos una medida para controlar dicha eficiencia, lo que denominaremos tasa de información y que como veremos será uno de los parámetros principales en la construcción de un esquema. Imaginemos que el secreto es 1 bit, no tendría sentido distribuir a cada participante 100 bits para recuperar un secreto que tenga un sólo bit de información. La tasa de información va a establecer entonces una relación entre la cantidad de información de las participaciones y la cantidad de información del secreto. Además, cuando el valor de esta razón sea la unidad diremos que se trata de un esquema ideal.

Después de haber explicado brevemente el núcleo de nuestro trabajo, podemos establecer que nuestro objetivo será estudiar los esquemas de reparto de secretos y dar diferentes ejemplos de ellos procedentes de muy distintas ramas matemáticas.

Una vez introducido el objetivo de este trabajo y una idea general de en qué consiste un problema de reparto de secretos, vamos a dar a continuación una definición más precisa y detallada del mismo.

Considerando un grupo de n participantes $P_1, P_2, \dots, P_n$, un esquema de reparto de secretos pretende distribuir un secreto k de forma que cada participante P_i obtenga cierta información s_i . A cada una de estas partes correspondientes a cada participante las denominaremos participaciones. Por otra parte, deberá ser posible recuperar el secreto siempre y cuando dispongamos de las agrupaciones de participantes adecuadas y no deberá de poder conocerse ningún tipo de información sobre el mismo en caso contrario. Para ello, un esquema de reparto de secretos constará de dos partes:

- Distribución de las participaciones: se designa un gestor D para calcular y distribuir las distintas participaciones s_i entre los participantes. Dicho gestor, por lo tanto, es el único conocedor del secreto.
- Reconstrucción del secreto: si un conjunto de participantes P_i tiene las participaciones necesarias procederá a recuperar el secreto combinándolas convenientemente. A este conjunto de participantes se le denominará agrupación autorizada y como hemos hecho notar, sólo las agrupaciones autorizadas serán capaces de calcular el secreto a partir de sus participaciones.

El trabajo sigue la siguiente estructura:

- En el Capítulo 2 definiremos los conceptos básicos que van a ser necesarios a lo largo del trabajo en cuanto a reparto de secretos se refiere. Daremos una definición formal de los conceptos que se han presentado en la introducción, como son un esquema de reparto de secretos, una estructura de acceso, un esquema perfecto o esquema umbral y la tasa de información entre otros [18].

- En el Capítulo 3 vamos a realizar un breve resumen de la teoría de códigos correctores, describiendo lo que es un código lineal, la distancia mínima de un código y las matrices generadora y de control entre otros resultados. Además, introduciremos los tipos de códigos correctores que sean necesarios para alguna de las construcciones posteriores de reparto de secretos, así como los conceptos de capacidad correctora y detectora de los códigos correctores. Cabe destacar que la mayoría de estos términos han sido estudiados con anterioridad en la asignatura del Grado de matemática discreta y por lo tanto, no vamos a desarrollar este capítulo con mayor detalle.
- En el Capítulo 4 vamos a presentar distintos ejemplos de esquemas de reparto de secretos, cada uno de ellos procedentes a distintas ramas de la matemática. Comenzaremos con el primer esquema de reparto por excelencia, el Esquema de Shamir y a continuación daremos una construcción geométrica basada en el método de Blakley, por lo que veremos dos puntos de vista diferentes de un esquema umbral. Una vez hayamos presentado dos construcciones basadas en esquemas umbrales, introduciremos un esquema monótono que puede describir una estructura de acceso cualquiera de la mano de Ito, Saito y Nishizeki. Además daremos otros esquemas, como pueden ser una construcción vectorial y otra construcción basada en grafos. Por último, realizaremos una construcción basada en códigos correctores.

Capítulo 2

Reparto de Secretos

En este capítulo introduciremos distintas nociones necesarias para el desarrollo de este trabajo, definiendo formalmente lo que es una estructura de acceso asociada a un esquema de reparto de secretos y algunas de sus propiedades principales [2] [14]. También describiremos conceptos como el modelo general de un esquema de reparto de secretos, o la tasa de información del mismo. Una vez que quedan definidos todos estos conceptos entenderemos como se realiza un esquema de reparto mediante un modelo de reglas de distribución.

2.1. Estructuras de acceso

Realizar un esquema de reparto de secretos es equivalente a repartir un secreto k dividiéndolo en trozos de información o participaciones y repartiéndolas entre un grupo de participantes de forma que ciertos subgrupos de éstos sean capaces de recuperar la información completa y otros no, a partir de las participaciones que les han asignado.

Cabe destacar que de ahora en adelante, cuando se hable de repartir o dividir un secreto, no estamos entendiéndolo como partir el secreto en trozos de información, si no que hace referencia al hecho de que a partir del secreto se generan las participaciones que posteriormente se repartirán entre los distintos participantes (de hecho, las mismas participaciones no tienen porqué ser parte del secreto).

Dado un conjunto de participantes $\mathcal{P} = \{P_1, \dots, P_n\}$, una estructura de acceso consiste en un conjunto Γ tal que $\Gamma \subseteq P(\mathcal{P})$. En estos términos, se dice que el esquema realiza la estructura de acceso Γ :

- $D \notin \mathcal{P}$: gestor del esquema.
- $\mathcal{K}$: conjunto de secretos.
- $\Gamma \subseteq P(\mathcal{P})$: subconjunto de partes de $\mathcal{P}$ cuyos elementos son las denominadas agrupaciones autorizadas (subconjuntos de partes de $\mathcal{P}$ que pueden calcular el secreto).

- $\Delta \subseteq P(\mathcal{P})$: subconjunto de partes de $\mathcal{P}$ cuyos elementos son las agrupaciones no autorizadas o agrupaciones que no pueden obtener ningún tipo de información acerca del secreto.

Es fácil observar que $\Gamma \cap \Delta = \emptyset$ y que $\Gamma \cup \Delta = P(\mathcal{P})$.

- $\mathcal{S}$: conjunto de las participaciones a repartir.
- $s_i \subseteq \mathcal{S} \ \forall i \in \{1, \dots, n\}$: conjunto de las participaciones que recibe cada participante P_i .

Cuando se desea obtener un secreto $k \in \mathcal{K}$ a partir de un esquema, el gestor D repartirá a cada participante $P_i \in \mathcal{P}$ una participación $s_i \in \mathcal{S}$ mediante unas reglas de distribución.

En resumen, la finalidad de una estructura de acceso es describir que conjuntos de participantes pueden reconstruir o no el secreto cuando se unen. Cuando un esquema se realiza a partir de una estructura de acceso Γ cuyos conjuntos son los subconjuntos de $P(\mathcal{P})$ compuestos por al menos t participantes se dice que es un (t, n) -esquema umbral. Es decir, basta con que t participantes cualesquiera compartan sus participaciones para obtener el secreto k .

A lo largo de este trabajo intentaremos realizar construcciones de forma que cuando se reúnan todos los participantes de una agrupación autorizada recuperen el secreto completamente y si por el contrario una agrupación no autorizada pone en común sus participaciones no sea capaz de obtener ningún tipo de información adicional. En la práctica, no siempre es posible realizar este tipo de esquemas, a los que denominaremos esquemas perfectos, pero si buscaremos aproximarnos todo lo posible a ellos.

Definición 2.1 *Un esquema para repartir secretos que realiza una estructura de acceso (Γ, Δ) se dice que es perfecto si:*

1. *Dado un conjunto autorizado $A \in \Gamma$, éste puede determinar el secreto $k \in \mathcal{K}$ si computa sus participaciones $s_{i_A} \in \mathcal{S}$.*
2. *Dado un conjunto no autorizado $B \in \Delta$, éste no puede obtener ninguna información del secreto poniendo en común sus participaciones.*

Es claro que, a partir de un conjunto P , existen numerosas estructuras de acceso, (de hecho $\mathcal{P}(P)$), pero si no se tiene en cuenta "la diferencia entre participantes" pueden existir varias equivalentes.

Definición 2.2 *Sean Γ_1, Γ_2 dos estructuras de acceso definidas sobre un mismo P . Si existe alguna permutación σ en el conjunto de índices $\{1, \dots, n\}$ de forma que:*

$$\{P_{i_1}, \dots, P_{i_r}\} \in \Gamma_1 \iff \{P_{\sigma(i_1)}, \dots, P_{\sigma(i_r)}\} \in \Gamma_2$$

*se dice que ambas estructuras son **equivalentes**, o lo que es lo mismo, que una de las dos estructuras es generada por la anterior renombrando a sus participantes.*

Ejemplo 2.1 Sea $\mathcal{P} = \{P_1, P_2, P_3, P_4\}$. Las estructuras $\Gamma_1 = \{\{P_1, P_2\}, \{P_2, P_4\}, \{P_1, P_3, P_4\}, \{P_1, P_2, P_3, P_4\}\}$ y $\Gamma_2 = \{\{P_1, P_4\}, \{P_4, P_2\}, \{P_1, P_2, P_3\}, \{P_1, P_2, P_3, P_4\}\}$ son equivalentes, ya que podemos renombrar $P_{\sigma(4)} = P_2$ y $P_{\sigma(2)} = P_4$.

Si volvemos a la definición de (t, n) -esquema umbral, cuando denominamos agrupación autorizada a todas aquellas que constan de más de t participantes, es fácil observar que si por ejemplo tenemos la agrupación $\{P_1, P_2, \dots, P_t\}$, entonces $\{P_1, P_2, \dots, P_t, P_{t+1}\}$ será también una agrupación autorizada. Partiendo de este caso concreto, vamos a definir formalmente que tipo de estructuras son aquellas donde se da esta circunstancia.

Definición 2.3 Sea Γ una estructura de acceso, se dice que es una **estructura de acceso monótona** si cumple que $\emptyset \notin \Gamma$, y dado $B \in \Gamma$ y $B \subseteq C \subseteq P$, entonces $C \in \Gamma$.

En otras palabras, que una estructura de acceso sea monótona quiere decir que al añadir participantes nuevos a cualquiera de los subconjuntos autorizados de Γ éstos no perderán dicha condición y que si, por el contrario, eliminamos algún participante de uno de los subconjuntos no autorizados, el subconjunto modificado tampoco será un subconjunto autorizado.

Ejemplo 2.2 Sea $P = \{P_1, P_2, P_3, P_4\}$. Es fácil observar que

$$\Gamma_1 = \{\{P_1, P_2\}, \{P_2, P_3\}, \{P_1, P_2, P_3\}, \{P_1, P_2, P_4\}, \{P_2, P_3, P_4\}, P\}$$

es una estructura de acceso monótona mientras que

$$\Gamma_2 = \{\{P_1, P_2\}, \{P_2, P_3\}, \{P_1, P_2, P_4\}\}$$

no lo es ya que, sin ir más lejos, $\{P_1, P_2, P_3\}$ no es una agrupación autorizada y $\{P_2, P_3\}$ si lo es.

En la mayoría de las construcciones de esquemas de reparto y en todos los ejemplos que veremos en este trabajo en particular, las estructuras que se presentan son estructuras monótonas, ya que parece lógico pensar, que si una agrupación puede recuperar el secreto, esta condición no va a cambiar porque se añada una persona más al grupo.

Al considerar que a partir de ahora vamos a utilizar estructuras monótonas, no tardamos mucho en darnos cuenta de que puede que existan conjuntos minimales que sean agrupaciones autorizadas, es decir, aquellos conjuntos mínimos a partir de los cuales si añadimos algún participante construiremos nuevas agrupaciones autorizadas y si por el contrario eliminamos a alguno de sus participantes los restantes no serán capaces de recuperar el secreto. A continuación, veremos como son estos conjuntos y que utilidad tienen a la hora de trabajar con una estructura de acceso.

Definición 2.4 Sea Γ una estructura de acceso monótona y $A \in \Gamma$. Se dice que A es un **conjunto minimal** si $\forall B \in \Gamma - A$, $A \not\subseteq B$.

Si tenemos un conjunto de participantes donde su cardinal sea bastante grande, al querer dar una definición total de Γ , como hasta ahora venimos haciendo, puede resultarnos difícil trabajar con su estructura de acceso. Esto se debe a que al tratarse de estructuras monótonas a mayor número de participantes, hay un mayor número de agrupaciones autorizadas existentes. Para que resulte más sencillo trabajar con estructuras de acceso que tengan un número elevado de participantes, podemos definir un conjunto que contenga sólo a las agrupaciones minimales, de forma que, gracias a ciertos cálculos éstas determinen la estructura de acceso completa.

Definición 2.5 Sea Γ una estructura de acceso, denominamos base de Γ al subconjunto formado por todas sus agrupaciones minimales. Se denota por Γ_0 y se define de la siguiente forma:

$$\Gamma_0 = \sum_i \gamma_i, \gamma_i = P_{j_1} \cdot \dots \cdot P_{j_{r_i}}$$

Ejemplo 2.3 Si tenemos la estructura de acceso monótona Γ_1 del ejemplo anterior la base de esta estructura quedará determinada por

$$\Gamma_0 = \{P_1 \cdot P_2\} + \{P_2 \cdot P_3\}$$

ya que $\{P_1, P_2\}$ y $\{P_2, P_3\}$ son sus conjuntos minimales.

Además de para simplificar el manejo de una estructura de acceso con muchos participantes, podemos utilizar su base para definir otro tipo de estructuras, las estructuras de acceso duales a una dada, que como veremos a continuación están íntimamente ligadas a ellas.

Definición 2.6 Sea Γ una estructura de acceso. Se llama **estructura dual** de ésta a aquella obtenida sustituyendo en la definición formal de base las sumas por productos y los productos por sumas. Se denota por Γ^* y su base se define como

$$\Gamma_0^* = \sum_i \gamma_i^*$$

Por último, diremos que una estructura es autodual cuando dicha estructura es la misma que su estructura dual.

Otra definición alternativa de estructura dual que no la relaciona directamente con la base de una estructura de acceso pero que nos da una idea más intuitiva del aspecto de la misma es la siguiente.

Definición 2.7 Sea Γ una estructura de acceso, podemos definir también su estructura dual como una agrupación de conjuntos de participantes $A_i \in P$ donde $A_i \in \Gamma^*$ si $A^c \notin \Gamma$.

Ejemplo 2.4 Utilizando el ejemplo anterior, si tenemos la estructura de acceso Γ_1 definida por

$$\Gamma_1 = \{\{P_1, P_2\}, \{P_2, P_3\}, \{P_1, P_2, P_3\}, \{P_1, P_2, P_4\}, \{P_2, P_3, P_4\}, \mathcal{P}\}$$

cuya base y conjuntos minimales son respectivamente

$$\Gamma_0 = P_1 \cdot P_2 + P_2 \cdot P_3 \text{ y } \{P_1, P_2\}, \{P_2, P_3\}$$

Entonces podemos definir la base de su estructura dual como

$$\Gamma_0^* = (P_1 + P_2) \cdot (P_2 + P_3) = P_1 \cdot P_2 + P_1 \cdot P_3 + P_2 + P_2 \cdot P_3$$

Por la definición anterior, es fácil ver que $(\Gamma^*)^* = \Gamma$.

Al igual que hemos establecido en Γ unas agrupaciones minimales a partir de las cuales queda definido todo el conjunto de agrupaciones autorizadas, se puede definir en el conjunto de agrupaciones no autorizadas unos subconjuntos maximales a partir de los cuales, todas las agrupaciones contenidas en éstas estarán en Δ , (es decir, serán no autorizadas).

A continuación veremos como se pueden relacionar las distintas estructuras entre si y como podemos definir unas a partir de otras.

Es fácil ver que el siguiente resultado determina una base para las agrupaciones no autorizadas.

Teorema 2.1 Sea Γ una estructura de acceso. Podemos describir Δ a partir de la estructura dual de Γ de la siguiente forma. Para cada $\gamma_i^* \in \Gamma_0^*$, definimos $\overline{\gamma}_i \in \overline{\Gamma}_1$ como $\overline{\gamma}_i = P - \gamma_i^*$.

Demostracion: Puede verse en [6].

Entonces, el teorema anterior afirma que, si en algún caso disponemos de las agrupaciones no autorizadas, podemos obtener la base dual de Γ a partir de ellas.

Proposición 2.1 El conjunto de las agrupaciones no autorizadas Δ tiene como base $\Delta = \sum_i \overline{\gamma}_i$ con $\overline{\gamma}_i$ el producto de los elementos de una agrupación maximal de Δ .

Demostracion: Dicha base se obtiene de forma trivial aplicando el resultado anterior.

Ejemplo 2.5 Sea $\Gamma_0^* = (P_1 + P_2) \cdot (P_2 + P_3)$ la base de la estructura de acceso dual del ejemplo 2.4. A partir de sus conjuntos $\gamma_1^* = P_2$, $\gamma_2^* = (P_1, P_3)$, $\gamma_3^* = \{P_2, P_3\}$ y $\gamma_4^* = \{P_1, P_2\}$ podemos calcular las agrupaciones maximales de Δ :

$$\begin{aligned} \overline{\gamma}_1 &= \mathcal{P} - \gamma_3^* = \{P_1, P_2, P_3, P_4\} - \{P_2\} = \{P_1, P_3, P_4\} \\ \overline{\gamma}_2 &= \mathcal{P} - \gamma_2^* = \{P_1, P_2, P_3, P_4\} - \{P_1, P_3\} = \{P_2, P_4\} \\ \overline{\gamma}_3 &= \mathcal{P} - \gamma_4^* = \{P_1, P_2, P_3, P_4\} - \{P_2, P_3\} = \{P_1, P_4\} \\ \overline{\gamma}_4 &= \mathcal{P} - \gamma_1^* = \{P_1, P_2, P_3, P_4\} - \{P_1, P_2\} = \{P_3, P_4\} \end{aligned}$$

Durante la introducción, cuando describimos las estructuras de acceso hablamos del caso concreto del (t, n) -esquema umbral donde todos los subconjuntos formados por t o más participantes son suconjuntos autorizados. A continuación veremos un resultado que nos hará notar cómo la estructura dual de un esquema umbral es a su vez estructura de acceso de un esquema umbral diferente cuya demostración puede verse en [2].

Teorema 2.2 *La estructura dual Γ^* de un (t, n) -esquema umbral es a su vez la estructura de acceso Γ de un $(n-t+1, n)$ -esquema umbral. Es fácil ver entonces que la estructura de acceso Γ de un (t, n) -esquema umbral es autodual si, y solo si, $2t = n + 1$.*

Ejemplo 2.6 *Sea Γ la estructura de acceso correspondiente a un $(3, 4)$ -esquema umbral. Entonces*

$$\Gamma = \{\{P_1, P_2, P_3\}, \{P_1, P_2, P_4\}, \{P_1, P_3, P_4\}, \{P_2, P_3, P_4\}, \{P_1, P_2, P_3, P_4\}\}$$

A partir del teorema anterior, la estructura dual de Γ realiza un $(2, 4)$ -esquema umbral. En efecto, sea

$$\Gamma_0 = P_1P_2P_3 + P_1P_2P_4 + P_1P_3P_4 + P_2P_3P_4$$

la base de Γ y estructura dual

$$\Gamma_0^* = (P_1 + P_2 + P_3) \cdot (P_1 + P_2 + P_4) \cdot (P_1 + P_3 + P_4) \cdot (P_2 + P_3 + P_4)$$

Con un poco de detenimiento es fácil deducir que

$$\{P_1, P_2\}, \{P_1, P_3\}, \{P_1, P_4\}, \{P_2, P_3\}, \{P_2, P_4\}, \{P_3, P_4\}$$

son sus agrupaciones minimales autorizadas.

Por lo tanto, Γ_0^ es una estructura de acceso que describe un $(2, 4)$ -esquema umbral.*

2.2. Modelo general de un esquema de reparto de secretos

Una vez definida lo que es una estructura de acceso necesitamos un protocolo o modelo, que será de conocimiento público, para repartir la información de una determinada manera. Para esto hay muchos tipos de esquemas, como veremos en el capítulo 4, y por ahora simplemente vamos a realizar una descripción en términos generales.

Un esquema para repartir secretos tiene como objetivo repartir pedazos de información (participaciones) entre los distintos participantes. Para ello se representa como un conjunto de reglas de distribución donde cada una de ellas es una aplicación $f : P \cup D \longrightarrow \mathcal{S} \cup \mathbf{K}$ que cumple que $f(D) = k$ y $f(P_i) = s_i$ con $1 \leq i \leq n$

Donde $f(D)$ es el secreto a repartir y $f(P_i)$ las participaciones asignadas a cada participante P_i .

Si denotamos al conjunto total de reglas de distribución por $\mathcal{J}$, para cada secreto k su conjunto de reglas de distribución se denota como $J_K = \{f \in \mathcal{J} : f(D) = k\}$ y cuando el gestor quiera repartir un secreto k elegirá aleatoriamente una de las J_k .

Una vez definido lo que es un conjunto de reglas de distribución, vamos a introducir el siguiente resultado que caracteriza los esquemas perfectos (Definición 2.1) en términos de dichas reglas, cuya demostración puede verse en [11].

Teorema 2.3 *sea Γ una estructura de acceso y sea $\mathcal{J}$ un conjunto de reglas de distribución. Si se satisfacen las siguientes propiedades*

1. *Sea $A \in \Gamma$ y $f, g \in \mathcal{J}$. Si $f(P_i) = g(P_i)$ para todo participante P_i de A se tiene que*

$$f(D) = g(D)$$

2. *Sean $B \in \Delta$ y $f \in \mathcal{J}$. Entonces $\exists \lambda \in \mathbb{N}$ (que depende de f y de B) de forma que para cada $k \in \mathbf{K}$*

$$|\{g \in J_K : g(P_i) = f(P_i) \forall P_i \in B\}| = \lambda$$

Podemos decir que la estructura de acceso Γ realiza un esquema perfecto.

Lo que afirma la primera condición es que las participaciones que se reparten a un conjunto autorizado A determinan de forma única el secreto. Por el contrario, la segunda parte asegura que cuando damos participaciones a un conjunto no autorizado hay $\lambda(f, B)$ posibles valores que determinan el secreto.

2.3. Medida de la información y entropía

El núcleo de la noción de reparto de secretos se basa en gestionar información. Ya inicialmente, "dividimos un secreto en fragmentos del mismo", que son trocitos de información y mediante una regla de distribución los repartimos entre los participantes. A partir de ahí, es fácil ver que la mayoría de los conceptos que introducimos son necesarios para realizar esta labor con éxito, y que desde un principio estamos tratando dicha información de manera cuantitativa, es decir, tomamos la información como una unidad de medida. Para tratar la información de forma cuantitativa es necesario poder medirla, y como veremos a lo largo de esta sección es posible calcular la cantidad de información que nos aporta cada suceso mediante el cálculo de probabilidades.

Además de nuestra necesidad por medir la información que estamos tratando de manera general, ya durante la introducción hablábamos de la **eficiencia** de un esquema de reparto de secretos. Dicha eficiencia establecía la relación entre la cantidad de información de las participaciones y la cantidad de información del secreto. Esta relación puede medirse a través de la **tasa de información**.

Para poder definir la tasa de información de un esquema, debemos ver primero como se calcula la cantidad de información del secreto y de las participaciones repartidas a los participantes.

- Un secreto k puede representarse como una cadena de bits. Se define su longitud como $\log_2 |k|$.
- La longitud de cada participación es $\lambda(s_i) = \log_2 |s_i|$ y la media de la longitud de todas las participaciones de cada participante $\lambda(s_j)_{j \in \Gamma}$.

Una vez que sabemos como cuantificar la información existente en un secreto y sus participaciones, podemos definir la tasa de información de la siguiente manera.

Definición 2.8 *Se define la tasa de información de cada participante P_i como la proporción entre la longitud del secreto y la longitud de su participación:*

$$\rho_i = \frac{\log_2 |k|}{\lambda(s_i)}$$

Además, denominaremos tasa de información del esquema a la menor tasa de información de entre sus participantes:

$$\rho = \min \{\rho_i / i \in 1, \dots, n\}$$

De la definición anterior, es fácil observar que cuanto mayor sea la cantidad de información que hay que repartir a cada participante P_i , mayor será la longitud de sus participaciones y menor será la tasa de información. Nuestro objetivo a la hora de realizar una construcción será obtener el valor más alto posible para la tasa de información, ya que esto significará que la longitud de las participaciones de cada participante es menor, y que por tanto no ha sido necesario repartir a cada participante demasiada cantidad de información. Además, vamos a introducir otro resultado que relaciona la tasa de información con los esquemas para repartir secretos perfectos e ideales.

Lema 2.1 *Siempre que tengamos un esquema perfecto, $p \leq 1$. En el caso de que $p = 1$ diremos que el esquema es ideal.*

Dejando momentáneamente a un lado la eficacia de un esquema para repartir secretos, otro de los objetivos de esta sección es determinar la cantidad de información que tiene una agrupación cualquiera y compararla con la del secreto para comprobar si va a ser posible o no recuperarla en su totalidad.

Para ello, va a ser necesario introducir un concepto que mida la cantidad de información que nos aporta un suceso, su entropía, como puede ser en nuestro caso la puesta en común de las participaciones de una agrupación autorizada. Al igual que hablamos de medir la cantidad de información, también podemos referirnos a la incertidumbre que produce un suceso, ya que es equivalente a la información que nos aporta una vez se ha realizado.

Definición 2.9 : Sea $p : \mathbf{X} \mapsto [0, 1] \subset \mathbb{R}$ una función de distribución que está definida sobre una variable aleatoria $\mathbf{X}$. Definimos la **entropía de Shannon** como:

$$H(p) = - \sum_{x \in \mathbf{X}} p(x) \cdot \log_2 p(x)$$

En otras palabras, la entropía de Shannon mide la información o la incertidumbre media sobre qué elemento de $\mathbf{X}$ ha sido elegido siguiendo la función de probabilidad $\{p(x)\}_{x \in \mathbf{X}}$. Es fácil observar que la entropía depende únicamente de la función de distribución, y como en nuestro caso todos los elementos tendrán la misma probabilidad, podemos hacernos una idea de el número medio de bits que tienen los elementos de $\mathbf{X}$.

Una vez conocido el concepto de entropía, que ha sido explicado con detalle en [18] y nos ha resultado de gran ayuda para familiarizarnos con el mismo, vamos a ver cómo cuantificar la incertidumbre que tiene un elemento conociendo un valor de otro suceso de antemano.

Definición 2.10 : Sean (X, Y) dos variables aleatorias. La entropía condicionada de X conociendo el valor de Y es la siguiente:

$$H(X/Y) = \sum_{y \in Y: p_Y(y) > 0} p_Y(y) \cdot H(X/Y = y)$$

donde

$$H(X/Y = y) = H(p_{X/Y=y}) = - \sum_{x \in X} p(x/y) \cdot \log_2 p(x/y)$$

que es la incertidumbre media sobre el valor $x \in X$ que se haya tomado al conocer previamente un valor $y \in Y$.

Por lo tanto, la entropía condicionada mide la cantidad de información que nos aporta conocer el valor de una variable sobre otra desconocida.

La definición de esquema perfecto dada nada más comenzar el capítulo establecía la cantidad de información que obtienen las agrupaciones de una estructura de acceso que realiza este tipo de esquemas. Esta descripción realizada en términos generales puede ser interpretada en términos de entropía, lo que nos aporta una idea más directa de cómo afecta la información de las agrupaciones al secreto.

Definición 2.11 Sea Γ una estructura de acceso, P un conjunto de participantes y $\mathbf{k}$ un conjunto de secretos los cuales tienen todos la misma probabilidad. Sea $P = \{P_1, \dots, P_n\}$ el conjunto de participantes, definiremos al conjunto de participaciones $S = (k, S_1, \dots, S_n)$ como un vector de variables aleatorias definidas en el mismo espacio probabilístico donde cada participación S_i se corresponderá con el participante P_i . Entonces, para cada $A \in \Gamma$ definimos $H(S_{i_1}, \dots, S_{i_{r_A}})$ como la entropía del conjunto de participantes de dicho grupo autorizado $(P_{i_1}, \dots, P_{i_{r_A}})$ y lo denotaremos por $S_A = (S_{i_1}, \dots, S_{i_{r_A}})$. Además diremos que $H(\mathbf{k})$ es la entropía del secreto. Como resultado de esta notación, definiremos $H(\mathbf{k}|S_{A_i})$ como la incertidumbre sobre el valor del secreto conociendo que elemento de A se ha tomado.

Para construir un **esquema perfecto** de reparto de secretos se deben de cumplir las siguientes condiciones:

- *Uniformidad del secreto*

$$H(k) = \log_2 |k| \geq 1$$

- *Reconstrucción del secreto en términos de entropía*

Sea $A \in (\Gamma, \Delta)$, diremos que A es un conjunto autorizado si:

$$H(k|S_A) = 0$$

Esto puede traducirse en que computando las participaciones del grupo autorizado se determina el secreto totalmente, o en términos de entropía con probabilidad 1.

- *Privacidad del secreto*

Sea $B \in (\Gamma, \Delta)$, diremos que B es un conjunto no autorizado si:

$$H(k|S_B) = H(k)$$

Lo que significa que el secreto y $S_B = S_{j_1} \times \dots \times S_{j_{r_B}}$ (las participaciones del conjunto no autorizado computadas) son independientes y que el conocimiento de éstas no aporta ningún tipo de información sobre el secreto.

Capítulo 3

Teoría de Códigos Correctores

En el año 1948 el matemático Shannon C.E. publica su trabajo, *A Mathematical Theory of Communication* [23]. Es el primer autor en tratar de dar un sentido matemático al concepto de información. En dicho artículo, presenta un esquema general que consta de una fuente, un transmisor y un receptor, un canal y un destino para la comunicación y pretende definir así cualquier sistema de comunicaciones. A este hecho se le considera el nacimiento de la Teoría de la Codificación y de los códigos correctores, ya que a partir de ahí numerosos autores comenzarían a elaborar distintos tipos de códigos y a resolver el problema de la decodificación. La construcción de la mayoría de códigos ya fiables comenzó en la década de los 50, por ejemplo Hamming creó los llamados códigos de Hamming [12], que permitían corregir y detectar errores de un bit. También en 1954, Reed [20] y Muller [17] crearon los códigos Reed-Muller que estaban basados en lógica mayoritaria, en 1955 se presentan los códigos convolucionales y en 1957, se introducen los códigos cíclicos. A finales de los años 50 Reed y Solomon elaboran un esquema de codificación para códigos bloque [21].

Durante los años 60 cabe destacar el trabajo de Bose y Chaudhuri [9] y Hocqenghem [13] que descubrieron de forma independiente unos códigos correctores de errores múltiples, que ahora denominamos códigos BCH y por supuesto la labor de Berlekamp [5] al dar algoritmos de decodificación algebraica para los códigos BCH.

Todos estos adelantos en el campo de los códigos correctores no fueron puestos en práctica hasta mediados de los años 70, cuando aparecieron las primeras máquinas con capacidad suficiente para llevar a cabo los cálculos necesarios. Hoy en día, podemos encontrar ejemplos de códigos correctores en cualquier sistema de comunicación como pueden ser las llamadas de teléfono digitales, los DVD's, las cuentas bancarias, los códigos de barras o el código ISBN de los libros.

Aún con la importancia que tienen la codificación y la corrección en numerosos ámbitos de la sociedad, el desconocimiento es bastante generalizado, por lo que para poder dar una generalización de esquemas de reparto en estos términos será necesario introducirlos en este capítulo. Además, podemos hacer notar que muchos de los conceptos que vamos a tratar han sido vistos en distintas asignaturas del Grado, principalmente en Matemática Discreta, por lo que no entraremos en ellos muy en detalle.

Una idea general de los códigos correctores de errores se basa en añadir a las palabras que queremos codificar un número de bits redundantes de chequeo que serán los encargados de detectar y corregir los errores que se producen en la transmisión. Existen numerosas técnicas para añadir dicha información y cada una de ellas dará lugar a códigos correctores distintos. En nuestro caso, nos centraremos en los códigos correctores de errores lineales.

Los códigos lineales se caracterizan por obtenerse a partir de una función de codificación lineal donde cualquier combinación lineal de palabras código da lugar a una nueva palabra código. Además, como cualquier código corrector, si se producen errores durante la transmisión del mensaje, los códigos lineales tienen la capacidad de que el receptor sea capaz de detectar o corregir algunos de éstos de forma directa. Veamos a continuación un poco de teoría de códigos lineales.

Definición 3.1 Sea $\mathbb{F}_q$ un cuerpo finito de q elementos. Diremos que un código $\mathcal{C}$ es lineal de longitud n siempre y cuando sea un subespacio vectorial de $\mathbb{F}_q^n$ de dimensión k y se denota como $[n, k]$ -código lineal.

Observaciones:

- Cualquier elemento de un código $\mathcal{C}$ puede expresarse como combinación lineal de elementos de una base de $\mathbb{F}_q^n$.
- Un código lineal $\mathcal{C}$ de longitud n y dimensión k sobre $\mathbb{F}_q$ tiene q^k elementos o palabras código.

3.1. Distancia mínima de un código y tasa de información

$\mathcal{C}$ es un subespacio vectorial sobre $\mathbb{F}_q$, y por lo tanto $|\mathcal{C}| \leq q^n$. Si además establecemos que, exactamente, $|\mathcal{C}| = M$, la longitud será como mínimo $\log_q(M)$. Esto se traduce en que en un código de longitud n , $\log_q(M)$ es el número de bits que contienen información determinante en cada palabra mientras que el resto de ellos contendrán solamente información redundante.

Como acabamos de explicar, el número de bits que contienen información correcta, $|\mathcal{C}|$ es diferente del número de bits que contiene cada palabra que se envía, n , donde se encuentran todos los que contienen información relevante y los redundantes. La tasa de información medirá la relación entre ambos, es decir, entre el número de bits del mensaje original $|\mathcal{C}|$ y el número de bits de la palabra codificada.

Definición 3.2 : Si $\mathcal{C}$ es un $[n, k]$ -código lineal, entonces su tasa de información será:

$$R = \frac{\log_q |\mathcal{C}|}{n} = \frac{\log_q q^k}{n} = \frac{k}{n}, 0 \leq R \leq 1$$

A continuación, definiremos lo que es la distancia mínima de un código, concepto que es necesario para poder determinar la capacidad de detección y corrección del mismo, es decir, el número de errores que es capaz de detectar y corregir respectivamente. Para ello, primero vamos a considerar la distancia entre elementos del espacio vectorial $\mathbb{F}_{q^n}$.

Definición 3.3 Sean $x, y \in \mathbb{F}_{q^n}$, se define la distancia de Hamming entre ambos elementos a la cantidad

$$d(x, y) = \# \{i/1 \leq i \leq n, x_i \neq y_i\}.$$

Además denominaremos peso de Hamming de $x \in \mathbb{F}_{q^n}$ como la distancia de un elemento al 0, es decir, $w(x) = d(\bar{x}, \bar{0})$.

Asimismo también se puede definir la distancia en un código lineal y cómo se verá más adelante su relación con la capacidad correctora. Dado un código $\mathcal{C}$ $[n, k]$ -lineal, se llama distancia del código a la menor distancia entre dos palabras del código cualesquiera:

$$d(\mathcal{C}) = \min \{d(x, y) / x, y \in \mathcal{C}\} \text{ con } x, y \in \mathcal{C}$$

Proposición 3.1 Si $\mathcal{C}$ es un (n, k) -código lineal, entonces su distancia mínima cumple que

$$d(\mathcal{C}) = \min \{w(v) / v \neq 0, v \in \mathcal{C}\}$$

Demostración: Al ser $\mathcal{C}$ un código lineal,

$$d(\mathcal{C}) = \min_{x \neq y \in \mathcal{C}} d(x, y) = \min_{x \neq y \in \mathcal{C}} w(x - y) = \min_{x \neq 0, x \in \mathcal{C}} w(x)$$

Una vez explicado lo que es la distancia mínima de un código quedan definidos completamente los tres parámetros básicos de un código lineal. Además, también está la tasa de información R , que es un parámetro secundario del código que mide la eficiencia de éste durante el envío de los mensajes. En lo que sigue, se denotará $\mathcal{C}$ como un $[n, k, d]$ -código lineal. En general, una vez quede establecida la longitud del código, buscaremos que su dimensión y su distancia mínima sean de tamaño grande, ya que esto nos permitirá transmitir un mayor número de mensajes y aumentar su capacidad de detección y corrección. Sin embargo al tratarse de parámetros inversamente proporcionales será necesario buscar un balance entre ambos.

Por otra parte, es lógico pensar que al intentar transmitir una palabra c del código, intentaremos que el número de errores que contenga sea lo más pequeño posible. Sin embargo, si el receptor que recibe c , recibe exactamente la misma palabra va a suponer que dicha palabra código no ha sufrido ninguna modificación, pero ¿que ocurrirá en el caso de haber recibido una palabra distinta? Razonablemente, una vez se hayan establecido los parámetros correctamente, podemos confiar en la capacidad de transmisión del código y suponer que no se han producido muchos errores. Por lo tanto, una primera forma de codificar es la siguiente:

Veamos a continuación, cómo se mide la capacidad de detección y corrección de un código. Sea $\mathcal{C}$ un código lineal. Entonces una palabra $y \in \mathcal{C}$ se decodifica en $x' \in \mathcal{C}$ si:

$$d(y, x') < d(y, z) \forall z \in \mathcal{C}, z \neq x'$$

Esta forma de decodificar un código se denomina **decodificación por distancia mínima**. Además, por el resultado anterior es lógico observar que esta decodificación solo funciona si la palabra más próxima es única. A continuación, estableceremos una serie de parámetros que determinan la capacidad de corrección y detección de un código.

Definición 3.4 : Si $\mathcal{C}$ es un código lineal, diremos que detecta s errores si al enviar $x \in \mathcal{C}$ se recibe y tal que $d(x, y) = s$, es decir si $r \leq s$ entonces $y \notin \mathcal{C}$.

Definición 3.5 : Si $\mathcal{C}$ es un código, se dice que $\mathcal{C}$ corrige t errores si puede corregir todos los errores de peso t o menor y se le denomina t -corrector si no corrige $t+1$ errores.

El siguiente resultado relaciona la distancia de un código con su capacidad detectora y correctora.

Proposición 3.2 : Sea $\mathcal{C}$ un código lineal con distancia mínima d . Entonces:

1. $\mathcal{C}$ es un código s -detector si, y sólo si $d = s + 1$.
2. $\mathcal{C}$ es un código t -corrector si, y sólo si $d = 2t + 1$.
3. $\mathcal{C}$ corrige r borrones si $r < d$.

Demostración:

1. Sea $B(c, r) = \{x \in \mathbb{F}_q^n / d(c, x) \leq r\}$ una bola de centro c y radio r . Entonces $\mathcal{C}$ detecta errores de peso menor o igual que r si y solo si dicha bola no contiene ninguna palabra codificada excepto c , lo que significa que solo detecta r errores o menos si $r < d$. Además si $d(c, c') = d$ con $c, c' \in \mathcal{C} \Rightarrow c + e = c'$ y $w(e) = d$ y c no detecta dicho error.
2. El motivo de tomar el valor $2t + 1$ para establecer la distancia mínima de $\mathcal{C}$ es que las bolas $B(c, t)$ con $t \in \mathcal{C}$ son disjuntas entre sí y por lo tanto, si se da esta condición $\mathcal{C}$ corregirá t errores. Además, si $c, c' \in \mathcal{C}$ y $c + e = c'$ con $w(e) = t + 1$, entonces $B(c, t + 1) \cap B(c', t + 1) \neq \emptyset$ y por lo tanto no se podría determinar cual era la palabra código inicial.
3. Si la palabra x contiene r borrones y $r < d$, podemos asignarles las posiciones $x_1, \dots, x_r$. Por tanto, existe una única palabra del código, c que coincide con x en $x_{r+1}, \dots, x_n$. Si existiera otra palabra $y \in \mathcal{C}$ que también coincidiera con x en estas posiciones, $d(c, y) \leq s < d$ y tendríamos que $c = y$.

3.2. Matrices generadoras de códigos

Si $\mathcal{C}$ es un código lineal con $f : \mathbb{F}_2^k \rightarrow \mathbb{F}_2^n$ su función de codificación, como $\mathcal{C} \subseteq \mathbb{F}_2^n$ y $\mathbb{F}_2^k$ es a su vez un subespacio vectorial de $\mathbb{F}_2^n$ de dimensión k , podemos establecer en dicho subespacio una serie de operaciones para definir el código $\mathcal{C}$ y llevar a cabo la labor de codificación. Esto será posible gracias a la construcción de las matrices generadora y de control que describiremos a continuación. En primer lugar, como $\mathcal{C}$ es un subespacio vectorial de $\mathbb{F}_2^n$ y se tiene una base del mismo en $\mathcal{C}$ y cualquier palabra $c \in \mathcal{C}$ será combinación lineal de los elementos de la base. Por esto, se puede definir lo que es una matriz generadora como la imagen de una aplicación inyectiva $f : \mathbb{F}_2^k \rightarrow \mathbb{F}_2^n$.

Definición 3.6 : *Denominaremos **matriz generadora** del código lineal $\mathcal{C}$ a aquella matriz G de dimensiones $k \times n$ asociada a la siguiente aplicación*

$$f : \mathbb{F}_2^k \longrightarrow \mathbb{F}_2^n \text{ inyectiva}$$

donde cada una de sus filas son una base de $\mathcal{C}$.

Por lo tanto, una matriz generadora define tanto un código como una codificación. Tomando cualquier elemento $\bar{a} \in \mathbb{F}_2^k$ se codifica por $\bar{a}G \in \mathbb{F}_2^n$ y se define el código como $\mathcal{C} = \{\bar{a}G / \bar{a} \in \mathbb{F}_2^k\}$. Además, es fácil de observar que al no poseer $\mathcal{C}$ una base única, la matriz generadora de $\mathcal{C}$ tampoco lo será, pero si podemos definir una matriz generadora estándar, $G = [I_k | A]$. Las palabras codificadas por dicha matriz tendrán la siguiente estructura:

$$\bar{x} = (x_1 \dots x_k x_{k+1} \dots x_n)$$

donde las k primeras coordenadas serán las que proporcionen información mientras que las $n - k$ siguientes serán únicamente símbolos de comprobación. En este caso, tanto el código como la codificación se denominan sistemáticos, algo que será explicado con mayor detalle a continuación:

Definición 3.7 : *Un $[n, k]$ -código lineal $\mathcal{C}$ se dice **sistemático** si existen k coordenadas $i_1, \dots, i_k$ de manera que si se reducen todas las palabras del código únicamente a estas coordenadas obtenemos todas las palabras de longitud k .*

Proposición 3.3 *Todo código lineal $\mathcal{C}'$ generado por otro $[n, k, d]$ -código lineal $\mathcal{C}$ es sistemático en las primeras k coordenadas.*

Demostración: *Basta con reducir cualquier matriz que genera el código lineal $\mathcal{C}'$ a una de la forma $[I_k | A]$ donde I_k es la matriz identidad de dim k y A una matriz cualquiera de dim $n - k \times k$.*

Una vez que hemos definido lo que es un código sistemático, vamos a explicar el significado de la equivalencia de códigos, concepto que nos resulta tremendamente útil, ya que podemos encontrar para cualquier código uno equivalente y sistemático.

Definición 3.8 Sean $\mathcal{C}, \mathcal{C}'$ dos $[n, k, d]$ -códigos lineales. Se dice que son equivalentes si existe una permutación $\sigma \in \{1, \dots, n\}$ de forma que:

$$\mathcal{C}' = \{ (a_{\sigma(1)} \dots a_{\sigma(n)}) / (a_1 \dots a_n) \in \mathcal{C} \}$$

Proposición 3.4 Para cualquier código $\mathcal{C}$ $[n, k, d]$ -lineal, existe un código equivalente $\mathcal{C}'$ que es sistemático en k coordenadas.

Demostración: Puede verse en [18].

Además de los códigos equivalentes, vamos a introducir también el concepto de código dual de un código $\mathcal{C}$, ya que van a ser imprescindibles en la labor de decodificación de cualquier palabra código, y posteriormente a la hora de recuperar el valor de un secreto s en los esquemas de reparto basados en códigos correctores.

Definición 3.9 : Sea $\mathcal{C}$ un $[n, k, d]$ -código lineal, diremos que el conjunto

$$\mathcal{C}^\perp = \{ \bar{u} / \bar{u} \cdot \bar{v} = 0 \forall \bar{v} \in \mathcal{C} \}$$

es un código $[n, n-k]$ -lineal denominado código dual de $\mathcal{C}$.

Los códigos duales son necesarios para la construcción de las matrices de control, que son las responsables de comprobar si una palabra codificada es o no una palabra código. Además, como comprobaremos a continuación hay una relación muy estrecha entre las matrices generadora y de control en un código lineal $\mathcal{C}$ y su dual.

Por otro lado, la matriz generadora G nos define las ecuaciones paramétricas de un código $\mathcal{C}$. Al ser un subespacio vectorial, $\mathcal{C}$ tendrá a su vez un conjunto de ecuaciones implícitas y dichas ecuaciones definen su matriz de control H .

Definición 3.10 : Sea H una matriz de dimensiones $(n - k) \times n$ y donde $\text{rg}(H) = n - k$. Se dice que H es una **matriz de control** si:

$$\forall \bar{x} \in \mathbb{F}_2^n, \bar{x} \in \mathcal{C} \iff H\bar{x}' = \bar{0}$$

Una vez definida lo que es una matriz de control H , si volvemos a la definición de código dual de $\mathcal{C}$ se puede afirmar que:

Definición 3.11 Siendo $\mathcal{C}^\perp$ un código dual de $\mathcal{C}$, éste será un código $[n, n - k]$ -lineal cuyas matrices generadora y de control serán H y G respectivamente.

Teorema 3.1 Sean G, H las matrices generadora y de control de un $[n, k, d]$ -código $\mathcal{C}$, entonces

1. $GH^t = 0$
2. Si G es una matriz generadora de $\mathcal{C}$, es a su vez la matriz de control de $\mathcal{C}^\perp$.

Demostracion:

1. Sean $x \in \mathcal{C}$, $\hat{x} \in \mathcal{C}^\perp$, entonces $x = Gy^\perp$, $\hat{x} = Hz^\perp$ con $y \in \mathbb{F}_q^k$, $z \in \mathbb{F}_q^{n-k}$. Por lo tanto, $x \cdot \hat{x} = 0$ si, y sólo si $Gy^\perp \cdot Hz^\perp = Hz^\perp \cdot Gy^\perp = zH \cdot Gy^\perp = 0$. por último, como $(HG)_{ij} = e_i(HG)e_j^\perp$, lo anterior solo se cumple si, y sólo si $HG = 0$.
2. Es trivial ya que H^t es la matriz de control de $\mathcal{C}$ y $H^t = (G^t)^t = G$.

Definición 3.12 Sea H la matriz de control de un código $\mathcal{C}$. Si $H = [-A^t | Id_k]$ se denomina matriz de control en forma estándar.

Como acabamos de ver, las matrices G y H quedan relacionadas por el teorema anterior. Pero si además, G se encuentra en forma estándar, el cálculo de H es trivial.

Proposición 3.5 Si $G = [Id_k | A]$ es la matriz generadora de un código lineal $\mathcal{C}$, entonces su matriz de control $H = [A^t | Id_{n-k}]$ es la matriz generadora de su código dual $\mathcal{C}^\perp$.

Demostracion: Tenemos que $GH^t = I_k(-A) + AI_{n-k} = 0$, por lo tanto las columnas de H son ortogonales a las columnas de G y por definición, $\text{span}(H) = \mathcal{C}^\perp$

Como vimos anteriormente, la matriz generadora además de dar una codificación de $\mathcal{C}$ también daba una definición del mismo. Lo mismo sucede con la matriz de control, que además de dar una decodificación de $\mathcal{C}$ define dicho código a partir de su propia definición.

Proposición 3.6 Si H es la matriz de control de un $[n, k]$ -código $\mathcal{C}$, podemos definir dicho código como:

$$\mathcal{C} = \{x \in \mathbb{F}_q^n / Hx^t = 0\}$$

Demostracion: Sea $y \in \mathbb{F}_q^n$, $y = Gx^t$ con $x \in \mathbb{F}_q^k$ y G la matriz generadora de $\mathcal{C}$. Entonces, $Hy^t = HGx^t = 0$ y así $\mathcal{C} \subseteq \{x \in \mathbb{F}_q^n / Hx^t = 0\}$. Además, como este conjunto es el subespacio formado por las soluciones de un sistema de $n - k$ ecuaciones con n incógnitas y rango $n - k$ y su dimensión es $\dim = (n - k) - n = k = \dim(\mathcal{C})$, entonces $\mathcal{C} = \{x \in \mathbb{F}_q^n / Hx^t = 0\}$

El siguiente resultado muestra cómo determinar la distancia mínima de un código $\mathcal{C}$ en función del número de columnas linealmente independientes de su matriz de control.

Teorema 3.2 Sea $\mathcal{C}$ un $[n, k]$ -código lineal y H su matriz de control. Entonces el código tiene distancia mínima d si todas las familias de $d - 1$ columnas de H son linealmente independientes y sin embargo alguna de las familias de d columnas es linealmente dependientes.

Demostracion: Sea H una matriz de control de $\mathcal{C}$ y $H_1, \dots, H_n$ las columnas de H y sea

$$d = \min \{r | r \text{ es el número de columnas linealmente independientes de } H\}$$

$$c \in \mathcal{C} \subseteq \mathbb{F}_q^n \Leftrightarrow Hc^\perp = 0 \Leftrightarrow H_1c_1 + \dots + H_nc_n = 0$$

Pero por otra parte, si $c \in \mathcal{C}$, c tiene peso r si, y sólo si, hay una familia de r columnas linealmente dependientes de H y como $r \geq d$ no puede haber $d - 1$ columnas linealmente independientes.

3.3. Códigos MDS

El último teorema introducido en la sección anterior se utiliza de manera general en la construcción de $[n, k, d]$ -códigos lineales donde la distancia mínima d quede fijada de antemano. A partir de este teorema podemos relacionar entre sí los tres parámetros de un código y construir un tipo de códigos denominados códigos MDS o códigos de máxima distancia.

Para construir estos códigos es imprescindible introducir la cota de Singleton, ya que para denominar a un código MDS, deberá alcanzar dicha cota, que definiremos a continuación.

Teorema 3.3 Sea $\mathcal{C}$ un código $[n, k, d]$ -lineal, se tiene entonces la **cota de Singleton** que establece que

$$n - k \geq d - 1$$

Demostracion: Por el teorema anterior, en la matriz de control H de $\mathcal{C}$ cualquier familia de $d - 1$ columnas es linealmente independiente. Además, como $H \in \mathcal{M}_{(n-k) \times n}$, entonces:

$$d - 1 \leq n - k \Leftrightarrow d \leq n - k + 1$$

Definición 3.13 se denominan códigos MDS o códigos de máxima separación, ya que la distancia mínima del código va a ser la mayor posible, a los códigos que alcanzan el igual en la cota de Singleton.

Además, es posible determinar que códigos son MDS mediante los siguientes resultados.

Teorema 3.4 Sea $\mathcal{C}$ un $[n, k, d]$ -código lineal, entonces son equivalentes:

- $\mathcal{C}$ es MDS.
- todas las familias de $n - k$ columnas de una matriz de control H de $\mathcal{C}$ son linealmente independientes.
- $\mathcal{C}^\perp$ es MDS.

Demostracion: Véase la siguiente referencia [16]

Capítulo 4

Ejemplos de esquemas para repartir secretos

Distintos autores han abarcado este tema desde diferentes puntos de vista utilizando múltiples herramientas matemáticas. Todos estos intentos por solucionar dicho problema dan lugar a distintos esquemas totalmente diferentes entre sí, algunos de los cuales estudiaremos a continuación.

4.1. Esquema de Shamir

El esquema de Shamir [22] para compartir secretos, que fue el primero de los numerosos esquemas para repartir secretos que podemos encontrar, se basa en la construcción de polinomios sobre un cuerpo finito $\mathbb{F}_q$. Dicho esquema, como veremos a continuación es un (t, n) -esquema umbral. Para que sea posible realizar esta construcción es necesario que el cardinal del cuerpo $\mathbb{F}_q$ sea mayor al número de participantes, y para ello por simplicidad asumimos que $\mathbb{F}_q = \mathbb{F}_p$, con p primo y $p > n$. La idea principal de este esquema, como aparece en [3] y [2], es la construcción de un polinomio aleatorio de grado fijado t , $f(x) \in \mathbb{F}[X]$ de forma que $f(0) = k$ (el secreto) y donde a cada participante se le repartirá el valor del polinomio en un punto. El esquema consiste en lo siguiente:

Sea $\Gamma = \{A \subseteq \{P_1, \dots, P_n\} : |A| \geq t \text{ con } t < n\}$ una estructura de acceso y sean t, n, p tres enteros no negativos. Se quiere repartir un secreto $k \in \mathbb{F}_p$. Para ello:

1. El gestor toma t elementos $a_0 = k, a_1, \dots, a_{t-1} \in \mathbb{F}_p$ de forma aleatoria.
2. Con estos elementos y el secreto se construye un polinomio de grado $t - 1$

$$f(x) = k + \sum_{i=1}^{t-1} a_i x^i$$

3. Se distribuye el secreto k en participaciones, que se reparten a cada participante de la siguiente forma.

A cada participante P_j se le reparte una participación $s_j = f(j)$ con $j \in \{1, \dots, t\}$ y $1 \leq j \leq n$

Una vez repartidas todas las participaciones s_j , se utiliza el Teorema de la Interpolación de Lagrange para recuperar el secreto.

Teorema 4.1 *Teorema de Interpolación de Lagrange*

Sea $\mathbb{F}$ un cuerpo. Dados dos conjuntos de puntos distintos $x_0, x_1, \dots, x_n$, $y_0, y_1, \dots, y_n$ hay un único polinomio $q(x)$ de grado a lo sumo n sobre $\mathbb{F}$ tal que:

$$q(x_i) = y_i \quad \forall i \in \{0, \dots, n\}$$

Se puede aplicar el teorema de Interpolación al esquema de Shamir reconstruyendo el polinomio $f(x)$ de nuestro esquema y a continuación calculando el secreto $k = f(0)$.

Sea $A = \{P_{i_1}, \dots, P_{i_t}\}$ una agrupación autorizada de t participantes. Entonces los participantes de A tendrán t puntos distintos $s_{i_j} = f(i_j)$ con $1 \leq j \leq t$ del polinomio $f(x)$. Podremos reconstruir $f(x)$ con estos puntos mediante el *Teorema de Interpolación de Lagrange* y una vez sea conocido el polinomio, recuperamos el secreto calculando $f(0) = k$.

Comenzamos construyendo un polinomio $h(x)$ con las participaciones de A de la siguiente manera:

$$h(x) = \sum_{l=0}^{t-1} s_{i_l} \delta_l(x)$$

donde $\delta_l(x)$ es un polinomio de grado $t - 1$ definido como:

$$\delta_l(x) = \prod_{0 \leq j \leq t-1, j \neq l} \frac{i_j - x}{i_j - i_l}$$

Es fácil observar que si $x = i_l$, $\delta_l(i_l) = 1$ y $\delta_j(i_l) = 0$ para todo $j \neq l$. Por lo tanto, tendremos que $h(i_l) = s_{i_l}$ para $1 \leq l \leq t$. Como además habíamos tomado los puntos i_j para calcular las participaciones de A de forma que $s_{i_j} = f(i_j)$ se tiene que $h(i_l) = s_{i_l} = f(i_l)$ y aplicando el Teorema de Lagrange, como ambos polinomios son de grado máximo $t - 1$ y coinciden en t puntos distintos, se tiene que $h(x) = f(x) \forall x$.

Una vez que el polinomio está construido veamos como se recupera el secreto k .

Como $k = f(0)$ y acabamos de ver que $h(x)$ y $f(x)$ son el mismo polinomio, bastará calcular $h(0)$ para obtener el valor del secreto k .

$$k = h(0) = \sum_{l=1}^t s_{i_l} \beta_l$$

siendo β_l la siguiente expresión:

$$\beta_l = \prod_{1 \leq j \leq t, j \neq l} \frac{i_j}{i_j - i_l} \text{ con } 1 \leq l \leq t$$

Si reagrupamos los términos como

$$k = \sum_{l=1}^t \beta_l s_{i_l}$$

es fácil observar que reconstruir el secreto k es equivalente a realizar una combinación lineal de las participaciones de A donde los β_l son todos valores conocidos por la agrupación autorizada y dependen únicamente de la misma. Por todo esto, dada una agrupación autorizada A , que tendrá t o más participantes, será capaz de obtener el valor de k a partir de sus participaciones.

Veamos ahora que ocurre cuando se tiene un conjunto no autorizado.

Como vamos a tomar una agrupación no autorizada, podemos elegir cualquier conjunto con a lo sumo $t - 1$ participantes. En nuestro caso, se toma $B = \{P_{i_1}, \dots, P_{i_{t-1}}\}$ y se da a cada participante P_{i_j} una participación $s_{i_j} = f(i_j)$. Por lo tanto se tienen $t - 1$ valores distintos del polinomio $f(x)$ de grado $t - 1$ y junto con todos los posibles valores del secreto, pueden determinar de forma única dicho polinomio.

Por todo esto, si tenemos los $t - 1$ valores de una agrupación no autorizada cualquiera de cardinal t , para cada posible valor $a \in \mathbb{F}_p$ del secreto, aplicando el Teorema de Interpolación se obtiene un único polinomio $f_a(x)$ de grado $t - 1$ con $f_a(0) = k$ y $f_a(i_j) = s_{i_j}$ con $1 \leq j \leq t - 1$. Lo que quiere decir que una agrupación no autorizada B de cardinal $t - 1$ obtiene q^{t-1} valores posibles para el valor del secreto.

A continuación, vamos a introducir un $(3, 4)$ -esquema umbral de Shamir.

Ejemplo 4.1 Sea $\mathcal{P} = \{P_1, P_2, P_3, P_4\}$ nuestro conjunto de participantes y sea

$$\Gamma_0 = \{\{P_1, P_2, P_3\}, \{P_1, P_3, P_4\}, \{P_1, P_2, P_4\}, \{P_2, P_3, P_4\}\}$$

la base de una estructura de acceso Γ .

Por lo tanto, $t = 3$ y $n = 4$. Además vamos a tomar $p = 23$, por lo que trabajaremos en el cuerpo finito $\mathbb{F}_{23}$.

Una vez introducidos todos los datos necesarios vamos a comenzar con el reparto de participaciones:

1. El gestor elige tres elementos de $\mathbb{F}_{23}$ al azar

$$a_0 = k = 4, a_1 = 18, a_2 = 19$$

2. Construimos un polinomio $f(x)$ de grado 2 con a_0, a_1, a_2

$$f(x) = 4 + 18x + 19x^2$$

3. Para distribuir k , el gestor reparte a cada P_i con $i \in \{1, 2, 3, 4\}$ su participación de la siguiente manera.

Siendo $x_1 = 1, x_2 = 2, x_3 = 3, x_4 = 4$, se calcula $f(x_i)$:

$$f(1) = s_1 = 18, f(2) = s_2 = 1, f(3) = s_3 = 22, f(4) = s_4 = 12,$$

y repartimos a cada P_i la participación s_i

Una vez repartidas las participaciones veamos como una agrupación autorizada puede recuperar el secreto.

Sea $\{P_1, P_2, P_3\} \in \Gamma_0$. Para recuperar el secreto k , calculamos el valor de $h(0)$ de la siguiente manera:

$$k = h(0) = \sum_{l=1}^t s_{i_l} \prod_{1 \leq j \leq t, j \neq l} \frac{x_{i_j}}{x_{i_j} - x_{i_l}} \text{ con } 1 \leq l \leq t$$

Sustituyendo con las participaciones de P_1, P_2, P_3 se obtiene

$$k' = 18 \cdot \frac{2}{1} \cdot \frac{3}{2} + 1 \cdot \frac{1}{22} \cdot \frac{3}{1} + 22 \cdot \frac{1}{21} \cdot \frac{2}{22} = 8 + 20 + 22 = 4 = k$$

Ahora que hemos visto como recupera el secreto una agrupación autorizada, veamos que sucede con las agrupaciones no autorizadas.

Sea $\{P_1, P_2\} \notin \Gamma$. Las participaciones correspondientes a P_1, P_2 son las siguientes:

$$x_1 = 1, x_2 = 2 \text{ y } s_1 = 18, s_2 = 1$$

Repitiendo los cálculos anteriores obtendrían

$$k' = 18 \cdot \frac{2}{1} + 1 \cdot \frac{1}{22} = 13 + 11 = 1 \neq 4 = k$$

Acabamos de observar que una agrupación no autorizada no obtiene el valor de k . Sin embargo, vamos a ver que resultado obtenemos con las participaciones de otra agrupación:

Tomando $\{P_1, P_3\} \notin \Gamma$, al juntar sus participaciones

$$x_1 = 1, x_2 = 3 \text{ y } s_1 = 18, s_2 = 22$$

obtendrían el siguiente valor

$$k' = 18 \cdot \frac{3}{2} + 22 \cdot \frac{1}{21} = 4 + 12 = 16 \neq 4 = k$$

Es fácil observar, que como hemos afirmado en este capítulo, una agrupación no autorizada va a obtener cualesquiera de los valores posibles para el secreto, en este caso cualquier elemento de $\mathbb{F}_{23}$. Además, al tener la misma probabilidad para cada uno de ellos, tampoco van a poder incrementar las posibilidades de obtener el verdadero valor de k ni de saber si el obtenido es correcto o no.

4.2. Construcción geométrica

El esquema umbral que se ha introducido en el apartado anterior fue desarrollado por Shamir [22] en 1979. A su vez, durante ese mismo año Blakley [7] también estudiaba el problema de reparto de secretos, y acabó describiendo un esquema umbral de manera totalmente diferente a Shamir. En este método Blakley realiza una construcción geométrica en un espacio afín de dimensión n sobre un cuerpo finito. En esta sección describiremos la construcción de Blakley, basándonos también en el esquema descrito por Blanco [8], cuya idea principal es tomar una recta pública y construir un hiperplano de manera que su intersección sea el secreto.

Sean $t, n \in \mathbb{N}$ con $t \leq n$ y $q = p^r$ potencia de un número primo, y sean $\mathbb{F}_q$ un cuerpo finito y $\mathbb{F}_q^t$ un espacio afín euclídeo.

Entonces el gestor D toma una recta pública V_D en el espacio afín $\mathbb{F}_q^t$ definida por un punto p_0 y el vector $e \in \mathbb{F}_q^t$, es decir:

$$V_D = \{p_0 + \mu e / \mu \in \mathbb{F}_q\}$$

D elige un punto $p \in V_D$ de forma aleatoria y privada que se corresponderá con el secreto k . Al encontrarse el secreto en la recta pública V_D , podemos escribir $p = p_0 + \mu_0 \cdot e$ con $\mu_0 \in \mathbb{F}_q$. Una vez dada la recta V_D y elegido el punto que se corresponde con el secreto k , pasaremos a contruir el hiperplano cuya intersección con la recta pública hará a una agrupación autorizada conocedora del secreto. Es fácil observar que, para cada vector u no ortogonal a la recta V_D podemos construir un hiperplano con distinta dirección y que corte a la recta en un único punto.

Para ello tomamos un vector $u \in \mathbb{F}_q^t$ cualquiera no ortogonal al vector director e de la recta V_D . Para que se corten en un único punto, basta con que el vector de la recta no esté en la dirección del plano. Al haber tomado un vector u para construir el hiperplano H_p no ortogonal al vector director de V_D , es decir que $u \notin \langle e \rangle^\perp$, la recta pública V_D ni está contenida en H_p ni son paralelos entre sí, por lo que dicho vector no está en la dirección del hiperplano, y entonces se cortan en un punto. Como además por construcción sabemos que $p \in V_D$ y $p \in H_p$, la intersección entre ambos será el punto p , que contiene el valor del secreto.

Entonces, para cada elemento a de $\mathbb{F}_q$, existe un hiperplano diferente de vector normal u , que podemos construir de la siguiente manera:

$$H_{u,a} = \{(x_1, \dots, x_t) = x / u \cdot x = a\}$$

es decir,

$$u_1 \cdot x_1 + u_2 \cdot x_2 + \dots + u_t \cdot x_t = a$$

cuya dirección es $H_u \equiv u \cdot x = 0$.

A continuación, para cada secreto p el gestor toma de forma privada un hiperplano $H_p = p + H_u$ que se corresponde con uno de los hiperplanos $H_{u,a}$ que pasa por el punto p y que por tanto tiene la misma dirección H_u y ecuación $ux = a$ con $a = up_0 + \mu_0 ue$.

Para repartir un secreto entre un conjunto de participantes de una agrupación $A = \{P_1, \dots, P_n\}$, el gestor D elige n puntos al azar $\{s_1, \dots, s_n\}$ del hiperplano H_p de foma que cada familia de t puntos sea linealmente independiente. Posteriormente reparte un punto s_i a cada participante P_i .

Es trivial afirmar, que el conjunto de puntos asociados a los participantes de una agrupación serán las participaciones del esquema y que a partir de éstas dichos participantes tratarán de recuperar el secreto p .

Ahora que ha quedado explicado como se realiza la contrucción geométrica vamos a ver que este esquema de reparto es un esquema umbral y como se establecen y se comportan las distintas agrupaciones que pretenden recuperar el secreto.

Proposición 4.1 *Este esquema es un (t, n) -esquema umbral.*

Demostracion: Veamos primero lo que ocurre en las agrupaciones con t o más participantes. Como hemos visto, a cada participante le corresponde un punto s_i del hiperplano H_p y por tanto, en una agrupación de al menos t participantes tendremos t puntos linealmente independientes. Entonces, la mínima variedad afín que los contiene tendrá dimensión $t - 1$, al igual que el hiperplano H_p , por lo que será posible construir dicho hiperplano con las participaciones de esta agrupación. Por último, sólo nos quedaría realizar la intersección con la recta pública V_D y calcular el valor del secreto.

Si por el contrario tenemos una agrupación con menos de t participantes, al unir sus participaciones se construye una variedad afín de dimensión t' siempre menor que $t - 1$. Por lo tanto, para cada punto $p' \in V_D$ hay al menos un hiperplano de dimensión $t - 1$ que contiene a p' y a la variedad afín construida con el conjunto de participaciones, por lo que cualquier punto de la recta podría ser el secreto.

Acabamos de probar que esta construcción describe un (t, n) -esquema umbral. Por tanto, podemos afirmar que todo conjunto con t o más participantes será una agrupación autorizada. A continuación, veamos como calcular secreto p al unir todas sus participaciones.

Sea $A = \{P_1, \dots, P_t\} \in \Gamma$. Por lo tanto si asociamos un punto s_i a cada participante P_i de A y ponemos en común los t puntos $s_1, \dots, s_t$ de la agrupación se construye un hiperplano de dimensión $t - 1$ de la siguiente manera.

Mediante el conjunto de puntos $\{s_1, s_2, \dots, s_t\}$ se puede construir una variedad

$$L \equiv s_1 + \langle \overrightarrow{s_2 s_1}, \overrightarrow{s_3 s_1}, \dots, \overrightarrow{s_t s_1} \rangle$$

donde renombramos $s_{j+1} - s_j = v_j$ con $j = 1, \dots, t - 1$ cuya ecuación paramétrica será

$$H_p \equiv s_1 + \lambda_1 v_1 + \dots + \lambda_{t-1} v_{t-1}$$

y eliminando los parámetros λ_i se obtiene la siguiente ecuación implícita:

$$A_1 x_1 + A_2 x_2 + \dots + A_t x_t = B$$

Una vez obtenidas las ecuaciones del hiperplano, al realizar la intersección con $V_D = p_0 + \mu_0 e$ se obtiene el secreto p al ser el único punto común de ambos. Para ello pasamos ambas variedades a sus ecuaciones paramétricas, las ponemos en común y resolvemos el sistema obtenido.

$$\begin{cases} x_1 &= s_{11} + \lambda_1 v_{11} + \lambda_2 v_{21} + \dots + \lambda_{t-1} v_{t-1,1} \\ x_2 &= s_{12} + \lambda_1 v_{12} + \lambda_2 v_{22} + \dots + \lambda_{t-1} v_{t-1,2} \\ \dots &\dots \dots \\ x_t &= s_{1t} + \lambda_1 v_{1t} + \lambda_2 v_{2t} + \dots + \lambda_{t-1} v_{t-1,t} \end{cases}$$

y las ecuaciones paramétricas de la recta que son:

$$\begin{cases} x_1 &= p_{01} + \lambda \mu_0 e_1 \\ x_2 &= p_{02} + \lambda \mu_0 e_2 \\ \dots &\dots \dots \\ x_t &= p_{0t} + \lambda \mu_0 e_t \end{cases}$$

Ejemplo 4.2 Vamos a realizar un ejemplo de la construcción geométrica basado en un $(3, 4)$ -esquema umbral.

Sean $\mathcal{P} = \{P_1, P_2, P_3, P_4\}$, $t = 3$, $n = 4$, $q = 5$.

Por lo que vamos a trabajar sobre el espacio afín euclídeo $\mathbb{F}_5^3$.
Veamos como reparte el gestor el secreto.

- Primero, se toman al azar los siguientes parámetros

$$\mu = 1, e = (2, 3, 1), p_0 = (2, 1, 0)$$

- Ahora, el gestor construye la recta pública $V_D = (2, 1, 0) + \langle (2, 3, 1) \rangle$ y calcula el valor de $k = p_0 + \mu * e = (2, 1, 0) + (2, 3, 1) = (4, 4, 1)$ correspondiente al secreto.
- Después de que se han elegido los parámetros anteriores, se construye un hiperplano privado $H_{a,u}$ tomando cualquier vector $u \in \mathbb{F}_5^3$ no ortogonal a e . Siendo $u = (3, 1, 0)$, se calcula

$$H_{a,(3,1,0)} = \{x / (3, 1, 0) \cdot x = a \text{ con } a \in \mathbb{F}_5\}$$

cuyas ecuaciones vectoriales son

$$H_{a,(3,1,0)} \equiv 3x_1 + x_2 = a$$

- A partir de $H_{a,u}$, se calcula otro hiperplano H_p con la misma dirección y parámetro

$$a = up_0 + \mu_0 ue = (3, 1, 0) (2, 1, 0) + (3, 1, 0) (2, 3, 1) = 1$$

cuyas ecuaciones vectoriales son

$$H_p \equiv 3x_1 + x_2 = 1$$

Una vez tenemos la recta pública y el hiperplano privado, el gestor reparte las participaciones necesarias para recuperar k entre los participantes.

- En nuestro caso, como se trata de un $(3, 4)$ -esquema umbral, se reparte un puntos del hiperplano H_p a cada participante de forma que los cuatro puntos sean linealmente independientes entre sí tres a tres.

$$P_1 : s_1 = (0, 1, 0) \quad P_2 : s_2 = (0, 1, 1)$$

$$P_3 : s_3 = (1, 3, 0) \quad P_4 : s_4 = (1, 3, 1)$$

- Sea $A = \{P_1, P_2, P_3\}$ una agrupación minimal. Al poner en común las participaciones de A , s_1, s_2, s_3 podemos construir la siguiente variedad

$$L = (0, 1, 0) + \langle (0, 0, 1), (1, 2, 0) \rangle$$

cuyas ecuaciones paramétricas son

$$\begin{cases} x_1 &= \mu_2 \\ x_2 &= \mu_1 + 2 \cdot \mu_2 \\ x_3 &= \mu_1 \end{cases}$$

y eliminando los parámetros obtenemos la siguiente ecuación vectorial

$$L \equiv 3x_1 + x_2 = 1$$

Es fácil observar que $L = H_p$, así que la agrupación autorizada ha sido capaz de construir el hiperplano privado H_p . Por lo tanto, sólo nos quedará intersectar dicho hiperplano con la recta V_D para obtener el valor de k .

Las ecuaciones paramétricas de V_D son

$$\begin{cases} x_1 &= 2 + 2\lambda \\ x_2 &= 1 + 3\lambda \\ x_3 &= \lambda \end{cases}$$

y al realizar la intersección con H_p obtenemos

$$3(2 + 2\lambda) + (1 + 3\lambda) = 1$$

y al despejar se tiene que $\lambda = 1$. Por último, al sustituir el valor de λ en las ecuaciones paramétricas de V_D hemos obtenido el punto $(x_1, x_2, x_3) = (4, 4, 1) = k$, y se observa cómo el esquema es correcto.

Sea ahora $B = \{P_3, P_4\}$ una agrupación no autorizada. Al poner en común sus participaciones $s_3 = (1, 3, 0)$ y $s_4 = (1, 3, 1)$ podemos construir la recta

$$r \equiv (1, 3, 0) + \langle 0, 0, 1 \rangle$$

cuyas ecuaciones paramétricas son

$$\begin{cases} x_1 &= 1 \\ x_2 &= 3 \\ x_3 &= \mu \end{cases}$$

Es fácil observar que al intersecar la recta anterior con la recta pública V_D de ecuaciones paramétricas

$$\begin{cases} 1 &= 2 + 2\lambda \\ 3 &= 1 + 3\lambda \\ \mu &= \lambda \end{cases}$$

al resolver el sistema obtenido vemos que es incompatible ya que tenemos dos valores diferentes para un mismo parámetro, $\lambda = 4$, $\lambda = 2$. Por lo que al tratarse de un sistema incompatible no existe ningún punto de corte entre ambas rectas y no será posible determinar el secreto.

4.3. Construcción de Ito, Saito y Nishizeki

En 1987, Ito, Saito y Nishizeki [15] mostraron una construcción que define un esquema para repartir secretos en términos generales para todo tipo de estructuras de acceso y en particular para estructuras de acceso monótonas, que es la que se muestra en esta sección.

Este tipo de construcción está basada en un circuito monótono que reconoce la estructura de acceso y establece un sistema para repartir las participaciones entre los participantes. Sin embargo, como veremos a continuación se trata de un esquema bastante ineficiente, ya que sólo es óptimo (en el sentido de la tasa de información) para estructuras de acceso que no tengan muchas agrupaciones minimales.

Antes de comenzar con la construcción, cuyas distintas versiones presentes en [2], [3] y [24], nos han servido desarrollar la nuestra, vamos a detenernos brevemente en definir lo que es un circuito monótono y cómo se opera en el, para poder manejar con mayor facilidad este tipo de esquemas.

Sea $B = \{0, 1\}$ el conjunto de dos elementos que representa los valores que puede tomar un bit. Se denomina **circuito monótono** a cualquier circuito booleano con n entradas, $x_1, \dots, x_n$ (que tomarán un valor de B) y una salida x dependiendo de éstos formado sólo por compuertas 'or' ($\vee$) y compuertas 'and' ($\wedge$) cuyas operaciones son:

$\vee$	0	1
0	0	1
1	1	1

$\wedge$	0	1
0	0	0
1	0	1

Además, se tiene en el conjunto B la operación siguiente:

$\oplus$	0	1
1	0	1
0	1	0

En resumen, en este esquema se pretende construir un circuito monótono que reconozca la estructura de acceso y una vez esté construido, distribuir las participaciones a partir del mismo. A cada participante se le va a repartir un bit por cada agrupación autorizada que lo contenga. Por lo que si tenemos una estructura de acceso donde haya un número elevado de agrupaciones autorizadas minimales, se repartirá mucha cantidad de información a cada uno de los participantes y en consecuencia, su tasa de información será muy baja.

Veamos cómo se realiza el esquema de reparto.

Sea Γ una estructura de acceso monótona y sea $A \in \Gamma$, entonces el gestor D "repartirá" el secreto $k \in \{0, 1\}$ de la siguiente manera:

Siendo A una agrupación autorizada cualquiera con $\{P_{i_1}, \dots, P_{i_l}\}$

- El gestor toma $l - 1$ bits $r_1, \dots, r_{l-1}$ de forma aleatoria.
- A continuación, calcula $r_l = k \oplus r_1 \oplus \dots \oplus r_{l-1}$
- Por último, asocia a cada participante P_{i_j} el bit r_j correspondiente a la participación s_{A_j} .

Es decir, cada participante recibirá tantas participaciones como el número de agrupaciones autorizadas a las que pertenece. Ahora que se ha establecido el esquema de reparto, es fácil observar que se trata de considerar para cada agrupación minimal A un (l_A, l_A) -esquema umbral donde $l_A = |A|$.

Para realizar esta construcción era necesario definir previamente un circuito monótono que se corresponderá con la base de una estructura de acceso a partir de la cual asociaremos a cada participante sus participaciones de una determinada manera. Para ello, a partir de la base de una estructura de acceso cualquiera

$$\Gamma_0 = \{\{P_{1_1}, \dots, P_{1_r}\}, \dots, \{P_{t_1}, \dots, P_{t_{r'}}\}\} \text{ con } r, r', t \leq n$$

podemos construir el siguiente circuito monótono:

$$\bigvee_{A \in \Gamma_0} \bigwedge_{P_i \in A} P_i$$

Rápidamente se observa que los elementos de la base quedan definidos por el circuito monótono que acabamos de construir. La unión de todas las agrupaciones minimales queda definida por la compuerta "or" de nuestro circuito monótono y los participantes de cada agrupación minimal se encuentran dentro de la compuerta "and", ya que es necesario la presencia de todos ellos para que se cumpla la condición de minimalidad.

Por lo tanto, a partir de la base de la estructura de acceso, es decir, de nuestro conjunto de agrupaciones minimales autorizadas, se construye un circuito monótono. En dicho circuito, una vez el gestor reparta las participaciones a cada participante, primero, éste, gracias a dicho circuito distinguirá que participación se corresponde con cada agrupación autorizada a la que pertenece y una vez hayan hecho todos los participantes esta distinción pondrán en común las participaciones correctas en cada caso para la recuperación del secreto. Para recuperar el secreto bastará con calcular el "o exclusivo" de todas las participaciones:

$$s_{A_{i_1}} \oplus s_{A_{i_2}} \oplus \dots \oplus s_{A_{i_j}} = k \text{ para todo } s_{A_{i_j}} \in A \in \Gamma$$

Si tenemos una agrupación no autorizada B , entonces no contiene a ninguna agrupación minimal y el gestor repartirá de forma aleatoria todos los bits correspondientes a estos participantes. Por lo tanto, al no seguir el esquema de reparto, tendremos una pérdida de al menos un bit de información, por lo que la probabilidad de obtener cualquiera de los dos valores del secreto k es la misma.

Ejemplo 4.3 Sea $\mathcal{P} = \{P_1, P_2, P_3, P_4, P_5\}$ y

$$\Gamma_0 = \{\{P_1, P_2\}, \{P_3, P_4\}, \{P_2, P_3, P_5\}\}$$

que se construye a partir del circuito monótono

$$(P_1 \wedge P_2) \vee (P_3 \wedge P_4) \vee (P_2 \wedge P_3 \wedge P_5)$$

por lo que asociamos a cada participante las siguientes participaciones

$$P_1 = \{s_{A_{1,1}}\} \quad P_2 = \{s_{A_{1,2}}, s_{A_{3,2}}\} \quad P_3 = \{s_{A_{2,3}}, s_{A_{3,3}}\} \quad P_4 = \{s_{A_{2,4}}\} \quad P_5 = \{s_{A_{3,5}}\}$$

El gestor toma de manera aleatoria $k = 0$ y les asocia los siguientes valores a las participaciones siguiendo el esquema de reparto:

$$s_{A_{1,1}} = 1, s_{A_{1,2}} = k \oplus 1 = 0, s_{A_{3,2}} = 1, s_{A_{2,3}} = 0 \oplus 1 = 0, s_{A_{3,3}} = 0, s_{A_{2,4}} = 0, s_{A_{3,5}} = k \oplus 0 \oplus 0 = 0$$

Una vez repartidas las participaciones, veamos cómo recuperan el secreto distintas agrupaciones autorizadas:

- Sea $A_1 = \{P_1, P_2\}$. Poniendo en común sus participaciones se obtiene:

$$s_{A_{1,1}} \oplus s_{A_{1,2}} = 1 \oplus 0 = 0 = k$$

- Sea $A_2 = \{P_1, P_2, P_4\}$ una agrupación autorizada. El gestor ha asociado a P_4 la participación

$$s_{A_{1,4}} = k \oplus s_1 \oplus s_2 = 0 \oplus 1 \oplus 0 = 1$$

y se calcula

$$s_{A_{1,1}} \oplus s_{A_{1,2}} \oplus s_{A_{1,4}} = 1 \oplus 0 \oplus 1 = 0 = k$$

- Sea $A_3 = \{P_1, P_2, P_3, P_4\} \in \Gamma$. Se tienen las participaciones $s_{A_{1,1}}$, $s_{A_{1,2}}$, $s_{A_{1,3}}$. Por último el gestor ha asociado a $P_{A_{1,3}}$ la participación

$$s_{A_{1,3}} = 0 \oplus 1 \oplus 0 = 1$$

$$s_{A_{1,1}} \oplus s_{A_{1,2}} \oplus s_{A_{1,3}} \oplus s_{A_{1,4}} = 1 \oplus 0 \oplus 1 \oplus 1 = 0 = k$$

- Veamos ahora que sucede si una agrupación es no autorizada.

Vamos a tomar la agrupación $B_1 = \{P_2, P_3\}$ no autorizada contenida en $A_3 = \{P_2, P_3, P_5\}$. Por lo tanto, tenemos las participaciones $s_{A_{3,2}}$ y $s_{A_{3,3}}$ y al ponerlas en común se obtiene

$$1 \oplus 0 = 0 = k$$

Sin embargo, si juntamos las participaciones de la agrupación $B_2 = \{P_3, P_5\}$ no autorizada también contenida en A_3

$$0 \oplus 0 = 1 \neq k$$

el valor calculado no se corresponde con el secreto.

Por último, si tenemos una agrupación no autorizada $B_3 = \{P_1, P_4, P_5\} \notin \Gamma$, no se les ha asociado ninguna participación correspondiente a esta agrupación a los participantes. Por lo tanto, las participaciones de P_1 y P_2 pueden tomar cualquier valor. Entonces,

$$s_{B_{3,1}} = 0, s_{B_{3,4}} = 1, s_{B_{3,5}} = 0$$

$$s_{B_{3,1}} \oplus s_{B_{3,4}} \oplus s_{B_{3,5}} = 0 \oplus 1 \oplus 0 = 1 \neq k$$

Sin embargo, si tomamos $s_{B_{3,1}} = 0, s_{B_{3,4}} = 1, s_{B_{3,5}} = 1$

$$s_{B_{3,1}} \oplus s_{B_{3,4}} \oplus s_{B_{3,5}} = 0 \oplus 1 \oplus 1 = 0 = k$$

Lo que nos indica que si una agrupación no es autorizada, va a obtener $k = 0$ o $k = 1$ dependiendo del valor que se de a las participaciones, que se asignan de manera aleatoria.

4.4. Construcción vectorial

La siguiente construcción fue descrita por Brickell [10] y da lugar a esquemas de reparto de secretos donde asociamos a cada participante y al gestor un vector de un espacio vectorial definido sobre un cuerpo finito. En estas construcciones, una agrupación será autorizada siempre y cuando el vector asociado al gestor sea combinación lineal de los vectores de los participantes de dicha agrupación.

Para poder realizar nuestra propia construcción vectorial hemos tenido en cuenta los resultados que aparecen en [3] y [19]. Además, previamente debemos introducir los siguientes conceptos.

Definición 4.1 Denominamos sistema generador monótono a una terna $(\mathbb{F}_q, M, \rho)$ siendo $\mathbb{F}_q$ un cuerpo finito, ρ una aplicación $\rho : \{1, \dots, a\} \rightarrow \{P_1, \dots, P_n\}$ y M una matriz $n \times a$ donde $a \in \mathbb{F}$ y con n igual al número de participantes.

Este sistema generador será necesario para describir la estructura de acceso que realiza el sistema de reparto, como veremos próximamente. La matriz M es la matriz correspondiente a todos los participantes de una estructura de acceso mientras que la aplicación ρ asocia a cada participante ciertas columnas de M , pudiendo repartir la misma columna a distintos participantes y sin ser necesario que todos ellos obtengan el mismo número de columnas.

Veamos a continuación cómo el sistema generador describe una estructura de acceso Γ .

Definición 4.2 Sea M una matriz $n \times a$ de un sistema generador y A un conjunto de participantes. Denominaremos M_A a la matriz de tamaño $n \times m$, obtenida al restringir la matriz M a las columnas asociadas a los participantes de A , que se corresponden con m .

Una vez conocida la matriz restringida de M a un grupo de participantes, podemos presentar el siguiente resultado que establece cómo es una agrupación de participantes a partir de su matriz correspondiente.

Definición 4.3 Diremos que una agrupación A es "válida" si las columnas de su matriz M_A generan el vector $e_1 = (1, 0, \dots, 0) \in \mathbb{F}_q^n$.

Aunque a priori pueda parecer que esta definición no aporta ningún tipo de información acerca del secreto, más adelante veremos que existe una equivalencia entre éstas agrupaciones y las agrupaciones autorizadas de una estructura de acceso. Por tanto, este resultado describirá cómo es nuestra estructura de acceso y determinará el valor del secreto.

A continuación, probaremos que los conjuntos "válidos" de un sistema generador se corresponden con los conjuntos autorizados de nuestra estructura de acceso, y que por lo tanto la terna antes descrita realiza dicha estructura.

Proposición 4.2 Sea $(\mathbb{F}_q, M, \rho)$ un sistema generador, A una agrupación válida, M_A su matriz restringida y sea Γ una estructura de acceso. Se dice que $(\mathbb{F}_q, M, \rho)$ describe la estructura de acceso Γ si se cumple la siguiente condición:

$$A \in \Gamma \Leftrightarrow \mathbf{e}_1 \in \{M_A c / c \in \mathbb{F}_q^{|A|}\}$$

donde $\mathbf{e}_1 = (1, 0, \dots, 0) \in \mathbb{F}_q^n$.

En otras palabras, la equivalencia anterior afirma que una agrupación A es autorizada si, y sólo si el vector $\mathbf{e}_1$ es combinación lineal de las columnas de M_A .

Veamos a continuación cómo se reparte el secreto entre los participantes y cómo recuperarlo:

En primer lugar, describiremos el posible sistema de reparto para probar después que cumple la condición a evaluar.

- Se tiene como entrada el secreto $k \in \mathbb{F}_q$ únicamente conocido por el gestor.
- El gestor toma un vector aleatorio $r \in \mathbb{F}_q^a$ de la siguiente manera:
se toman $n-1$ elementos $r_i \in \mathbb{F}_q$ con $i = \{2, \dots, n\}$ aleatorios y el gestor construye el vector $r = (k, r_2, \dots, r_n)$.
- El gestor calcula las participaciones $s = (s_1, \dots, s_a) = rM$.
- A continuación se distribuye a cada participante P_j las coordenadas j -ésima del vector s correspondientes a la columna a_j de la matriz M asociadas a dicho participante. Es decir, cada P_{j_i} recibe las coordenadas $\{s_i / P(i) = j\}$.

Por lo tanto, es fácil observar que del vector completo de las participaciones s cada participante sólo conocerá las coordenadas que le corresponden.

Ahora que tenemos definido nuestro posible esquema de reparto, vamos a tomar los conjuntos válidos del sistema generador inicial, y probaremos que dichos conjuntos se corresponden con las agrupaciones autorizadas de Γ .

Sea A un conjunto válido de cardinal m , M_A su matriz asociada y s_A las participaciones correspondientes a sus participantes. Como A es un conjunto válido, las columnas de M_A generan el vector e_1 , y por lo tanto existe un vector $c \in \mathbb{F}_q^m$ tal que $e_1 = M_A c \in \mathbb{F}_q^m$. Por otra parte, sabemos que las participaciones de A son de la forma $s_A = rM_A$.

Veamos si A puede recuperar el secreto k :

- Primero se calcula el vector c resolviendo el sistema $M_A \cdot c = e_1$ ya que podemos construir M_A a partir de las columnas asociadas a los participantes de A .
- A continuación, sacamos el valor del vector r resolviendo el sistema $s_A = rM_A$ ya que s_A es el vector de las participaciones de A .
- Una vez son conocidos c y r , podemos recuperar el valor de k calculando $s_A \cdot c$:

$$s_A \cdot c = (r \cdot M_A) c = r \underbrace{(M_A \cdot c)}_{e_1} = r \cdot e_1 = k$$

Por tanto, todo conjunto válido A es una agrupación autorizada. Así queda probado que nuestro sistema generador describe una estructura de acceso Γ al definir sus agrupaciones autorizadas. Una vez observado que todo conjunto válido es una agrupación autorizada, veamos que además dichas agrupaciones sólo pueden ser conjuntos válidos.

Para ello, probaremos que cualquier conjunto de nuestro sistema generador que no sea válido dará lugar a una agrupación no autorizada, y que por tanto no será capaz de recuperar ningún tipo de información a cerca del secreto.

Supongamos que B es un conjunto no válido. Entonces, las columnas de su matriz restringida M_B no pueden generar e_1 , por lo que es fácil observar que

$$\text{rg}(M_B) < \text{rg}(M_B \mid e_1) \Rightarrow |\text{Ker}(M_B)| > |\text{Ker}(M_B \mid e_1)|$$

y que además, existe un vector $w \in \mathbb{F}_q^n$ de forma que $w \cdot M_B = 0$ y $w \cdot e_1 \neq 0$ por lo que podemos tomar $w_1 = 1$.

Ahora, vamos a ver porqué estos conjuntos no válidos se corresponden con agrupaciones no autorizadas, y para ello nuestro conjunto B tratará de recuperar un secreto k a partir de sus participaciones.

Al fijar un vector $r = (0, r_2, \dots, r_n)$ que genere las participaciones para $k = 0$ a partir de la matriz restringida de B se calcula

$$r \cdot M_B = (s_{i_1}, \dots, s_{i_B})$$

quedando así generadas las participaciones para $k = 0$. Vamos a ver ahora como son las participaciones para cualquier valor de k .

Para un $k \in \mathbb{F}_q$ cualquiera, tomamos el vector $r' = r + k \cdot w$ generado a partir de r . Como $r'_1 = k$, tenemos que r' genera a su vez las participaciones de k . Calculemos ahora dichas participaciones:

$$r' \cdot M_B = (r + k \cdot w) \cdot M_B = r \cdot M_B + k \cdot \underbrace{w \cdot M_B}_0 = (s_{i_1}, \dots, s_{i_B})$$

Como acabamos de probar, para todos los posibles valores de k se genera el mismo vector de participaciones, lo que significa que a partir de las participaciones de dichos conjuntos, la probabilidad de obtener un valor determinado de k es la misma para cada valor. Por lo tanto no es posible obtener información del secreto a partir de B , y queda probado que todos los conjuntos no válidos son agrupaciones no autorizadas.

Con esto, vemos como al cumplirse la equivalencia anterior el sistema generador describe una estructura de acceso Γ definiendo sus agrupaciones autorizadas a partir de las columnas de su matriz de restricción. De esta forma, podemos decir que una construcción vectorial realiza un esquema perfecto en un estructura (Γ, Δ) .

Ejemplo 4.4 Sea $\mathcal{P} = \{P_1, P_2, P_3, P_4\}$ un conjunto de participantes, $(\mathbb{F}_7, M, \rho)$ un sistema generador con $\rho : \{1, 2, 3, 4, 5\} \rightarrow \{P_1, P_2, P_3, P_4\}$

$$M = \begin{pmatrix} 1 & 0 & 3 & 1 & 0 \\ 3 & 0 & 2 & 5 & 1 \\ 0 & 6 & 0 & 0 & 2 \\ 4 & 1 & 4 & 2 & 0 \end{pmatrix}$$

y $s = (s_1, s_2, s_3, s_4, s_5)$ el vector de las participaciones.

Asociando las columnas de M a los participantes mediante ρ tenemos:

$$\rho(1) = \rho(2) = P_1, \rho(4) = \rho(3) = P_2$$

$$\rho(5) = P_3, \rho(3) = P_4$$

Por lo que a cada P_i con $i \in \{1, 2, 3, 4\}$ le corresponden las siguientes participaciones:

$$P_1 = \{s_1, s_2\}, P_2 = \{s_4, s_5\}, P_3 = \{s_5\}, P_4 = \{s_3\}$$

Ahora que las coordenadas del vector de participaciones han quedado asignadas, el sistema de reparto para calcular el valor de ellas es el siguiente.

- El gestor toma $k = 1$ de manera aleatoria.
- A continuación, se construye un vector $r = (k, r_2, r_3, r_4)$ tomando tres elementos aleatorios de $\mathbb{F}_7$:

$$r = (1, 0, 1, 2)$$

- Seguidamente se calcula el vector de participaciones $s = (s_1, s_2, s_3, s_4, s_5)$, $s = rM$:

$$s = (1, 0, 1, 2) \begin{pmatrix} 1 & 0 & 3 & 1 & 0 \\ 3 & 0 & 2 & 5 & 1 \\ 0 & 6 & 0 & 0 & 2 \\ 4 & 1 & 4 & 2 & 0 \end{pmatrix} = (2, 1, 4, 5, 2)$$

- Se reparten a cada participante los valores correspondientes a las coordenadas que se le han asignado:

$$P_1 = \{2, 1\}, P_2 = \{4, 5\}, P_3 = \{2\}, P_4 = \{5\}$$

Ahora que cada participante tiene sus participaciones asociadas, vamos a ver como una agrupación autorizada es capaz de calcular el secreto.

- Sea $A \in \{P_1, P_3, P_4\}$. Podemos afirmar que es una agrupación autorizada ya que su rango es máximo:

$$M_A = \begin{pmatrix} 1 & 0 & 1 & 0 \\ 3 & 0 & 5 & 1 \\ 0 & 6 & 0 & 2 \\ 4 & 1 & 2 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 5 & 0 \\ 0 & 0 & 5 & 2 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

entonces, $\text{rg}(A) = 4$ y A es una agrupación autorizada pues generan $(1, 0, 0, 0)$.

- Como a los participantes se le han asociado la primera, segunda, cuarta y quinta columna de M , el vector de participaciones de A será

$$s_A = (s_1, s_2, s_4, s_5) = (2, 1, 5, 2)$$

- A continuación se calcula el vector c resolviendo el sistema $M_A \cdot c = e_1$

$$\begin{pmatrix} 1 & 0 & 1 & 0 \\ 3 & 0 & 5 & 1 \\ 0 & 6 & 0 & 2 \\ 4 & 1 & 2 & 0 \end{pmatrix} \begin{pmatrix} c_1 & c_2 & c_3 & c_4 \end{pmatrix} = (1, 0, 0, 0)$$

Por lo que $(c_1, c_2, c_3, c_4) = (6, 0, 2, 0)$.

- Se calcula r resolviendo el sistema $s_A = r \cdot M_A$

$$(2, 1, 5, 2) = (k, r_2, r_3, r_4) \begin{pmatrix} 1 & 0 & 1 & 0 \\ 3 & 0 & 5 & 1 \\ 0 & 6 & 0 & 2 \\ 4 & 1 & 2 & 0 \end{pmatrix}$$

y se obtiene $r = (1, 0, 1, 2)$.

- Por último, se calcula el valor de k al resolver $r \cdot M_A \cdot c = r \cdot e_1$

$$(1, 0, 1, 2) \cdot (1, 0, 0, 0) = 1 = k$$

Veamos ahora que sucede cuando una agrupación no es autorizada.

Sea $B = \{P_2, P_3, P_4\}$ una agrupación de participantes. Siguiendo el ejemplo anterior tenemos el vector de participaciones $s_B = (s_3, s_4, s_5) = (4, 5, 2)$ y la matriz asociada al código

$$M_B = \begin{pmatrix} 3 & 1 & 0 \\ 2 & 5 & 1 \\ 0 & 0 & 2 \\ 4 & 2 & 0 \end{pmatrix}$$

Ahora, la agrupación tratará de calcular c y para ello debe de resolver el sistema $M_B \cdot c = e_1$

$$\begin{pmatrix} 3 & 1 & 0 \\ 2 & 5 & 1 \\ 0 & 0 & 2 \\ 4 & 2 & 0 \end{pmatrix} \begin{pmatrix} c_1 & c_2 & c_3 \end{pmatrix} = (1, 0, 0, 0)$$

Al resolver este sistema se obtendría

$$\left(\begin{array}{ccc|c} 3 & 1 & 0 & 1 \\ 2 & 5 & 1 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 1 & 2 & 0 \end{array} \right) \rightarrow \left(\begin{array}{ccc|c} 3 & 1 & 0 & 1 \\ 0 & 1 & 4 & 2 \\ 0 & 0 & 2 & 2 \\ 0 & 0 & 0 & 2 \end{array} \right)$$

que, como podemos ver, es un sistema incompatible, por lo que no existe ningún vector c que genere e_1 . El vector e_1 era necesario para la recuperación del secreto, ya que hay que resolver $r \cdot M_A \cdot c = r \cdot e_1 = k$ por lo que nuestra agrupación no autorizada no será capaz de calcular su valor al no poder resolver $M_A \cdot c$.

4.5. Construcción mediante grafos

En el intento de construir distintos tipos de esquemas para repartir secretos, muchos autores, siendo Benaloh y Rudich [4] los principales y a quienes debemos la siguiente construcción, han optado por realizar esquemas de reparto mediante grafos para describir estructuras de acceso en las que las participaciones correspondientes a cada participante consten de sólo un par de vértices.

En esta sección, estudiaremos una construcción que realiza un esquema de reparto de secretos a partir de grafos completos no dirigidos, [3] donde diremos que un conjunto de aristas es una agrupación autorizada, si a partir de ellas se puede construir un camino entre el primer v_1 y el último vértice v_n . Además, es fácil percatarse de que si una agrupación autorizada es un conjunto de aristas determinado, los participantes del esquema se corresponderán con las distintas aristas del grafo completo K_n , es decir, que $P_{ij} = (v_i, v_j)$. Además, podemos establecer que el número de participantes es $\binom{n}{2}$.

Antes de comenzar con el reparto de participaciones, veamos que operación se utiliza para calcular su valor:

$\oplus$	0	1
0	0	1
1	1	0

A continuación, presentamos el esquema de reparto:

Sea $k \in \{0, 1\}$ el secreto.

- El gestor D debe tomar n bits $r_1, \dots, r_n$.
- Para ello, fija de antemano los valores del primer y del último bit tomando $r_1 = k$ y $r_n = 0$. A continuación elige de forma aleatoria los $n - 2$ bits restantes $r_2, \dots, r_{n-1}$.
- Se calculan las participaciones $s_{i,j} = r_i \oplus r_j$.
- Por último, reparte a cada participante (v_i, v_j) la participación $s_{i,j}$ que le ha sido asociada.

Una vez repartidas las participaciones, veamos que este esquema de reparto es correcto. Para ello, empezaremos probando que todo conjunto de aristas que realice un camino de v_1 a v_n es una agrupación autorizada.

Supongamos entonces que tenemos un conjunto de aristas A , donde dichas aristas definen un camino de v_1 a v_n .

Por lo tanto, vamos a considerar el siguiente conjunto de vértices

$$v_1 = v_{i_1}, v_{i_2}, \dots, v_{i_{l-1}}, v_{i_l} = v_n$$

que definirá un camino de v_1 a v_n y donde cada participante de A puede definirse como $P_{ij} = (v_i, v_j)$. Además, a cada participante P_{ij} se le asocia la participación $s_{ij} = r_i \oplus r_j$.

Teniendo estas pautas previas en cuenta podemos poner en común éstas participaciones y calcular el valor de k :

$$(r_1 \oplus r_2) \oplus (r_2 \oplus r_3) \oplus \dots \oplus (r_{l-2} \oplus r_{l-1}) \oplus (r_{l-1} \oplus r_l) = r_1 \oplus r_l = r_1 \oplus r_n = k \oplus 0 = k$$

Entonces, queda probado que la construcción del esquema es autorizada en el sentido usual, es decir, que cualquier conjunto de aristas que definen un camino de v_1 a v_n es capaz de recuperar el secreto. Veamos ahora que en el caso de tener un conjunto no autorizado, éste no puede obtener ninguna información a cerca del mismo.

Como vamos a tomar un conjunto no autorizado B , no existe ningún camino de v_1 a v_n en su conjunto de aristas y podemos definir un conjunto de vértices V_1 de forma que $v_i \in V_1$ si existe un camino entre v_i y v_1 .

Es trivial ver que $v_1 \in V_1$ y $v_n \notin V_1$ y que además si una arista pertenece a B , sus vértices v_i, v_j estarán ambos en V_1 o no.

Teniendo esto en cuenta, vamos a comprobar que al poner en común todas las participaciones de una agrupación no autorizada, el número de vectores que generan el conjunto de participaciones $\{s_{i,j}\}_{(i,j) \in B}$ es el mismo para $k = 0$ y para $k = 1$.

1. Comenzaremos con $k = 0$.

Primero fijamos un vector de bits aleatorios $(r_1, r_2, \dots, r_{n-1}, r_n)$ con $r_1 = k = 0$ y $r_n = 0$ que genere las participaciones necesarias para calcular k .

Una vez construido el vector de bits aleatorios necesario para calcular $\{s_{i,j}\}_{(i,j) \in B}$ para $k = 0$, vamos a definir el vector que genera las participaciones para $k = 1$.

2. Ahora vamos a tomar un vector $(r'_1, \dots, r'_n)$ de forma que:

$$r'_i = \bar{r}_i \text{ si } v_i \in V_1 \text{ o}$$

$$r'_i = r_i \text{ si } v_i \notin V_1$$

es trivial ver que $r'_1 = 1$ y $r'_n = 0$ y que además $k' = 1$.

Como $\{r'_1, \dots, r'_n\}$ es el conjunto de bits que genera las participaciones $s_{i,j} = r_i \oplus r_j$ debemos considerar dos casos:

- Los vértices del participante (v_i, v_j) pertenecen a V_1 . Entonces:

$$r'_i \oplus r'_j = \bar{r}_i \oplus \bar{r}_j = r_i \oplus r_j = s_{i,j}$$

- Los vértices del participante (v_i, v_j) no pertenecen a V_1 .

$$r'_i \oplus r'_j = r_i \oplus r_j = s_{i,j}$$

Entonces, el número de vectores que generan el conjunto de participaciones $\{s_{i,j}\}_{(i,j) \in B}$ es el mismo para $k = 0$ y para $k = 1$, y por lo tanto, al poner una agrupación no autorizada en común sus participaciones tendrán las mismas probabilidades de obtener ambos valores para un sólo valor de k .

Podemos concluir esta sección afirmando que la construcción de Benaloh y Rudich realiza una estructura de acceso a partir de grafos completos tomando como agrupaciones autorizadas aquellas que contienen un camino de v_1 a v_n . Por esto, la presencia de ambos vértices en la arista de alguno de sus participantes, será una condición necesaria la realización del esquema de reparto.

Ejemplo 4.5 *Tenemos un grafo completo K_5 y $\mathcal{P} = \{P_{1,2}, P_{1,3}, P_{1,4}, \dots, P_{4,5}\}$ un grupo de diez participantes. El gestor toma como secreto $k = 0$ y las participaciones de cada uno de los participantes serán una arista de dicho grafo. Veamos como una agrupación autorizada recupera el valor de k .*

Vamos a establecer primero el esquema de reparto:

- *El gestor elige 3 bits aleatorios $r_2 = 0$, $r_3 = 0$, $r_4 = 1$ y se toma $r_1 = k = 0$, $r_5 = 0$ de antemano.*
- *Además, reparte, entre otras, las siguientes participaciones:*

$$P_{1,2} : (v_1, v_2) \text{ , } s_{1,2} = r_1 \oplus r_2 = 0 \oplus 0 = 0$$

$$P_{2,4} : (v_2, v_4) \text{ , } s_{2,4} = r_2 \oplus r_4 = 0 \oplus 1 = 1$$

$$P_{4,5} : (v_4, v_5) \text{ , } s_{4,5} = r_4 \oplus r_5 = 1 \oplus 0 = 1$$

$$P_{4,3} : (v_4, v_3) \text{ , } s_{4,3} = r_4 \oplus r_3 = 1 \oplus 0 = 1$$

$$P_{1,3} : (v_1, v_3) \text{ , } s_{1,3} = r_1 \oplus r_3 = 0 \oplus 0 = 0$$

- *Sea $A = \{P_{1,2}, P_{2,4}, P_{4,5}\} \in \Gamma$. Como A es una agrupación autorizada, existe un camino de v_1 a v_5 . Por lo tanto, al poner los participantes de A sus participaciones en común se obtiene $k = 0$:*

$$s_{1,2} \oplus s_{2,4} \oplus s_{4,5} = 0 \oplus 1 \oplus 1 = 0$$

Veamos que sucede cuando tenemos agrupaciones no autorizadas.

Sean $B_1 = \{P_{1,2}, P_{2,4}, P_{4,3}\}$ y $B_2 = \{P_{2,4}, P_{1,3}\}$ dos agrupaciones no autorizadas. Al poner en común sus participaciones respectivamente se obtiene

$$s_{1,2} \oplus s_{2,4} \oplus s_{4,3} = 0 \oplus 1 \oplus 1 = 0 = k$$

$$s_{2,4} \oplus s_{1,3} = 1 \oplus 0 = 1 \neq k$$

Por lo que observamos, que para una agrupación no autorizada podemos obtener indistintamente ambos valores de k .

4.6. Construcción basada en códigos

Como ya sabemos, para resolver el problema de compartir un secreto entre un grupo de personas se han utilizado diferentes tipos de herramientas y en este apartado vamos a introducir una construcción basada en la teoría de Códigos Correctores.

No vamos a incidir en las propiedades y los conceptos básicos de Códigos Correctores ya que han sido explicados con detenimiento en el capítulo 3 y por lo tanto, como cabía esperar, vamos a utilizar muchos de esos resultados para poder definir nuestra construcción con éxito.

A continuación, veamos como se realiza una construcción basada en códigos para repartir secretos en el caso ideal, es decir, donde cada uno de sus participantes recibe sólo una participación.

Sean $\mathcal{P}$ el conjunto de los n participantes, Γ una estructura de acceso con Γ_0 una base de la estructura de acceso Γ , F_q^m cuerpo que contiene al conjunto de secretos, $\mathcal{C}$ un código $[n, k, d]$ -lineal con $k \geq m$ y $G \in \mathcal{M}_{k \times n}$ una matriz generadora del código lineal $\mathcal{C}$, que será de conocimiento público.

Para realizar una construcción general, asociamos una o más columnas de G a cada uno de los participantes. El número de columnas asociadas varían en función del participante (en nuestro caso, tomaremos una sólo columna por participante) y las columnas no se asocian de forma exclusiva, es decir, que podremos asociar la misma columna a distintos participantes.

Por otra parte, para asociar las columnas a los participantes construimos ciertos conjuntos de índices que indiquen qué columna o columnas le corresponden a cada quién, J_D y J_i , dónde J_D será el conjunto de índices correspondiente al gestor y J_i el conjunto de índices asociados a cada participante P_i .

En el caso de tener una agrupación autorizada, debemos de tomar las columnas asociadas al gestor de forma que sean linealmente dependientes a las columnas asociadas a sus participantes. Si por el contrario la agrupación a la que queremos repartir la información es no autorizada, las columnas de sus participantes deberán de ser linealmente independientes a las columnas del gestor.

Notemos que esta es la idea que subyace en la construcción vectorial, sin embargo, no debe de causarnos ningún tipo de confusión, ya que como veremos en adelante no se trata del mismo tipo de esquema aunque partan de una idea común.

Antes de presentar el esquema de reparto, tengamos en cuenta un par de observaciones.

- Si A es una agrupación de participantes, denotamos su conjunto de índices por $J_A = \cup_{P_i \in A} J_i$.
- A lo largo de los diferentes tipos de esquema, hemos denotado al secreto como k . Sin embargo, durante esta sección vamos a establecer un cambio de notación, denominando s al secreto y k a la dimensión del código.

Algoritmo para repartir el secreto s

- El gestor D toma un secreto $s \in \mathcal{K}$, $s = (s_1, \dots, s_m)$.
- Se amplía s hasta obtener un vector con k coordenadas, de forma que $s' = (\bar{s}, \bar{a})$ con $a \in \mathbb{F}^{k-m}$ aleatorio.
- Se codifica el vector s' mediante la matriz generatriz de conocimiento público para obtener una palabra del código

$$(\bar{s}, \bar{a}) G = (c_1, \dots, c_n) = \bar{c} \in \mathcal{C}$$

- Se reparten las coordenadas de $\bar{c}$ correspondientes a cada uno de los participantes y al gestor mediante sus conjuntos de índices:

$$c_i \text{ será una participación de } P_i \Leftrightarrow i \in J_i \text{ con } i = 1, \dots, n$$

$$c_i \text{ será una participación de } D \Leftrightarrow i \in J_D \text{ con } i = 1, \dots, n$$

Por tanto, las participaciones repartidas a cada participante serán las coordenadas de $\bar{c}$ que ocupen un lugar correspondiente a su conjunto de índices. Ahora que hemos visto cómo el gestor reparte un secreto s entre los distintos participantes, veamos cómo puede recuperarse.

Nota: Para que la tarea de recuperación nos resulte más sencilla, se pueden tomar las columnas de G que se asocian a cada participante de forma que sean linealmente independiente entre ellas o construyendo la matriz G de la siguiente forma:

$$G = \begin{pmatrix} Id_m & G' \\ 0 & G'' \end{pmatrix}$$

Recordamos la noción de código sistemático del capítulo anterior.

Sea A una agrupación autorizada. Para la posterior recuperación de un secreto s , A construye una matriz generatriz con las columnas de G correspondientes a los conjuntos de índices J_D y J_A de forma que decodifica la palabra $\bar{c}$ con las componentes que le han sido asociadas a sus participantes y tomando las coordenadas desconocidas como errores de tipo borrón.

Por lo tanto, es necesario que nuestro código lineal $\mathcal{C}$ tenga la capacidad de corrección suficiente para que una agrupación autorizada pueda decodificar $\bar{c}_A$ a partir únicamente de las coordenadas que conoce. Para esto, es importante recordar que un código tiene la capacidad de corregir $d - 1$ borrones.

Si tenemos en cuenta la nota anterior, al codificar

$$(\bar{s}, \bar{a}) G = \bar{c} \text{ se obtiene}$$

$$\bar{c} = (s_1, \dots, s_m, c_{m+1}, \dots, c_n)$$

donde, como podemos observar, las m primeras coordenadas son las correspondientes al vector s .

Proposición 4.3 *Con las notaciones anteriores, para cada secreto $s \in \mathbb{F}_q^m$ se tiene el conjunto de reglas de distribución*

$$J_s = \{\bar{c} \in \mathcal{C} / \bar{c}(J_D) = s\}$$

donde $\bar{c}(J_D)$ son las coordenadas de la palabra código $\bar{c}$ con los índices correspondientes al gestor J_D y

$$s_i = \{\bar{c}(J_i) / \bar{c} \in \mathcal{C}\}$$

el conjunto de participaciones para cada participante P_i con $i = 1, \dots, n$.

Entonces, el esquema anterior es perfecto realiza una estructura de acceso Γ .

Demostración: *Al haber definido un conjunto de reglas de distribución de forma que $J_s \subseteq J = \mathcal{C}$, bastará demostrar que dicho conjunto satisface las condiciones del teorema 2.3 y así quedará demostrado que $\mathcal{C}$ realiza la estructura de acceso Γ .*

1. *Si tenemos un secreto $s \in \mathbb{F}_q^m$, al codificar la palabra completa $(s, \bar{a})$ se obtiene $(s, \bar{a})G \in J_s \subseteq \mathcal{J}$ y si además existe una palabra código $\bar{c}$ para la que existe una regla de distribución de forma que*

$$\bar{c}(J_D) = (s, \bar{a})G(J_D)$$

entonces $\bar{c} \in J_s$.

Veamos ahora cómo las participaciones de una agrupación autorizada establecen el secreto de forma única.

Supongamos que $A \in \Gamma$ y que existen dos palabras código $\bar{c}_1, \bar{c}_2$ con $\bar{c}_1 \neq \bar{c}_2$ tales que $\bar{c}_1(J_A) = \bar{c}_2(J_A)$. Además como $G(J_D)$ es combinación lineal de las columnas $G(J_A)$, entonces para cada componente i -ésima de $G(J_D)$ existen $a_j \in \mathbb{F}_q$ de forma que

$$G(J_D)_i = \sum a_j G(J_A)_j \text{ entonces,}$$

$$\bar{c}_1(J_D)_i = (s, \bar{a}_1)G(J_D)_i = \sum a_j (s, \bar{a}_1)G(J_A)_j$$

$$= \sum a_j (s, \bar{a}_2)G(J_A)_j = (s, \bar{a}_2)G(J_D)_i = \bar{c}_2(J_D)_i \forall i$$

.

Por lo tanto, si dos palabras código generan las mismas participaciones, $\bar{c}_1(J_D) = \bar{c}_2(J_D)$ se tiene que $s_1 = s_2$.

2. *Por último, comprobaremos que se trata de un esquema privado al probar que se satisface la segunda condición del teorema.*

Como $B \notin \Gamma$, tomamos cualquier secreto $s \in \mathbb{F}_q^m$ y fijamos una palabra código $c_1 \in \mathcal{C}$. Como la familia de columnas asociadas al gestor $G(J_D)$ es linealmente independiente a las columnas asociadas a B , $G(J_B)$ aplicando el Lema 1.7 de [2] se obtiene:

$$\begin{aligned} & |\{\bar{c} \in \mathcal{C}_k : \bar{c}(J_B) = \bar{c}_1(J_B)\}| \\ &= |\{\bar{c} \in \mathcal{C} : \bar{c}(J_B) = \bar{c}_1(J_B) \text{ y } \bar{c}(J_D) = ((\bar{k}, \bar{a})G)(J_D)\}| \\ &= q^{k - (rg(G(J_B)) + |J_D|)} > 0 \end{aligned}$$

Acabamos de dar una serie de reglas de distribución a partir de la cuales nuestro código $\mathcal{C}$ puede describir una estructura de acceso. El siguiente resultado aportará una condición suficiente y necesaria que determina si una agrupación es o no autorizada en función del rango de su matriz asociada.

Teorema 4.2 *Sea Γ una estructura de acceso y A y B dos agrupaciones de participantes. Se tiene que*

- $A \in \Gamma \Leftrightarrow rg(G(J_A \cup J_D)) = rg(G(J_A))$
- $B \notin \Gamma \Leftrightarrow rg(G(J_B \cup J_D)) = rg(G(J_B)) + rg(G(J_D))$

Demostracion: *Puede verse en [1]*

Ahora que sabemos qué condición debe de cumplir una agrupación para ser autorizada, vamos a describir un subcódigo de $\mathcal{C}$ generado a partir de sus participantes. Esta construcción será de especial importancia, ya que si tenemos una agrupación que es autorizada recuperará el secreto a partir de dicho subcódigo.

Definición 4.4 *Sea $\mathcal{C}$ un código lineal que realiza una estructura de acceso Γ , G su matriz generadora, $\mathcal{P}$ un conjunto de participantes y A una agrupación cualquiera. Diremos que un subcódigo de $\mathcal{C}$ es un subcódigo generado por A si su matriz generadora es $G(J_A \cup J_D)$ cuyas columnas serán todas las columnas linealmente independientes de los participantes de A y todas las columnas asociadas al gestor. A este subcódigo se le denota como $\mathcal{C}_A$*

En nuestro caso, vamos a contemplar únicamente la opción en la que a cada P_i del conjunto de participantes se le asocie una única columna de la matriz G , y por lo tanto le corresponderá también una única coordenada del vector de participaciones. Además, al gestor obtiene también una sola columna de G y por lo tanto el valor del secreto se encontrará en la primera coordenada del vector de participaciones. A este tipo de esquema se le denomina esquema ideal y la idea principal es la siguiente.

Dado una agrupación autorizada A , al unir sus participaciones no obtenemos palabras con todas las coordenadas completas, ya que solo se conocen las componentes asociadas a los participantes de A . Sin embargo, aunque a parte de la coordenada del gestor correspondiente al secreto existan otras coordenadas desconocidas, gracias a una

serie de resultados que presentaremos a continuación se puede probar que es posible construir un subcódigo C_A que tiene distancia 2. Lo que nos va a permitir corregir una componente de nuestro vector de participaciones, que será la correspondiente al valor de s considerada un borrón, sin que sea necesario conocer el valor del resto de componentes desconocidas.

Proposición 4.4 *Sea C un código lineal que realiza una estructura de acceso Γ . Si $A = \{P_1, \dots, P_r\} \in \Gamma_0$, entonces el subcódigo generado por A es un subcódigo MDS de longitud $|A| + 1$, dimensión $|A|$ y distancia 2.*

Demostración: *Las columnas asociadas a los participantes de A en G_A son linealmente independientes por definición de C_A y entonces, las columnas de $G(J_A \cup J_D)$ también. Por lo tanto, cualquier familia de $|A|$ columnas de $G(J_A \cup J_D)$ es linealmente independiente y la dimensión de C_A es $|A|$.*

Además de esto, por la independencia de las columnas de $G(J_A \cup J_D)$ aplicando el Teorema 3.2 se obtiene que la distancia mínima del código $C_A^\perp$ es $|A| + 1$. Por lo tanto, dicho código es un código MDS de parámetros $(|A| + 1, 1, |A| + 1)$, y entonces aplicando el Teorema 3.4 el código C_A es otro código MDS con distancia mínima 2.

Como acabamos de ver, si A es una agrupación autorizada el subcódigo C_A asociado a ella tendrá distancia 2 y por lo tanto, corregirá un error, que se corresponde con la coordenada asociada al gestor. Gracias al resultado siguiente, podemos ver que sucede si una agrupación no es autorizada.

Proposición 4.5 *Sea C un código lineal que realiza una estructura de acceso Γ y sea $B = \{P_1, \dots, P_m\}$ una agrupación no autorizada. Entonces, el subcódigo C_B asociado a B es un código lineal de misma longitud y dimensión $r + 1$ con r el número de columnas linealmente independientes de B .*

Demostración: *Como B es una agrupación no autorizada, las columnas que son linealmente independientes entre sí de B también lo son con la columna asociada al gestor. Por lo tanto, si teníamos r columnas linealmente independientes en B , el código C_B tendrá longitud $r + 1$ y también la misma dimensión. Como tiene la misma longitud y la misma dimensión, el subcódigo genera todo el espacio vectorial y su distancia mínima es 1, por lo que no va a ser capaz de corregir ningún error.*

Para terminar, vamos a realizar un esquema de reparto basado en códigos correctores en el caso ideal.

Ejemplo 4.6 *Sea*

$$\Gamma_0 = \{\{P_1, P_2\}, \{P_4, P_5\}, \{P_2, P_3, P_4\}, \{P_1, P_3, P_5\}\}$$

la base de una estructura de acceso para 5 participantes y sea $\mathbb{F}_2$ nuestro conjunto de secretos. Además, tenemos la matriz generadora G

$$G = \begin{pmatrix} 1 & 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 1 & 1 & 0 \end{pmatrix}$$

del código $\mathcal{C}$. Asociamos las columnas de G a los distintos participantes de $\mathcal{P}$ de la siguiente manera; al gestor la primera columna y a los participantes las siguientes $P_1, P_2, \dots, P_5$ las columnas c_2, c_6, c_3, c_4 y c_5 respectivamente.

A continuación, se construye un vector aleatorio $r = (k, r_1, r_2)$ donde k es el secreto y r_1, r_2 dos bits tomados de forma aleatoria que se codifica mediante G , de la forma rG , para obtener todas las palabras del código a partir de las columnas asociadas a los participantes:

$$\begin{array}{ll} \text{Gestor} & a_0 \\ P_1 & a_0 + a_1 \\ P_3 & a_2 \\ P_4 & a_0 + a_1 + a_2 \\ P_5 & a_1 + a_2 \\ P_2 & a_1 \end{array}$$

Las posibles palabras del código son

$$\begin{aligned} & (0, 0, 0, 0, 0, 0), (1, 1, 0, 1, 0, 0), (0, 1, 0, 1, 1, 1), (0, 0, 1, 1, 1, 0) \\ & (1, 0, 0, 0, 1, 1), (1, 1, 1, 0, 1, 0), (0, 1, 1, 0, 0, 1), (1, 0, 1, 1, 0, 1) \end{aligned}$$

Veamos como se recupera el secreto a partir de las agrupaciones minimales.

Sea $\{P_1, P_2\}$ una agrupación autorizada. Entonces, dicha agrupación puede generar un subcódigo $\mathcal{C}_A$ cuya matriz generadora es

$$\begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{pmatrix}$$

y sus posibles palabras del código son

$$(0, 0, 0), (1, 1, 0), (0, 1, 1), (1, 0, 1)$$

A partir de la matriz generadora y las posibles palabras del código, es fácil observar que se trata de un código lineal de parámetros $[3, 3, 2]$. Además, como $d=2$, nuestro subcódigo $\mathcal{C}_A$ es capaz de corregir un error tipo borrón, por lo que será posible recuperar el valor de s .

Por ejemplo, si tomamos los valores $s = 0$, $a_1 = 1$ y $a_2 = 0$ y reunimos las participaciones de P_1 y P_2 obtenemos la palabra código $(s, 1, 1)$ que sólo puede asociarse a la palabra $(0, 1, 1)$ del código y por lo tanto se tiene que $s = 0$.

Otro caso puede ser la agrupación $B = \{P_2, P_3, P_4\} \in \Gamma_0$, que genera un subcódigo $\mathcal{C}_B$ cuya matriz generadora y posibles palabras del código son respectivamente

$$\begin{pmatrix} 1 & 0 & 1 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 1 & 0 \end{pmatrix}$$

y

$$(0, 0, 0, 0), (1, 0, 1, 0), (0, 0, 1, 1), (0, 1, 1, 0), (1, 0, 0, 1), (1, 1, 0, 0), (0, 1, 0, 1), (1, 1, 1, 1)$$

De la misma forma que con la agrupación anterior, podemos afirmar que nuestro subcódigo tiene parámetros $[4, 3, 2]$ y que por lo tanto, corrige un error de tipo borrón y será capaz de calcular el valor de s . Tomando los mismo valores que antes, $s = 0, a_1 = 1$ y $a_2 = 0$, al poner en común las agrupaciones de los participantes de B se construye la palabra código $(s, 0, 1, 1)$ que solo se puede corresponder con $(0, 0, 1, 1)$ y de donde se obtiene $s = 0$.

Si por el contrario se tiene una agrupación no autorizada como puede ser $C = \{P_2, P_3\}$ se puede construir un subcódigo $\mathcal{C}_C$ cuya matriz generadora es

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}$$

sus posibles palabras son

$$(0, 0, 0), (1, 0, 0), (0, 0, 1), (0, 1, 0), (1, 0, 1), (1, 1, 0), (0, 1, 1), (1, 1, 1)$$

y los parámetro asociados a el, $[3, 3, 1]$. Por lo tanto, al tener $d = 1$ no será capaz de corregir ningún borrón, lo que se traduce en que no podrá recuperar s .

Tomando nuevamente $s = 0, a_1 = 1$ y $a_2 = 0$, obtenemos el vector de participaciones $(s, 0, 1)$ que puede asociarse a la vez con $(0, 0, 1)$ y $(1, 0, 1)$ por lo que tenemos dos posibles valores para s .

Bibliografía

- [1] Abascal, P., Tena, J., Algoritmos de búsqueda de un Código Corrector de Errores realizando una Estructura de Acceso para Compartir Secretos. Actas de la V Reunión Española sobre Criptografía y Seguridad de la Información. 1998. pp, 279-283.
- [2] Abascal, P., Compartir secretos mediante esquemas basados en códigos correctores. Tesis Doctoral. Universidad de Oviedo. 1999.
- [3] Beimel, A. Secret-Sharing Schemes: A Survey. Department of Computer Science, Ben-Gurion University, Beer-Sheva, Israel. 2011, pp, 3-13.
- [4] Benaloh, J.C., Rudich, S. Private communication. 1989.
- [5] Berlekamp, E., On Decoding Binary Bose-Chaudhuri-Hocquenghem Codes. IEEE Trans. on Information Theory, vol. IT-11. 1965.
- [6] Bertilsson, M., Linear Codes and Secret Sharing. Linköping Studies in Science and Technology, Dissertation No. 299. 1993.
- [7] Blakley, G.R., Safeguarding cryptographic keys. AFIPS Conference Proceedings. 1979, 48, pp, 313-317.
- [8] Blanco Martín, M.F., Construcción Afín de un Esquema de Secreto Compartido. Universidad de Valladolid. 1997.
- [9] Bose, R.C., Ray-Chaudhuri, D.K., On a Class of Errorcorrecting Binary Group Codes. Inform. Control, vol. 3. 1960.
- [10] Brickell, E.F., Some ideal secret sharing schemes. Journal of Combin. Math. and Combin. Comput. 1989.
- [11] Brickell, E.F., Stinson, D.R., Some Improved Bounds on the Information Rate of Perfect Sharing Schemes. Lecture Notes in Computer Science. vol. 576. Berlin. 1992. pp, 242-252.
- [12] Hamming, R.W., Error Detecting and Error Correcting Codes. The Bell System Technical Journal. vol. 26. No. 2. 1950.
- [13] Hocquenghem, A., Codes Correcteurs d'Erreurs. Chiffres. vol. 2. 1959.

- [14] Deng, R. H., Feng, T., Information Security Practice and Experience. 9th International Conference, ISPEC 2013. Lanzhou. 2013.
- [15] Ito, M., Saito, A., Nishizeki, T. Secret sharing schemes realizing general access structures. Department of Electrical Communications. Tohoku University, Sendai. Japan.
- [16] McWilliams, F.J., Sloane, N.J.A., The Theory of Error-Correcting Codes. North Holland. 1988.
- [17] Muller, D.E., Application of Boolean Algebra to Switching Circuit Design and to Error Correction. IRE Trans. Electron. Computers, vol. EC-3. 1954.
- [18] Munuera, C., Tena, J., Codificación de la Información. Servicio de Publicaciones Universidad de Valladolid. 1997.
- [19] van de Pol, J., Smart, N., FHE-MPC Notes Lecture 7. 2011. pp, 2-3.
- [20] Reed, I., A class of Multiple-error-correcting Codes and the Decoding Scheme. IRE Trans. Inform. Theory. vol. 4, 1954.
- [21] Reed, I., Solomon, G., Polynomial Codes Over Certain Finite Fields. SIAM Journal of Applied Math. vol. 8, 1960.
- [22] Shamir, A., How to share a secret. Comm ACM. 1979.
- [23] Shannon, C.E., A mathematical Theory of Communication. Bell. Syst. J. vol. 27, 1948. pp, 379-423.
- [24] Stinson, D.R., An Explication of Secret Sharing Schemes. University of Nebraska, Lincoln. 1992. pp, 362-365.