

UNIVERSIDAD DE MÁLAGA
ESCUELA TÉCNICA SUPERIOR DE
INGENIEROS DE TELECOMUNICACIÓN



PROYECTO FIN DE CARRERA:

**DESARROLLO DE UN TUTORIAL CON
TECNOLOGÍAS WEB PARA LA ASIGNATURA DE
ELECTRÓNICA DE DISPOSITIVOS**

INGENIERÍA TELECOMUNICACIÓN

Málaga, 2001

Enrique Gómez García

UNIVERSIDAD DE MÁLAGA
ESCUELA TÉCNICA SUPERIOR DE INGENIEROS DE TELECOMUNICACIÓN

Titulación: Ingeniería de Telecomunicación

Reunido el tribunal examinador en el día de la fecha, constituido por:

D. _____

D. _____

D. _____

para juzgar el Proyecto Fin de Carrera titulado:

**DESARROLLO DE UN TUTORIAL CON
TECNOLOGÍAS WEB PARA LA ASIGNATURA DE
ELECTRÓNICA DE DISPOSITIVOS.**

del alumno D. Enrique Gómez García

dirigido por D. Eduardo Casilari Pérez

ACORDÓ POR: _____ OTORGAR LA CALIFICACIÓN DE

y, para que conste, se extiende firmada por los componentes de Tribunal, la presente diligencia.

Málaga, a ____de_____de 2001

El Presidente:

El Secretario:

El Vocal:

Fdo. _____

Fdo. _____

Fdo. _____

UNIVERSIDAD DE MÁLAGA
ESCUELA TÉCNICA SUPERIOR DE INGENIEROS DE TELECOMUNICACIÓN

DESARROLLO DE UN TUTORIAL CON TECNOLOGÍAS WEB PARA LA ASIGNATURA DE ELECTRÓNICA DE DISPOSITIVOS

REALIZADO POR:

Enrique Gómez García

DIRIGIDO POR:

Eduardo Casilari Pérez

DEPARTAMENTO DE: Tecnología Electrónica.

TITULACIÓN: Ingeniería de Telecomunicación

Palabras claves: Electrónica de Dispositivos, Tutorial, *applets* de Java, Internet

RESUMEN:

En este proyecto se desarrolla un tutorial para la asignatura de Electrónica de Dispositivos. Al utilizar los applets de Java podemos realizar una serie de simulaciones interactivas que ilustrarán los procesos físicos que tienen lugar en los semiconductores. Dentro del programa de la asignatura se intentan cubrir los aspectos más importantes entre los que se encuentran los semiconductores, diodos, transistores bipolares y dispositivos unipolares como el JFET y el MOS.

Málaga, Septiembre de 2001

ÍNDICE

Índice	3
CAPÍTULO 1: Introducción	1
1.1. LA PÁGINA PRINCIPAL	7
CAPÍTULO 2: Estructura de los applets	11
2.1. JAVA	11
2.2. Diseño e Implementación de los applets	19
CAPÍTULO 3: Descripción de los applets y manual de usuario	25
3.1. EVOLUCIÓN DEL NIVEL DE FERMI	25
3.1.1. Descripción del applet	35
3.2. Semiconductor homopolar en equilibrio con distribución no uniforme de impurezas	38
3.2.1. Descripción del applet	41
3.3. Unión PN en equilibrio y polarizada	43
3.3.1. Descripción del applet	46
3.4. LA LEY DE SHOCKLEY	55
3.4.1. Descripción del applet	57
3.5. EL DIODO VARACTOR	61
3.5.1. Descripción del applet	62
3.6. Problema del diodo varactor	66
3.6.1. Descripción del applet	68
3.7. DISTRIBUCIÓN DE CORRIENTES Y PORTADORES EN EL BJT	70
3.7.1. Descripción del applet	72
3.8. CONMUTACIÓN DEL TRANSISTOR	75
3.8.1. Descripción del applet	77
3.9. CANAL EN UNA ESTRUCTURA JFET	84
3.9.1. Descripción del applet	85
3.10. LA CAPACIDAD MOS	87
3.10.1. Descripción del applet	96
CAPÍTULO 4: Conclusiones y líneas futuras	101
BIBLIOGRAFÍA	105

CAPÍTULO 1: INTRODUCCIÓN

Internet ha incidido mucho en la ampliación de las posibilidades educativas. Las nuevas tecnologías *Web* aplicadas al campo de la pedagogía hacen posible una mayor integración dentro de la red.

En la asignatura de Electrónica de Dispositivos surge la necesidad de un procedimiento que ilustre los procesos físicos que tienen lugar en los semiconductores. El aprendizaje acerca de los materiales y dispositivos de estado sólido es una tarea ardua. Por eso disponer de una herramienta de aprendizaje dinámica y visual que ilustre los conceptos físicos abstractos puede ser una fuente útil para el proceso de aprendizaje. Los conceptos físicos que sirven de base para los materiales y dispositivos son inherentemente abstractos, y una comprensión adecuada es importante para el posterior aprendizaje de los principios de operación de los dispositivos y de la tecnología de fabricación de estos y de los circuitos.

Hasta el momento contábamos con una serie de herramientas de simulación asistida por computador que ayudaban al alumno en su proceso de investigación. Gracias a la utilización de *applets* podemos llevar a cabo una serie de simulaciones visuales interactivas de la física de los dispositivos semiconductores, que permitirán al alumno asentar sus conocimientos.

La simulación con *applets* tendrá un impacto arrollador en el proceso de aprendizaje de los estudiantes. Vamos a tratar algunos aspectos relacionados con el desarrollo de los *applets* de Java educativos, la utilización de los mismos, y la posible reorganización que se puede llevar a cabo. También es de gran importancia tener un buen material de soporte como tutoriales de introducción y ejemplos para trabajar con los *applets*.

La tecnología basada en *applets* de Java es muy conveniente para crear y difundir por la red Internet programas de simulación pequeños e interactivos con fines docentes. Java es un lenguaje de programación para Internet, y los *applets* están integrados naturalmente en el entorno de documentos con HTML (*HyperText Markup Language*). Los *applets* se pueden ejecutar directamente dentro de un explorador como Netscape 2.0, Internet Explorer 3.0 o versiones superiores.

Los esfuerzos para desarrollar herramientas multimedia para la enseñanza están encaminados para acelerar el aprendizaje del estudiante a través del uso de programas de software interactivo, ayudas visuales como imágenes estáticas y pequeños vídeos, y

documentos con hiperenlaces. Estas herramientas para la ingeniería pueden descargarse de Internet. Podemos encontrar varios cursos multimedia que proporcionan mejoras sustanciales en el aprendizaje del estudiante. Estos beneficios incluyen una considerable ganancia en tiempo y en la retención del estudiante.

En los últimos años, con la creciente difusión de la *World Wide Web* (WWW), unida al uso genérico de los exploradores que utilizan los documentos HTML por defecto, y la introducción en 1996 de un lenguaje orientado a objetos para Internet, Java, de Sun Microsystems inc., han proporcionado una nueva oportunidad para el desarrollo, reparto y distribución de utilidades interactivas docentes por la WWW.

La tecnología de los *applets* de Java es especialmente adecuada para las aplicaciones educativas por varias razones:

- (i) Permite distribuir el contenido del curso independientemente de la plataforma utilizada y de forma automática por la WWW. Esta independencia de la plataforma significa que códigos idénticos, tanto el código fuente como los compilados se pueden ejecutar en sistemas Windows, Unix, etc. Estos *applets* de Java, al descargarse de Internet, podrán verse en cualquier sistema si su explorador tiene habilitado el Java.
- (ii) La disponibilidad de los materiales del curso para los estudiantes desde su casa o desde una estación de trabajo en la Universidad sin necesidad de instalar ningún software.
- (iii) La posibilidad de realizar una clase práctica si se dispone de un aula con ordenadores.
- (iv) La existencia de extensas librerías de Java para trabajar con gráficos en el JDK (*Java Development Kit*) facilita en gran medida la programación de estas aplicaciones gráficas.
- (v) Fácil programación estructurada para animaciones gráficas.
- (vi) Integración natural de los *applets* con los documentos con HTML. Estos documentos son los más adecuados para proporcionar el contexto adecuado a los materiales del curso que se distribuyen.

Los *applets*, dependiendo de cómo se desarrollen, pueden contener diversos contenidos pedagógicos y de ingeniería. Como se ha comentado, estos programas se pueden utilizar tanto en una clase práctica como tarea propuesta para el alumno gracias a su propiedad de disponibilidad universal (independiente de la plataforma y accesible desde el WWW). Si se

quiere utilizar como un tutorial independiente donde los alumnos utilizan el programa como la herramienta principal de aprendizaje, es necesario acompañarlo de un texto para conducir al lector por el programa de una manera significativamente didáctica. Algunas lecturas complementarias e imágenes estáticas se pueden utilizar para introducir conceptos básicos, para completar el *applet* o para proporcionar los conocimientos previos necesarios. Un tutorial completo se compone de una página *Web* con hiperenlaces con los *applets* embebidos, con ecuaciones, imágenes estáticas, datos y todos los elementos multimedia necesarios. Este tutorial se podrá repartir de manera directa por la WWW desde un servidor que podrá ser accedido desde cualquier explorador con Java habilitado.

Los *applets* que se han desarrollado ilustran, de una manera interactiva y con gráficos animados, diversos conceptos de materiales y dispositivos semiconductores.

El software reutilizable, en forma de *Java applets*, proporciona una plataforma abierta, distribuida y expansible para un curso interactivo. Los *applets* poseen unas excelentes características como software reutilizable con propósitos docentes. Pueden ser fácilmente incluidos en una página *html* conjuntamente con otros elementos multimedia tales como imágenes, videos y sonidos, permitiendo una sencilla configuración del material multimedia dinámico destinado al aprendizaje. Esto se debe al hecho de que los *applets* se ejecutan embebidos en una página *Web*.

Una ventaja importante de los *Java applets* como software destinado a la docencia estriba en la facilidad para su reutilización por parte del personal docente y del alumnado. De tal manera que es extremadamente importante la facilidad para componer un curso por Internet y la sencillez de uso por parte del alumno para conseguir un amplio impacto. Así la sobrecarga incluida en el curso, tal como aprender a incluir los *applets* en él o cómo utilizar el propio curso, es prácticamente nula. Cualquier educador que sepa realizar una página *html* puede incluir fácilmente un *applet* en ésta. Los alumnos simplemente tendrán que visitar la página del Tutorial con cualquier explorador para comenzar inmediatamente su aprendizaje gracias a los *applets*. Éste es uno de los mayores activos que los *applets* proporcionan a las aplicaciones educativas. Otras ventajas son la independencia de la plataforma y la naturaleza distribuida del entorno de ejecución. Una herramienta para crear páginas *html* puede ser una ayuda adicional para los profesores a la hora de incluir los *applets* en sus páginas (transparencias, apuntes, ejercicios, etc.). Con esta herramienta podremos incluir los programas desde nuestro sistema local o bien desde un sitio *Web* remoto. Los usuarios de la página *Web* no encontrarán diferencia si se trata de un enlace

local o remoto salvo quizás por el tiempo de descarga propio de Internet de los códigos de las aplicaciones.

Para que el profesor consiga una máxima productividad, incluir un *applet* en una página *html* será tan fácil como incluir un fichero de imagen o un vídeo. Una simple operación de arrastre con el ratón podrá permitir poner un programa de gran calidad en las transparencias, apuntes o ejercicios que desee. Este documento podría servirse sobre la WWW desde la página Web inicial del profesor, o bien desde cualquier otro sitio de la red. La principal limitación es el reducido número de aplicaciones *applet* existentes con la calidad docente suficiente y la ausencia de esa herramienta que permita su integración tan sencilla para el autor. Es importante y oportuno desarrollar *applets* de gran calidad para varios temas de la asignatura. El beneficio redundará en la efectividad del aprendizaje y la productividad del educador que podrán compensar los altos costes involucrados en el esfuerzo del desarrollo del *applet*. Una vez desarrollada podremos crear una librería de *applets* que forme la colección base de la asignatura. Esta colección base puede continuar en un constante crecimiento con la aportación de las comunidades académica y de desarrolladores.

El propósito de nuestros *applets* es producir una visualización dinámica de los conceptos físicos de los materiales de estado sólido e ilustrar visualmente varios principios de los dispositivos. El énfasis se centra en la visualización interactiva de importantes conceptos que a menudo escapan de la atención del estudiante de los que es difícil dar una visión global con el material tradicional de aprendizaje solamente. Es lo que podríamos denominar una realización visual de la teoría.

Las páginas *html* se servirán desde un sitio Web del curso, desde la página del profesor, desde un sitio centralizado, o desde una combinación de los anteriores. Como el uso de los estudiantes de la WWW se ha convertido en una práctica común, el material instructivo basado en tecnologías Web puede ser integrado cómodamente dentro de la estructura de un curso tradicional compuesto normalmente por clases, prácticas de laboratorio y ejercicios.

Los *applets* tienen una gran efectividad en el aprendizaje de los estudiantes ya que su uso requiere que el alumno se involucre en el proceso. Como ejercicios propuestos, los *applets* deben acompañarse de un breve texto introductorio y de una hoja de trabajo que el alumno deberá completar mientras trabaja con el mismo. Utilizar otro tipo de programas en una presentación no beneficia tanto al alumno ya que hace que éste sea un mero observador. Si se hace esto como un paso previo al trabajo de cada estudiante con el *applet*, conseguiremos resultados muy positivos.

Un *applet* depende por naturaleza de la Interfaz gráfica de usuario para la entrada de datos y de una representación visual de los datos de la simulación. Por eso es una herramienta excelente para la enseñanza / aprendizaje de conceptos abstractos y principios de la ciencia. Los procesos de simulación y los resultados pueden ser presentados visualmente en tiempo real. Una variación dinámica del *applet* mantendrá un alto nivel de atención del alumno. Será deseable que un gran número de aplicaciones sirva para cubrir el temario de la asignatura en su totalidad.

Los sistemas de software destinado a la enseñanza se pueden clasificar de acuerdo con el modo que los estudiantes acceden (local y distribuido), si es expansible por los educadores / usuarios (cerrado y abierto), y el entorno de ejecución del software (ejecutable y embebido). Un curso debería permitir un acceso distribuido y una fácil expansión del sistema por los usuarios (educadores). Esto se puede llevar a cabo efectivamente si el software se ejecuta embebido en una página *html*.

Si dividimos el sistema entre cerrado y abierto, entonces un curso preempaquetado en una plataforma específica con un lenguaje como C será necesariamente un curso cerrado. Lo mismo podemos decir de una herramienta como Pspice. Los estudiantes tendrán acceso a ese curso cerrado bien mediante un terminal en el laboratorio o con un CD-ROM distribuido individualmente (p.e. con un libro de texto). Este sistema es rígido. Una ampliación o reorganización no es fácil.

Un curso formará un sistema abierto si es una colección de material con tecnología *Web* o si es compatible con alguna herramienta para hacer un curso ampliable. Un Tutorial con páginas *html* multimedia con imágenes gráficas, archivos de audio y vídeo es un ejemplo de curso abierto. Estos elementos multimedia son buenas herramientas de exposición pero tienen el peligro de hacer del alumno un simple observador pasivo en lugar de un participante activo. Por otro lado los *applets* de Java permiten al alumno interactuar más estrechamente con el material de aprendizaje, dándole un mayor control sobre éste. Los resultados de la simulación se obtienen en tiempo real. En este sentido los *applets* de Java añaden una nueva dimensión a los cursos multimedia basados en tecnologías *Web*. Tienen excepcionales características como herramientas explicativas y demostrativas. Un sistema es abierto en el sentido que permite un acceso distribuido a través de la WWW, los usuarios pueden añadir fácilmente nuevos objetos de software, la reorganización y modificación del sistema es fácilmente llevada a cabo, y la creación de nuevo material docente es rápida gracias a una librería de objetos de software.

De modo que como un elemento aislado, un *applet* permite mucha más interactividad y potencia de cálculo que cualquier otro elemento simple de un curso. Los *applets* de Java por consiguiente son muy adecuados para integrar muchos principios y conceptos relacionados en un mismo elemento. Cuando se integra con otros elementos multimedia dentro de un documento con hiperenlaces, que podrá verse con un explorador de Web, la tecnología de los *applets* de Java permite incrementar la calidad y eficiencia del aprendizaje de manera considerable.

Los *applets* de Java son más útiles cuando se trata de integrar conceptos relacionados y principios más que introducir nuevos conceptos para los que es más adecuado un tratamiento progresivo. Al enseñar un nuevo concepto es mejor desmenuzarlo en muchos conceptos sencillos e introducirlos poco a poco, en lugar de combinarlos todos en un solo programa. Debido al esfuerzo que requiere la programación y al tiempo consumido, es más difícil producir muchos *applets* que traten un concepto simple para combinarlos progresivamente. En este sentido y porque es una herramienta ideal para sintetizar principios y conceptos, los *applets* no son elementos lineales de un curso. No obstante, los programas constituyen una utilidad excelente para demostrar y visualizar conceptos difíciles, para ensamblar muchos principios y conceptos asociados en un simple elemento.

1.1. LA PÁGINA PRINCIPAL

El aspecto de la página principal del Tutorial lo podemos observar en la Figura 1.1.

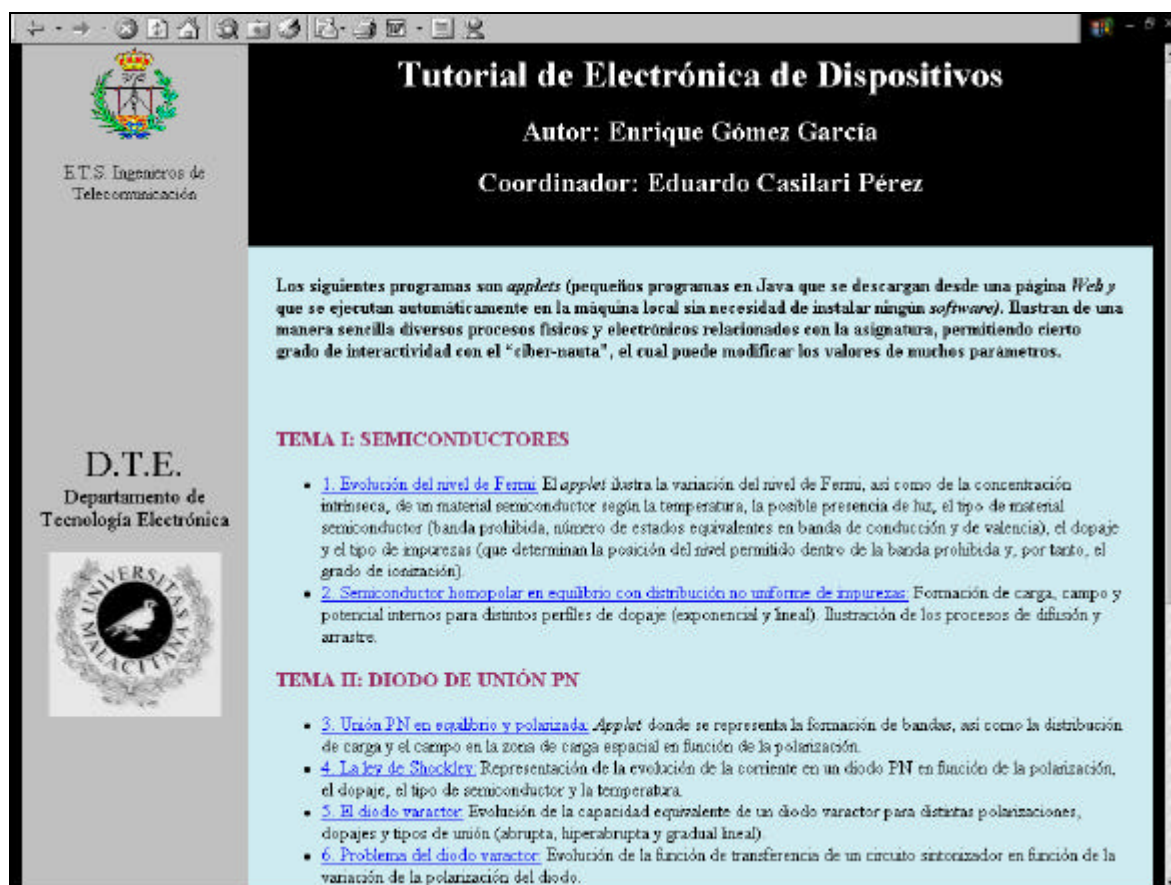


Figura 1.1. Página de inicio del Tutorial de Electrónica de Dispositivos

Esto es lo que primero se observa al acceder al Tutorial. Obviamente el aspecto general dependerá en menor medida de que aplicación que utilizemos para visualizar páginas *Web*, en este caso concreto se utiliza el Netscape® Communicator 4.72.

Como se puede ver a simple vista, tenemos una primera gran división de la materia. Dentro del programa de la asignatura de Electrónica de Dispositivos vamos a tratar los cuatro temas principales del temario. El primero de éstos es el de "Semiconductores", a él se dedican los dos primeros *applets*: Evolución del nivel de Fermi y Semiconductor homopolar en equilibrio con distribución no uniforme de impurezas. El segundo tema trata sobre el "Diodo de unión PN", para el cual se han dedicado cuatro *applets*: Unión PN en equilibrio y polarizada, la ley de Shockley, el diodo varactor y problema del diodo varactor. El tercer tema es "El transistor bipolar" sobre el que tratan los *applets* siete y

ocho: Distribución de corrientes y portadores y Conmutación del transistor. El último tema se titula “El transistor unipolar” y se estudia en los dos últimos *applets*: Canal en una estructura JFET y la capacidad MOS.

Cada *applet* está acompañado en su página *Web* correspondiente con unos textos de introducción en la materia y de guía de uso del mismo. Con esto el alumno podrá obtener una rápida visión global sobre el asunto en cuestión, y mediante el tutorial del *applet* podrá sacar al mayor provecho de la aplicación.

A continuación se describe pormenorizadamente la división de la asignatura presente en la página principal. Como se ha mencionado tenemos cuatro temas principales:

- TEMA I: SEMICONDUCTORES

1. Evolución del nivel de Fermi: El *applet* ilustra la variación del nivel de Fermi, así como de la concentración intrínseca, de un material semiconductor según la temperatura, la posible presencia de luz, el tipo de material semiconductor (banda prohibida, número de estados equivalentes en banda de conducción y de valencia), el dopaje y el tipo de impurezas (que determinan la posición del nivel permitido dentro de la banda prohibida y, por tanto, el grado de ionización).

2. Semiconductor homopolar en equilibrio con distribución no uniforme de impurezas: Formación de carga, campo y potencial internos para distintos perfiles de dopaje (exponencial y lineal). El programa ilustra los procesos de difusión y arrastre.

- TEMA II: DIODO DE UNIÓN PN

3. Unión PN en equilibrio y polarizada: *Applet* donde se representa la formación de bandas, así como la distribución de carga y el campo en la zona de carga espacial en función de la polarización.

4. La ley de Shockley: Representación de la evolución de la corriente en un diodo PN en función de la polarización, el dopaje, el tipo de semiconductor y la temperatura.

5. El diodo varactor: Evolución de la capacidad equivalente de un diodo varactor para distintas polarizaciones, dopajes y tipos de unión (abrupta, hiperabrupta y gradual lineal).

6. Problema del diodo varactor: Evolución de la función de transferencia de un circuito sintonizador en función de la variación de la polarización del diodo.

- TEMA III: EL TRANSISTOR BIPOLAR

7. Distribución de corrientes y portadores: Este *applet* permite observar la distribución de corrientes y portadores, así como los parámetros del transistor, de un transistor bipolar polarizado en función de la tensión aplicada. Los parámetros físicos básicos del transistor (dopaje, naturaleza del semiconductor) son programables también.

8. Conmutación del transistor: Ejemplificación del proceso de conmutación de un bipolar. Las gráficas representan las variaciones de la carga en base, las corrientes de colector y base y la tensión de salida, que se producen en un bipolar cuando cambia la tensión con la que está polarizado. El programa permite definir los parámetros del circuito de polarización así como los del propio transistor.

- TEMA IV: EL TRANSISTOR UNIPOLAR

9. Canal en una estructura JFET: Ilustración de la modificación que sufre el canal en un transistor JFET cuando se aplican tensiones en drenador. Se pueden alterar los parámetros físicos básicos del JFET.

10. La capacidad MOS: Representación de las variaciones y distribución espacial de carga, campo y potencial, así como de las bandas de energía, de una estructura MOS en función de la polarización. Los parámetros físicos básicos de la capacidad son programables.

Moverse por el tutorial es tan sencillo como lo es el manejo de las aplicaciones de visualización de objetos *Web*, popularmente llamadas exploradores o navegadores. Con los

hiper enlaces presentes en todas las páginas podremos ir al sitio deseado. Obviamente también disponemos de todas las facilidades que proporciona el explorador para ir a una página determinada utilizando los menús correspondientes, o bien el historial. Los hiper enlaces son las cadenas de texto que aparecen subrayadas, tienen colores distintos en función de si sus páginas *Web* enlazadas han sido visitadas anteriormente o no. Al hacer clic en cualquiera de los hiper enlaces listados anteriormente accederemos a la página que contiene el *applet* correspondiente. El manejo de los mismos se explicará con detalle más adelante en el Capítulo 3.

En general los *applets* se pueden utilizar para fijar conceptos si se usan como un ejemplo explicativo o como un resumen del tema tratado (acercamiento a posteriori). Los *applets* también pueden servir de vehículo introductorio para temas nuevos (acercamiento a priori). Y también se pueden utilizar en cualquier momento dentro del proceso de aprendizaje (acercamiento integral).

A posteriori se utilizan después de haber sido introducido un nuevo concepto de una manera tradicional. Seguramente esta es la forma más efectiva de usar los *applets* ya que de alguna manera se puede visualizar lo que realmente está sucediendo. En esta aproximación los *applets* vendrían al final de una lección o como tarea para realizar en casa. Servirían tanto para asegurar la comprensión del alumno de los nuevos conceptos como para ayudar en el aspecto cuantitativo. Para este último propósito es fundamental la interactividad y la computabilidad en tiempo real de los *applets*.

En la aproximación a priori los *applets* se utilizan antes de introducirse en un nuevo tema. Se pueden utilizar para mostrar una visión general del problema o para llamar la atención del alumno antes de proceder a una explicación más profunda.

Con el acercamiento integral al concepto se puede intercalar el trabajo con el *applet* con la lectura, los ejercicios propuestos y las explicaciones de clase. En este punto los *applets* se utilizan durante la lectura en el momento que el profesor considere más oportuno. Conforme se vayan haciendo más abundantes los temas tratados, esta última será la aproximación principal.

CAPÍTULO 2: ESTRUCTURA DE LOS *APPLETS*

2.1. JAVA

El lenguaje HTML permite describir la manera en que deben exhibirse las páginas estáticas de la *Web*, incluidas tablas y fotografías. Para hacer posible páginas *Web* altamente interactivas, se requiere un mecanismo como el lenguaje e intérprete Java^{MR}.

Java se originó cuando algunas personas de Sun Microsystems trataban de desarrollar un lenguaje nuevo adecuado para programar aparatos caseros orientados a información. Después se reorientó el lenguaje hacia la *World Wide Web*. Aunque Java toma prestadas muchas ideas y algo de la sintaxis de C y C++, es un lenguaje nuevo orientado a objetos, incompatible con ambos. A veces se dice que, en grande, Java es como Smalltalk pero, en pequeño, es como C o C++.

La idea principal de usar Java para páginas *Web* interactivas es que una página *Web* puede apuntar a un programa Java pequeño, llamado *applet* (que podría traducirse como "aplicacioncita"). Cuando el visualizador llega a ella, el *applet* se "baja" a la máquina cliente y se ejecuta ahí de una manera segura. Debe ser estructuralmente imposible que el *applet* lea o escriba archivos que no está autorizado a acceder. Debe también ser imposible que el *applet* introduzca virus o cause algún otro daño. Por estas razones, y para lograr transportabilidad entre máquinas, los *applets* se compilan para crear un código de bytes después de escribirse y depurarse. Son estos programas en código de bytes a los que apuntan las páginas *Web*, de manera parecida a como se apunta a las imágenes. Al llegar un *applet*, se ejecuta interpretativamente en un entorno seguro.

Antes de entrar en detalles del lenguaje Java, vale la pena decir algunas palabras sobre la utilidad del sistema Java y las razones por las que la gente quiere incluir *applets* Java en sus páginas de *Web*. Por una parte, los *applets* permiten que las páginas *Web* se vuelvan interactivas. Por ejemplo, una página de *Web* puede contener un tablero para jugar al ajedrez, y jugar un juego con el usuario. El programa de juego (escrito en Java) se baja junto con su página *Web*. Como segundo ejemplo, pueden presentarse formas complejas

(por ejemplo, hojas de cálculo), llenando los usuarios elementos y viendo instantáneamente los cálculos.

De esta forma es posible que, a la larga, el modelo de gente que compra programas, los instala y los ejecuta localmente será reemplazado por un modelo en el que la gente haga clic en las páginas *Web* y baje *applets* que trabajan para ella, posiblemente en colaboración con un servidor o base de datos remoto. En lugar de llenar la declaración de impuestos a mano o usando un programa especial, la gente podrá hacer clic en la página base de la Agencia Tributaria para bajar un *applet* de impuestos. Este *applet* podría hacer algunas preguntas, luego comunicarse con el patrón, banco y corredor de bolsa de la persona a fin de reunir la información necesaria sobre sueldos, intereses y dividendos, llenar el formulario y presentarlo para verificación y envío.

Otra razón para ejecutar *applets* en la máquina cliente es que hace posible la adición de animación y sonido a las páginas de *Web* sin tener que llamar a visualizadores externos. El sonido puede reproducirse cuando se carga la página, como música de fondo, o cuando ocurre algún suceso específico. Ocurre lo mismo con la animación. Como el *applet* se ejecuta localmente, aun si se está interpretando, puede escribir en cualquier parte de la pantalla de la manera que quiera y a muy alta velocidad.

El sistema Java tiene tres partes:

1. Un compilador de Java a código de bytes.
2. Un visualizador que entiende *applets*.
3. Un intérprete de código de bytes.

El programador escribe el *applet* en Java, luego lo compila, obteniendo un código de bytes. Para incluir este *applet* compilado en una página *Web*, se ha inventado una nueva etiqueta HTML, `<APPLET>`. Un uso típico es:

```
<APPLET CODE=fermi.class WIDTH=100 HEIGHT=200> </APPLET>
```

Cuando el visualizador ve la etiqueta `<APPLET>`, trae el *applet* compilado `fermi.class` de la instalación de la página actual de la *Web* (o, si está presente otro parámetro, `CODEBASE`, del URL que especifica). El visualizador entonces pasa el *applet* al intérprete local de código de bytes para su ejecución (o interpreta el *applet* él mismo, si tiene un intérprete interno). Los parámetros `WIDTH` y `HEIGHT` dan el tamaño de la ventana predeterminada del *applet*, en píxeles.

De alguna manera, la etiqueta `<APPLET>` es análoga a la etiqueta `<IMG>` que se utiliza para las imágenes estáticas. En ambos casos, el visualizador trae un archivo y luego lo

entrega a un intérprete (posiblemente interno) para presentarlo en un área limitada de la pantalla; luego continúa procesando la página *Web*.

Para algunas aplicaciones que requieren mucha computación, algunos intérpretes Java tienen la capacidad de compilar programas de código de bytes a lenguaje máquina sobre la marcha, según se requiera.

Como consecuencia de este modelo, los visualizadores basados en Java son extensibles de una manera que no lo son los visualizadores de primera generación. Éstos básicamente son intérpretes HTML que tienen módulos interconstruidos para manejar los diferentes protocolos necesarios, como HTTP 1.0, FTP, etc., al igual que decodificadores de varios formatos de imagen. Si alguien inventa o populariza un formato nuevo, como audio o MPEG-2, estos visualizadores viejos no serán capaces de leer las páginas que los contienen. Cuando mucho, el usuario tendrá que encontrar, bajar e instalar un visualizador externo adecuado.

Con un visualizador basado en Java, la situación es diferente. Al arranque, el visualizador de hecho es una máquina virtual Java vacía. Al cargar *applets* HTML y HTTP, se vuelve capaz de leer páginas de Web estándar. Sin embargo, a medida que se requieren protocolos y decodificadores nuevos, sus clases se cargan de manera dinámica, posiblemente a través de la red desde instalaciones especificadas en las páginas *Web*.

Por tanto, si alguien inventa un formato nuevo, todo lo que tiene que hacer la persona es incluir el URL de un *applet* para manejarla en una página *Web*, y el visualizador automáticamente traerá y cargará el *applet*. Ningún visualizador de primera generación es capaz de traer e instalar automáticamente visualizadores externos nuevos sobre la marcha. La capacidad de cargar visualizadores dinámicamente significa que la gente puede experimentar fácilmente con formatos nuevos sin primero tener que realizar interminables juntas de estandarización para llegar a un consenso.

Esta capacidad de extensión también se aplica a los protocolos. En algunas aplicaciones se requieren protocolos especiales, por ejemplo, protocolos seguros para aplicaciones bancarias y comerciales. Con Java, estos protocolos pueden cargarse dinámicamente según se requiera, y no hay necesidad de lograr una estandarización universal. Para comunicarse con la compañía X, simplemente bajamos su *applet* de protocolo. Para hablar con la compañía Y, traemos su *applet* de protocolo. No hay necesidad de que X y Y acuerden un protocolo estándar.

Los objetivos antes listados han conducido a un lenguaje de tipo seguro, orientado a objetos, con capacidad multihilos interconstruida y sin características no definidas o

dependientes del sistema. Lo que sigue es una descripción muy simplificada de Java, simplemente para dar una idea de su sabor. Se han omitido muchas características, detalles, opciones y casos especiales en aras de la brevedad. La especificación completa del lenguaje, y mucho más sobre Java, está disponible en la *Web* misma.

Como mencionamos antes, en pequeño, Java es parecido a C y C++. Las reglas léxicas, por ejemplo, son muy parecidas.

En los lenguajes de procedimientos como Pascal o C, un programa consiste en un conjunto de variables y procedimientos, sin ningún principio general de organización. En contraste, en los lenguajes orientados a objetos (casi) todo es un objeto. Un objeto normalmente contiene algunas variables de estado internas (es decir, escondidas), junto con algunos procedimientos públicos, llamados métodos, para acceder a ellos. Se espera (y puede obligarse a que) los programas que usan el objeto invoquen los métodos para manipular el estado del objeto. De esta manera, el escritor del objeto puede controlar la manera en que los programas usan la información del objeto. Este principio se llama encapsulamiento, y es la base de toda la programación orientada a objetos.

Java intenta tomar lo mejor de ambos mundos: puede usarse como lenguaje de procedimientos tradicional o como lenguaje orientado a objetos. Desde este punto de vista, un subgrupo de Java podría considerarse como una versión depurada de C. Sin embargo, para escribir páginas *Web*, es mejor considerar a Java como un lenguaje orientado a objetos, por lo que se analizará su orientación a objetos aquí.

Un programa en Java consiste en uno o más paquetes, cada uno de los cuales contiene algunas definiciones de clases. Los paquetes pueden accederse remotamente a través de una red, por lo que aquellos destinados para ser usados por mucha gente deben tener nombres únicos.

Una definición de clase es una plantilla para producir ejemplares de objetos, cada uno de los cuales contiene las mismas variables de estado y los mismos métodos que todos los otros ejemplares de objetos de su clase. Sin embargo, los valores de las variables de estado de los diferentes objetos son independientes. Las clases, por tanto, son como moldes para cortar galletas: no son las galletas realmente, pero sirven para producir galletas estructuralmente idénticas, produciendo cada molde de galletas una forma diferente de galleta. Una vez producidas, las galletas (objetos) son independientes entre sí.

Los objetos en Java pueden producirse dinámicamente durante la ejecución, por ejemplo mediante

```
objeto = new nombreClase()
```

Estos objetos se almacenan en el pila y los elimina el reciclador de memoria dinámica cuando ya no se necesitan. De esta manera, la administración de almacenamiento en Java es manejada por el sistema, sin necesidad de los temidos procedimientos en C *malloc* y *free*, ni tampoco de apuntadores específicos.

Cada clase se basa en otra clase. Se dice que una clase de nueva definición es una subclase de la clase en la que se basa, la superclase. Una (sub)clase siempre hereda los métodos de su superclase; puede o no tener acceso directo a las variables internas de la superclase, dependiendo de si la superclase lo quiere o no. Por ejemplo, si una superclase, A, tiene los métodos M1, M2 y M3, y una subclase B, define un método nuevo, M4, entonces los objetos creados en B tendrán los métodos M1, M2, M3 y M4. La propiedad por la que una clase automáticamente adquiere todos los métodos de su superclase se llama herencia, y es una propiedad importante de Java. La acción de agregar nuevos métodos a los métodos de la superclase se llama extender la superclase. Como nota, algunos lenguajes orientados a objetos permiten que las clases hereden los métodos de dos o más superclases (herencia múltiple), pero los diseñadores de Java pensaron que esta propiedad era demasiado desordenada e intencionadamente la dejaron fuera.

Dado que cada clase tiene una y sólo una superclase inmediata, el grupo de todas las clases de un programa en Java forma un árbol. La clase raíz del árbol se llama *Object*, y todas las demás clases heredan sus métodos. Cualquier clase cuya superclase no se menciona explícitamente en su definición es por omisión una subclase de la clase *Object*.

Además del lenguaje básico, los diseñadores de Java han definido e implementado unas 200 clases con la versión inicial. Los métodos contenidos en estas clases forman una especie de entorno estándar para los creadores de programas Java. Las clases están escritas en Java, por lo que son transportables a todas las plataformas y sistemas operativos.

Las 200 clases se agrupan en siete paquetes de tamaño desigual, cada uno de los cuales se enfoca hacia algún tema central. Los *applets* que necesiten un paquete en particular pueden incluirlo usando el enunciado *import* de Java. Los métodos contenidos pueden usarse según se requiera. Este mecanismo reemplaza la necesidad de incluir archivos de cabecera en C; también reemplaza la necesidad de bibliotecas, puesto que los paquetes se cargan dinámicamente durante la ejecución al ser invocados.

Brevemente se describen estos siete paquetes. El paquete *java.lang* contiene clases que pueden verse como parte del lenguaje, pero no lo son técnicamente. Éstas incluyen clases para administrar las clases mismas, hilos y manejo de excepciones. También están aquí las bibliotecas estándar de matemáticas y de cadenas.

Al igual que C, el lenguaje Java no contiene primitivas de E/S. La E/S se hace cargando y usando el paquete *java.io*, análogo a la biblioteca de E/S estándar de C. Se proporcionan métodos para leer y escribir cadenas, archivos de acceso aleatorio y hacer el formateo necesario para imprimir.

Muy relacionado con la E/S está el transporte por red. Los métodos que buscan y manejan las direcciones IP se localizan en el paquete *java.net*. El acceso a los sockets también forma parte de este paquete, al igual que la preparación de datagramas. La transmisión misma se maneja con *java.io*.

La siguiente clase es *java.util*; contiene las clases y métodos para las estructuras de datos comunes, como pilas y tablas de dispersión, de modo que los programadores no tengan que reinventar constantemente la rueda. También se lleva a cabo aquí la administración de fecha y hora.

El paquete *java.applet* contiene parte de la maquinaria básica de los *applets*, incluidos métodos para traer páginas Web comenzando por sus URL. También tiene métodos para presentar páginas Web y ejecutar segmentos de audio (por ejemplo, música de fondo). El paquete *java.applet* también contiene la clase *Object*. Todos los objetos heredan sus métodos, a menos que se reemplacen. Estos métodos incluyen hacer clones de un objeto, comparar la igualdad de dos objetos, convertir un objeto en una cadena, y otros más.

Por último, llegamos a *java.awt* y sus dos subpaquetes *image* y *peer*. AWT (Abstract Window Toolkit), herramientas de ventana abstracta, está diseñado para hacer que los *applets* sean transportables entre los sistemas de ventanas. Por ejemplo, ¿cómo debe un *applet* dibujar un rectángulo en la ventana de tal manera que pueda ejecutarse la misma versión compilada (en código de bytes) en UNIX, Windows y Macintosh, aunque cada uno tenga su propio sistema de ventanas? Parte del paquete se encarga de dibujar en la pantalla, por lo que hay métodos para colocar líneas, figuras geométricas, texto, menús, botones, barras de desplazamiento y muchos otros elementos en ella. Los programadores de Java llaman a estos métodos para escribir en la pantalla. Es responsabilidad del paquete *java.awt* hacer las llamadas adecuadas al sistema operativo local para llevar a cabo el trabajo. Esta estrategia significa que *java.awt* tiene que describirse para cada plataforma nueva, pero los *applets* entonces son independientes de la plataforma, lo que es mucho más importante.

Otra tarea importante de esta clase es la administración de eventos. La mayoría de los sistemas de ventanas fundamentalmente son controlados por eventos. Lo que quiere decir esto es que el sistema operativo detecta pulsaciones de teclas, movimientos del ratón,

presión y liberación de botones, y otros eventos, y los convierte en llamadas a procedimientos de usuario. En el caso de Java, se proporciona una gran biblioteca de métodos para manejar estos eventos en *java.awt*. Su uso simplifica la escritura de programas que interactúen con el sistema local de ventanas y aun así sean 100% transportables a máquinas con diferentes sistemas operativos y diferentes sistemas de ventanas.

Una parte del trabajo de este paquete se hace en *java.awt.image*, por ejemplo el manejo de imágenes, y en *java.awt.peer*, que permite el acceso al sistema subyacente de ventanas.

En el presente tutorial principalmente se hace uso de los paquetes *java.applet* y *java.awt* y en ocasiones se accederá a algunas clases del paquete *java.util*.

Uno de los aspectos más importantes de Java es sus características de seguridad. Cuando se trae una página *Web* de la red que contiene un *applet*, ésta automáticamente se ejecuta en la máquina del cliente. Idealmente, no debe provocar la caída ni problemas en la máquina del cliente.

Es más, no se requiere mucha imaginación para imaginar a un estudiante emprendedor preparando una página *Web* que contiene algún interesante juego nuevo, y luego haciendo publicidad mundialmente a su URL (por ejemplo, publicándola de manera cruzada en todos los grupos de noticias). Lo que no se menciona en la publicación es el pequeño detalle de que la página también contiene un *applet* que, al ejecutarse, de inmediato cifra todos los archivos del disco duro del usuario. Cuando ha terminado, el *applet* anuncia lo que ha hecho y menciona cortésmente que los usuarios que deseen comprar la clave de descifrado deben enviar 1000 dólares en billetes de baja denominación a cierto apartado postal de Panamá.

Además del esquema de millonario instantáneo anterior, hay otros peligros inherentes a permitir la ejecución de código extraño en una máquina. Un *applet* podría andar a la caza de información interesante (correo electrónico almacenado, el archivo de contraseñas, las cadenas del entorno local, etc.) y devolverlas o enviarlas por correo electrónico a través de la red. El *applet* podría también consumir recursos (por ejemplo, llenar discos), presentar fotografías obscenas o consignas políticas en la pantalla, o hacer ruidos irritantes usando la tarjeta de sonido.

Los diseñadores de Java estaban bien conscientes de estos problemas, y erigieron una serie de barreras contra ellos. La primera línea de defensa es un lenguaje seguro respecto a los tipos. Java tiene especificaciones detalladas de los tipos de datos, arreglos verdaderos con verificación de límites y no tiene apuntadores. Estas restricciones hacen imposible que

un programador de Java construya un apuntador para leer y escribir localidades arbitrarias de la memoria.

En pocas palabras, Java genera muchas posibilidades y oportunidades nuevas en la WWW; permite que las páginas *Web* sean interactivas y contengan animación y sonido; también permite que los visualizadores sean infinitamente extensibles. Sin embargo el modelo Java de descargar *applets* también genera algunos problemas de seguridad nuevos y graves que aún no se han resuelto por completo.

Otro tipo de problemas que se pueden presentar es la dependencia que pueden tener los *applets* de Java para el uso educativo con la tecnología, y particularmente con el soporte por parte de los exploradores populares de la versión de Java en la que el *applet* fue escrito. Esto quiere decir que si hemos programado el *applet* con la última versión del JDK (Java Development Kit), probablemente no se podrán ejecutar en la mayoría de los exploradores. La versión actual del JDK es la 1.3 pero si utilizamos una versión anterior nos ahorraremos ciertos problemas de incompatibilidad, esto es posible gracias a que los *applets* del tutorial no necesitan utilizar las últimas librerías desarrolladas ya que se trata de aplicaciones bastante sencillas. El SDK (*Software Development Kit*) se puede descargar de Internet de la página Web de Sun Microsystems Inc. Es el paquete con todas las aplicaciones necesarias para desarrollar programas en Java, aunque realmente se ha utilizado el software de Borland Jbuilder que se basa en el citado SDK por su potencia y características especiales, como el *debugger* para depurar aplicaciones que resultó ser de extrema utilidad.

2.2. DISEÑO E IMPLEMENTACIÓN DE LOS *APPLETS*

Dentro de un tema concreto, un *applet* puede estar limitado a un solo concepto, o puede tratar varios conceptos relacionados. La selección de un tema apropiado para un *applet* y una buena práctica diseñándolo son extremadamente importantes.

Al ser una herramienta visual para el aprendizaje, el aspecto del interfaz de usuario de la aplicación es importante. Un mal diseño del GUI (*Graphical User Interface*) puede llevarnos a un *applet* que no sea muy útil a pesar del gran esfuerzo realizado en su desarrollo. El GUI del *applet* se puede desarrollar por separado del resto del programa, haciendo uso de las ventajas de la programación orientada a objetos de Java. Para el GUI también podemos utilizar una herramienta de desarrollo visual, en el presente caso se ha utilizado el Jbuilder Enterprise 5 de Borland Software Corporation. Esta utilidad permite desarrollar rápidamente un prototipo de interfaz gráfico. Los otros contenidos del *applet* (es decir, los contenidos técnicos y la simulación real) se desarrollan aparte del GUI. Aquí de nuevo, una disciplinada aproximación al diseñar con el paradigma del diseño orientado a objetos va en beneficio para futuros mantenimientos del programa y para incrementar la productividad en la programación. El desarrollo de clases de Java reutilizables puede llevarnos a recompilar una librería de clases de Java en este particular área de los dispositivos semiconductores. Esta librería de clases facilitará sobremedida la reutilización del software e incrementará el rendimiento.

El propósito de esta memoria no es la de dar una exhaustiva descripción del código fuente, ni se pretende realizar un curso de programación. Tampoco se incluye código dentro del presente texto. A tal efecto junto con el tutorial completo se suministrarán los códigos fuentes en un soporte físico que se adjuntará al final del libro. Lo que se pretende es dar una idea general de la estructura común que constituye los cimientos de los *applets*. Obviamente habrá grandes diferencias entre ellos en función de la orientación tomada en cada caso, pero esto no es óbice para poder abstraer el soporte común.

Para escribir *applets* de Java hay que utilizar una serie de métodos. Incluso para el *applet* más sencillo. Son los que se usan para arrancar (*start*) y detener (*stop*) la ejecución del *applet*, para pintar (*paint*) y actualizar (*update*) la pantalla y para capturar la información que se pasa al *applet* desde el fichero *html* a través de la marca `APPLET`.

Los *applets* no necesitan un método *main()* como las aplicaciones Java, sino que deben implementar (redefinir) al menos uno de los tres métodos siguientes: *init()*, *start()* o *paint()*.

En el apartado anterior ya se habló del AWT, se trata de una biblioteca de clases Java para el desarrollo de Interfaces Gráficas de Usuario. Hay que decir también que es la parte más débil de todo lo que representa Java como lenguaje. El entorno que ofrece es demasiado simple, no se han tenido en cuenta las ideas de entornos gráficos novedosos.

A esta biblioteca de clases pertenecen la mayoría de las que se utilizarán en todos los *applets* del tutorial. Por citar algunas: Button, Canvas, Checkbox, Choice, Panel, etc.

Como ejemplo de estructura de los *applets* se utilizará el de “La ley de Shockley” ya que representa lo que podría denominarse el *applet standard*. Está compuesto de siete clases. La clase principal se llama “pnJunction.class” y si se abre directamente con el Jbuilder se obtiene información sobre su lista de importaciones, las variables que utiliza y los métodos que define. Esta información se recopila de la siguiente forma:

```
// JBuilder API Decompiler stub source generated from class file
// 28-sep-01
// -- implementation of methods is not available
// Imports
import java.lang.String;
import java.awt.Event;
import java.applet.Applet;
public class pnJunction extends Applet {
    // Fields
    boolean inAnApplet;
    MainPanel mainpanel;
    ControlPanel control;
    PropertyPanel property;
    // Constructors
    public pnJunction() { }
    // Methods
    public void init() { }
    public boolean handleEvent(Event p0) { }
    public void start() { }
```

```

public void stop() { }
public static void main(String[] p0) { }
}

```

Esta clase `pnJunction.class` extiende a la clase `Applet` como puede verse en la definición, e implementa los métodos típicos de un *applet* como `init()`, `start()` y `stop()`. En la lista de variables se puede observar que hay una variable `mainpanel` que es de la clase `MainPanel`, otra variable `control` de la clase `ControlPanel` y la variable `property` correspondiente a `PropertyPanel`. Esta es la estructura común a casi todos los *applets* del tutorial. La clase `MainPanel` define el panel central del programa donde se desarrolla la animación en cuestión. Luego existen normalmente dos paneles más de control, en la parte superior y en la inferior, que controlan la animación y las propiedades del dispositivo.

Estos paneles de control son una extensión de la clase `Panel` de la librería AWT. La información que proporciona el Jbuilder acerca de la clase `MainPanel` es la siguiente:

```

// JBuilder API Decompiler stub source generated from class file
// 28-sep-01
// -- implementation of methods is not available
// Imports
import java.awt.Panel;
import java.awt.Color;
import java.lang.StringBuffer;
import java.awt.Event;
import java.awt.Dimension;
import java.awt.Image;
import java.awt.Graphics;
class MainPanel extends Panel {
    // Fields
    Particle[] particles;
    // Constructors
    MainPanel() { }
    // Methods
    public void preParameter() { }
    public void init() { }
}

```

```
public boolean handleEvent(Event p0) { }  
public void start() { }  
public void Vstart() { }  
public void stop() { }  
public void suspend() { }  
public void resume() { }  
public void update(Graphics p0) { }  
public void paint(Graphics p0) { }  
public void paintApplet(Graphics p0) { }  
void paintCurrent(Graphics p0) { }  
void paintFixed() { }  
}
```

Se ha eliminado la mayoría de las variables para no tener un listado engorroso y fuera del objetivo de presentar la estructura del *applet*. Lo que se intenta es destacar que efectivamente la clase *MainPanel* extiende a la clase *Panel* y la utilización prácticamente exclusiva de las clases de la biblioteca AWT.

Ahora quedarán los paneles de control superior e inferior que regulan el funcionamiento de la animación y las propiedades del dispositivo respectivamente. Ambas clases vuelven a ser extensiones de la clase *Panel*, o lo que es lo mismo heredarán todas sus propiedades. Como ejemplo de estos dos paneles de control se presenta la composición de la clase *ControlPanel* a continuación.

```
// JBuilder API Decompiler stub source generated from class file  
// 28-sep-01  
// -- implementation of methods is not available  
// Imports  
import java.lang.Object;  
import java.awt.Panel;  
import java.awt.Component;  
import java.awt.GridBagConstraints;  
import java.awt.GridBagLayout;  
import java.awt.Button;  
import java.awt.Event;
```

```
import java.awt.Checkbox;
class ControlPanel extends Panel {
    // Fields
    MainPanel mainpanel;
    PropertyPanel property;
    pnJunction applet;
    Button Exit;
    Button Pause;
    Button Stop;
    Checkbox onebox;
    Checkbox twobox;
    Checkbox threebox;
    // Constructors
    public ControlPanel(MainPanel p0, PropertyPanel p1, pnJunction p2) { }
    // Methods
    void setadd(Component p0, GridBagLayout p1, GridBagConstraints p2) { }
    public boolean action(Event p0, Object p1) { }
}
```

En el listado de variables se pueden observar los controles de tipo botones y cuadros de selección que componen la barra de control del panel superior.

Esto es realmente el *applet*, sólo quedan algunas clases más que sirven para la representación de los números en pantalla que son las clases `Format` y `DataFormatException`. Y la última clase que completa el *applet* es la clase `Particle` que se encarga de simular el flujo de partículas que forman las corrientes de huecos y electrones. La estructura de la clase `Particle` es la siguiente:

```
// JBuilder API Decompiler stub source generated from class file
// 28-sep-01
// -- implementation of methods is not available
// Imports
import java.awt.Color;
import java.lang.Runnable;
import java.lang.Thread;
```

```
import java.awt.Graphics;
class Particle implements Runnable {
    // Constructors
    public Particle(MainPanel p0, int p1) { }
    // Methods
    public void draw_particle(int p0, int p1, Graphics p2) { }
    public void move_particle(int p0, int p1, int p2, Graphics p3) { }
    public void paint(Graphics p0) { }
    public void start() { }
    public void stop() { }
    public void suspend() { }
    public void resume() { }
    public void run() { }
}
```

En este caso se trata de una tarea ya que hay dos modos de conseguir tareas o “hilos” (*threads*) en Java. Una implementando la interfaz `Runnable` que es la del caso actual, y la otra extender la clase `Thread`.

La implementación de la interfaz `Runnable` es la forma habitual de crear tareas. Las interfaces proporcionan al programador una forma de agrupar el trabajo de infraestructura de una clase. Se utilizan para diseñar los requerimientos comunes al conjunto de clases que se implementan. La interfaz define el trabajo y la clase, o clases, que implementan la interfaz para realizar ese trabajo. De modo que los diferentes grupos de clases que implementen la interfaz tendrán que seguir las mismas reglas de funcionamiento.

En resumen se puede afirmar que aunque puedan existir *applets* dentro de el presente tutorial que difieran tanto en la forma como en el contenido con la aplicación utilizada de ejemplo, la estructura común de todos y la metodología aplicada está fielmente reflejada en este ejemplo.

CAPÍTULO 3: DESCRIPCIÓN DE LOS *APPLETS* Y MANUAL DE USUARIO

3.1. EVOLUCIÓN DEL NIVEL DE FERMI

Las impurezas donadoras o aceptadoras en los semiconductores participan de modo muy activo en el aporte de electrones y huecos a las bandas de conducción y valencia y modifican, por tanto, apreciablemente sus propiedades de conducción eléctrica. La adición controlada de estas impurezas constituye por ello el fundamento básico de la obtención de dispositivos semiconductores.

En los semiconductores con impurezas o extrínsecos, a diferencia de los intrínsecos, deja de verificarse la igualdad de concentraciones de electrones y huecos, y es otra relación más general la que, combinada con la ley de masas, permite determinar dichas concentraciones en equilibrio. En el caso particular de los semiconductores con distribución uniforme de impurezas esta nueva relación es de formulación simple y establece la neutralidad de carga en el entorno de cualquier punto interior del material.

Se analiza la ecuación de neutralidad de carga en relación con las propiedades estadísticas de equilibrio para encontrar la dependencia del nivel de Fermi y de las concentraciones de portadores con las de impurezas y con la temperatura.

Aunque estrictamente sólo son aplicables a semiconductores homogéneos, los resultados de este análisis mantienen un elevado grado de validez incluso en muchos casos de semiconductores no homogéneos.

Los electrones de conducción y huecos en un semiconductor se originan no solamente por rotura estadística de algunos enlaces de la red cristalina sino también por ionización de impurezas donadoras o aceptadoras cuando éstas están presentes en el material. La proximidad de los niveles de impureza a las bandas de conducción o valencia hace de hecho, a este último proceso más probable que el primero.

En general las impurezas pueden no encontrarse distribuidas uniformemente en un semiconductor. Salvo a temperaturas elevadas (de varios cientos de °C al menos) los

átomos o iones de impureza están fijos, enclavados en la red cristalina y no pueden desplazarse por ella.

Por el contrario los electrones y huecos son partículas de gran facilidad de movimientos, y en equilibrio tenderán a distribuirse homogéneamente en el interior del material de forma similar a las partículas de un gas clásico en el recinto que lo encierra. Sin embargo, a diferencia de los gases clásicos, los electrones y huecos son partículas cargadas y en su tendencia a la homogeneidad dejan tras de sí a las impurezas ionizadas inmóviles originándose una disociación de las cargas que, a su vez, da lugar a un campo eléctrico interno. Este campo eléctrico actúa sobre los electrones y huecos móviles impidiéndoles completar la homogeneidad de su distribución, de manera que en equilibrio todas las tendencias de sentidos contrarios se igualan.

Si la densidad de carga en equilibrio en cada punto es conocida, dispondremos de un medio para determinar la distribución espacial de las concentraciones de electrones y huecos.

En efecto, si N_D^+ y N_A representan las concentraciones de impurezas fijas ionizadas, donadoras y aceptadoras, en cada punto, la densidad de carga r se expresa como:

$$r = e (N_D^+ + p - n - N_A) \quad (3.1)$$

ya que las restantes cargas (núcleos, electrones internos y de valencia, etc), se encuentran localmente compensadas. Si r , N_D^+ y N_A son o pueden ser funciones conocidas, la Ecuación 3.1 y la ley de masas ($n \cdot p = n_i^2$) determinan individualmente las distribuciones n y p .

Conviene recalcar que la densidad de carga r se produce por disociación de las cargas de impurezas y de portadores y la Ecuación 3.1 es, por tanto, una condición local. Sin embargo globalmente estas cargas se compensan, es decir el material en su conjunto es eléctricamente neutro.

Las consideraciones anteriores se simplifican notablemente en el caso particular de distribuciones uniformes de impurezas. Un dopaje homogéneo da lugar directamente a concentraciones también homogéneas de impurezas ionizadas y de portadores y no existe disociación de cargas. La densidad de carga ρ es, por tanto, nula y el material localmente neutro. Esta condición se expresa matemáticamente por la ecuación de neutralidad de carga deducida de la Ecuación 3.1:

$$p + N_D^+ = n + N_A \quad (3.2)$$

Esta ecuación engloba, como caso particular, el caso de los semiconductores intrínsecos en que, por la inexistencia de impurezas, p y n son iguales. Equivalente al caso intrínseco, desde el punto de vista de las concentraciones de electrones y huecos, es el de un semiconductor en que las concentraciones de impurezas ionizadas donadoras y aceptadoras sean iguales, el cual recibe el nombre de semiconductor intrínseco compensado.

Por el contrario, si $N_D^+ > N_A^-$ de la Ecuación 3.1 resulta $n > p$ y el semiconductor se dice extrínseco de tipo n . Análogamente si $N_A^- > N_D^+$ y, por tanto, $p > n$ el semiconductor es extrínseco de tipo p . A estas denominaciones se puede añadir la de parcialmente compensado para distinguir el caso de presencia de ambos tipos de impurezas ($N_D^+ \neq 0$ y $N_A^- \neq 0$) de aquel en que existe un solo tipo de ellas ($N_A^- = 0$ o $N_D^+ = 0$).

En los semiconductores extrínsecos los portadores cuya concentración es dominante reciben el nombre de mayoritarios y los del tipo contrario minoritarios. Es de notar que, en virtud de la ley de masas, las concentraciones de portadores mayoritarios y minoritarios en equilibrio son, a su vez, respectivamente mayor y menor que la concentración intrínseca. Este resultado es una nueva manifestación de la propiedad natural de tender hacia el mínimo gasto energético: la mayor parte de los portadores mayoritarios proceden de la ionización de las impurezas, que precisa una energía equivalente a la diferencia entre el nivel de impureza y la banda correspondiente, mientras que la rotura de enlaces, que precisa la energía de la banda prohibida, se limita a la cantidad justa para producir los portadores minoritarios exigidos por la ley de masas. En términos simplificados se podría decir que la ecuación de neutralidad de carga regula fundamentalmente la concentración de portadores mayoritarios y la ley de masas la de minoritarios.

Cuantitativamente estos resultados pueden expresarse del modo siguiente:

El sistema de dos ecuaciones constituido por la de neutralidad de carga (Ecuación 3.1) y la ley de masas es equivalente al de las ecuaciones de 2º grado

$$n^2 - n(N_D^+ - N_A^-) - n_i^2 = 0 \quad (3.3)$$

$$p^2 - p(N_D^+ - N_A^-) - n_i^2 = 0 \quad (3.4)$$

cuya solución correcta (la concentración es una magnitud esencialmente positiva) es:

$$n = \frac{N_D^+ - N_A^- + \sqrt{(N_D^+ - N_A^-)^2 + 4n_i^2}}{2} \quad (3.5)$$

$$p = \frac{N_A^- - N_D^+ + \sqrt{(N_A^- - N_D^+)^2 + 4n_i^2}}{2} \quad (3.6)$$

Aunque ambas expresiones son válidas en todo caso, resulta cómodo utilizar la Ecuación 3.5 para semiconductores extrínsecos tipo n , especialmente si $N_D^+ - N_A^- \gg 2n_i$ en cuyo caso:

$$n \cong N_D^+ - N_A^- \quad (3.7)$$

y a partir de este resultado determinar p mediante

$$p = \frac{n_i^2}{n} \cong \frac{n_i^2}{N_D^+ - N_A^-} \quad (3.8)$$

o, análogamente, en el caso extrínseco tipo p si $N_A^- - N_D^+ \gg 2n_i$, la Ecuación 3.6 conducirá a:

$$p \cong N_A^- - N_D^+ \quad (3.9)$$

y, por la ley de masas:

$$n = \frac{n_i^2}{p} \cong \frac{n_i^2}{N_A^- - N_D^+} \quad (3.10)$$

A continuación se analiza más detalladamente el cálculo de concentraciones de electrones y huecos partiendo del de la posición del nivel de Fermi y suponiendo que el dopaje (concentraciones de impurezas totales) es conocido.

La ecuación de neutralidad de carga y las que de ella se han deducido anteriormente no establecen una relación directa entre concentraciones de portadores y de impurezas totales sino sólo entre concentraciones de portadores y de impurezas ionizadas. Sin embargo tanto unas como otras están relacionadas con la posición relativa del nivel de Fermi dentro de la banda prohibida. En efecto, si se considera válida la función de Fermi como factor de ocupación de los niveles de impureza, en el supuesto de un semiconductor no degenerado se tiene que:

$$M_v \exp \frac{E_v - E_F}{kT} + N_D \frac{1}{1 + \exp \frac{E_F - E_D}{kT}} = M_c \exp \frac{E_F - E_c}{kT} + N_A \frac{1}{1 + \exp \frac{E_A - E_F}{kT}} \quad (3.11)$$

que es una expresión implícita de E_F , como única variable, cuando se conoce el semiconductor de base y la temperatura (M_c , M_v , $E_c - E_v$) y sus impurezas (N_D , N_A , $E_c - E_D$, $E_A - E_v$). La determinación de E_F , para cada caso particular, por resolución de la Ecuación 3.11 representa, en definitiva, la resolución del problema completo puesto que las restantes magnitudes de equilibrio (concentraciones de portadores y de impurezas ionizadas) se

determinan por su relación con E_F . Sin embargo la Ecuación 3.11 no admite, en general, una solución analítica para E_F y se requieren métodos numéricos.

La característica de variación casi en escalón de la función de Fermi justifica la adopción del siguiente método iterativo para calcular las magnitudes de equilibrio:

- a) Inicializar el proceso de cálculo con una estimación de la posición de E_F . En la práctica no es necesario que esta estimación inicial sea muy precisa y puede suponerse una posición cualquiera (por ejemplo el centro de la banda prohibida).
- b) Calcular las concentraciones de impurezas ionizadas N_D^+ y N_A^- mediante los términos correspondientes de la Ecuación 3.11.
- c) Calcular n y p mediante las Ecuaciones 3.5 y 3.6 o, en su caso, con las Ecuaciones de la 3.7 a la 3.10.

- d) Reestimar E_F a partir de los resultados anteriores mediante $E_F = E_c - kT \ln \frac{M_c}{n}$ o

$$E_F = E_v + kT \ln \frac{M_v}{p}.$$

- e) Repetir el ciclo desde b) hasta que las variaciones entre dos ciclos consecutivos sean despreciables. En general una o dos iteraciones suelen ser suficientes.

Este es el algoritmo que se utiliza en el *applet*. Aunque este sea un método general no es fácil extraer de él conclusiones para casos particulares. Con este objeto se adopta un punto de vista algo diferente basado en el siguiente esquema de razonamiento:

Consideraciones cualitativas $\Rightarrow$ hipótesis simplificadoras $\Rightarrow$ Resultados cuantitativos $\Rightarrow$ comprobación de las hipótesis iniciales.

Para fijar las ideas se trata el caso de un semiconductor extrínseco tipo n en el que únicamente existan impurezas donadoras. Las conclusiones serán trasladables simétricamente a los semiconductores extrínsecos tipo p con sólo impurezas aceptadoras. El caso de semiconductores extrínsecos parcialmente compensados es, en algunos aspectos, diferente.

Los rangos de concentraciones de impurezas N_D se limitan al caso en que no tiene lugar la superposición de órbitas y, por tanto, los estados de impureza se agrupan en un solo nivel E_D separado de la banda de conducción. Para el silicio esto corresponde aproximadamente a concentraciones $N_D < 10^{19} \text{ cm}^{-3}$.

La ecuación de neutralidad de carga es entonces

$$p + N_D^+ = n \tag{3.12}$$

y expresada en función de E_F :

$$M_v \exp \frac{E_v - E_F}{kT} + N_D \frac{1}{1 + \exp \frac{E_F - E_D}{kT}} = M_c \exp \frac{E_F - E_c}{kT} \quad (3.13)$$

o bien:

$$n_i \exp \frac{E_i - E_F}{kT} + N_D \frac{1}{1 + \exp \frac{E_F - E_D}{kT}} = M_c \exp \frac{E_F - E_i}{kT} \quad (3.14)$$

Se pueden establecer distintos rangos de temperaturas donde el comportamiento es diferente:

- Temperaturas bajas. Son las que están próximas a 0°K. Algunas características son similares a las de ésta última. A 0°K la concentración intrínseca es nula y, por tanto, también n y p . Todas las impurezas donadoras deben de estar ocupadas, es decir no ionizadas ($N_D^+ = 0$), y el nivel de Fermi se encontrará por encima del nivel donador E_D y por debajo de el mínimo de la banda de conducción E_c (por tratarse de un semiconductor no degenerado). Al aumentar ligeramente la temperatura aumenta la probabilidad de ionización de las impurezas donadoras y, por tanto, n . Como el nivel de Fermi se encuentra muy alejado de la banda de valencia la concentración de huecos p es mucho menor que n , creciendo mucho más lentamente que ella y la concentración intrínseca lo hace con una rapidez intermedia entre las de n y p . El semiconductor es extrínseco de tipo n y la progresividad de esta situación implicará a la larga el acercamiento del nivel de Fermi al de impurezas para aumentar las concentraciones N_D^+ y n . Por temperaturas bajas entenderemos, por tanto, aquellas para las que el nivel de Fermi se encuentra por encima de E_D .
- Temperaturas intermedias. Si la progresión indicada en el caso anterior continúa, al crecer la temperatura, el nivel de Fermi atravesará la posición del nivel donador E_D Para situarse por debajo de él. El crecimiento de N_D^+ se hará más lento teniendo como límite el valor de N_D , y el de n_i incidirá fundamentalmente en el de la concentración de huecos, mientras la temperatura sea suficientemente baja como para que la concentración intrínseca no sea comparable a la de impurezas. Estas consideraciones son compatibles con la idea de que, a temperatura suficiente, las vibraciones térmicas de la red son fácilmente capaces de aportar la pequeña energía necesaria para la ionización de la práctica totalidad de las impurezas. El rango de temperaturas intermedias corresponderá, pues, a un nivel de Fermi situado por

debajo del nivel de impureza y por encima del nivel intrínseco y, por tanto, a un semiconductor extrínseco de tipo n .

- Temperaturas altas. Serán aquéllas para las que la concentración intrínseca haya crecido lo suficiente como para superar a la concentración de impurezas. Esta última adquiere cada vez menos importancia al aumentar la temperatura y, aunque todas las impurezas se encontrarán ionizadas, es sobre todo la variación de la concentración intrínseca la que da lugar a la de las concentraciones de electrones y huecos. El semiconductor tenderá por tanto a un comportamiento intrínseco a pesar de la presencia de impurezas.

Por otra parte en este *applet* se incluye la posibilidad de someter a la muestra de semiconductor a una iluminación externa. Cuando un semiconductor se encuentra sometido a alguna acción exterior se dice que está fuera del equilibrio. A esta acción exterior el material responde mediante una alteración de sus distribuciones internas de carga, potencial o campo eléctrico lo que en muchos casos se traduce en la aparición de una corriente eléctrica que lo atraviesa y que se manifiesta hacia el exterior a través de los contactos. Si la acción exterior y, por tanto, la respuesta del material son constantes en el tiempo se encontrará en un régimen permanente o estacionario. En caso contrario el régimen será dinámico. El de equilibrio es un caso particular de régimen permanente en que las acciones exteriores son nulas.

En un semiconductor en equilibrio las concentraciones de electrones y de huecos están relacionadas entre sí de modo que no sólo se encuentran en equilibrio los electrones entre sí y los huecos entre sí sino también el gas de electrones con el de huecos. Ya se ha establecido el concepto de nivel de Fermi, cuya posición relativa dentro del diagrama de bandas se relaciona directamente con las concentraciones de ambos gases. Desde un punto de vista estrictamente matemático la introducción del concepto de nivel de Fermi puede ser considerada como un útil cambio de variables que permite manejar magnitudes que pueden variar dentro de amplios márgenes (las concentraciones de electrones y huecos) por medio de una sola magnitud (el nivel de Fermi) cuya variación relativa es mucho más débil. No obstante el concepto de nivel de Fermi va mucho más lejos de lo que se deduce del puro aspecto matemático.

En desequilibrio las distribuciones de concentración de portadores son casi siempre distintas de las de equilibrio y, por otra parte, la ley de masas no se cumple ya necesariamente. En otras palabras, el desequilibrio puede darse en uno de los dos gases de

portadores o en la relación entre ambos. En principio las concentraciones de portadores pueden variar también dentro de amplios márgenes y, por paralelismo con el caso de equilibrio, se pueden definir magnitudes energéticas similares al nivel de Fermi cuyos márgenes de variación sean más reducidos. En general y dado que normalmente no se verificará la ley de masas se define una magnitud distinta para cada uno de ambos gases de portadores. Estas magnitudes se denominan pseudoniveles de Fermi de electrones E_{FE} y de huecos E_{FH} y quedan definidas matemáticamente como:

$$E_{FE} = E_i + kT \ln(n/n_i) \quad (3.15)$$

$$E_{FH} = E_i - kT \ln(p/n_i) \quad (3.16)$$

o, lo que es equivalente, las concentraciones n y p en desequilibrio se expresarán mediante:

$$n = n_i \exp(E_{FE} - E_i)/kT \quad (3.17)$$

$$p = n_i \exp(E_i - E_{FH})/kT \quad (3.18)$$

A diferencia del nivel de Fermi de equilibrio que es constante, los pseudoniveles E_{FE} y E_{FH} pueden ser espacialmente variables. El equivalente a la ley de masas se escribe ahora como producto de estas dos últimas ecuaciones en la forma:

$$np = n_i^2 \exp(E_{FE} - E_{FH})/kT \quad (3.19)$$

y puesto que en general E_{FE} y E_{FH} no tendrán porqué variar paralelamente el producto np en desequilibrio no será una magnitud constante.

Las expresiones de E_{FE} y E_{FH} pueden considerarse de validez general (en el caso de semiconductores no degenerados) que incluyen el caso particular de equilibrio en el que, por definición:

$$E_{FE} = E_{FH} = E_F = cte \quad (3.20)$$

El lenguaje de energías puede ser sustituido por el de potenciales dividiendo aquéllas por la carga del electrón ($-e = 1.6 \cdot 10^{-19}$ C). En tal caso puede hablarse de pseudopotenciales de Fermi $\Phi_e = -E_{FE}/e$ y $\Phi_h = -E_{FH}/e$ con lo que se obtienen las siguientes relaciones:

$$n = n_i \exp(j - \Phi_e)e/kT \quad (3.21)$$

$$p = n_i \exp(\Phi_h - j)e/kT \quad (3.22)$$

siendo $j = -E_i/e$ (o $\nabla j = -\nabla E_i/e = -\vec{e}$) el potencial electrostático.

Al producirse una generación neta de portadores en un semiconductor, como consecuencia de una acción exterior, el semiconductor reacciona oponiéndose a dicha generación mediante una recombinación neta tanto más importante cuanto mayor sea la

generación debida a la causa externa. Cuando ambas tendencias se contrarrestan mutuamente se alcanza un régimen estacionario, siendo el de equilibrio un caso particular de régimen estacionario en el que no existe acción exterior alguna. Si la acción exterior cambia con el tiempo ello dará lugar a un régimen dinámico en el que predomina una de las dos tendencias de acción o reacción. En el caso particular de una variación instantánea de la acción externa se produce un régimen transitorio entre estados estacionarios.

Las ecuaciones que rigen el proceso dinámico general se denominan ecuaciones de continuidad de las concentraciones de portadores. Mediante su resolución, se determina cómo tiene lugar la tendencia de retomo al equilibrio (es decir, cómo se establece la reacción frente a una acción externa) en el tiempo y en el espacio y cuáles son los parámetros básicos que caracterizan el proceso.

La variación con el tiempo de la concentración c de portadores, de uno u otro tipo, en un elemento de volumen puede deberse a dos causas distintas: la generación o recombinación netas en el citado elemento de volumen y el transporte o flujo variable de portadores a través de él.

La primera causa se expresa en la forma:

$$\left(\frac{dc}{dt} \right)_{g-r} = G - U \quad (3.23)$$

donde G representa la velocidad de generación debida exclusivamente a acciones externas y U la velocidad de recombinación neta interna. En el caso más general (régimen dinámico) G y U pueden tomar valores diferentes para electrones y huecos.

Esta será la única causa de variación de portadores en el *applet* ya que se tiene una muestra de semiconductor homogénea y sin polarizar. Por ejemplo, sea una muestra de material semiconductor de dopaje uniforme, de tipo n , y homogéneamente iluminada mediante una radiación de energía suficiente ($E > E_G$) para poder crear pares electrón-hueco en su interior. El coeficiente de absorción α y el espesor de la muestra W se supondrán suficientemente pequeños de modo que $\alpha W \ll 1$. En estas condiciones puede suponerse que la luz es absorbida también homogéneamente en el material dando lugar a una velocidad de generación G_L uniforme en todo el volumen. De estas hipótesis es inmediato deducir la ausencia de variación espacial de cualquier magnitud relacionada con las concentraciones de portadores, y en las ecuaciones de continuidad los términos de divergencia de densidades de corriente desaparecen. Como hipótesis complementarias, que contribuyen a simplificar aún más el problema, se añaden las de máxima eficacia de los centros de recombinación e inyección débil en todo momento. Con ello podemos asegurar

que la generación y la recombinación se producen siempre por pares electrón-hueco (factor de ocupación constante de los niveles de recombinación) y por tanto que las ecuaciones de continuidad para electrones y huecos son totalmente equivalentes. Para resolverlas utilizaremos la ecuación correspondiente a los portadores minoritarios (huecos en el caso supuesto), ya que en inyección débil la expresión de la velocidad de recombinación neta es muy simple en función del exceso de portadores minoritarios ($p_n - p_{neq}$), y del tiempo de vida de los huecos (t_h). Se tiene por tanto:

$$\frac{dp_n}{dt} = G_L - \frac{p_n - p_{neq}}{t_h} \quad (3.24)$$

Para el régimen estacionario ($dp_n/dt = 0$) con iluminación ($G_L > 0$) la solución se obtiene directamente:

$$p_{nest} - p_{neq} = p'_L (= n'_L) = G_L t_h \quad (3.25)$$

expresión de la que, además, es inmediato deducir sus condiciones de validez como solución correspondiente a un régimen de inyección débil:

$$(p'_L)_{iny. débil} < n_{neq} \Rightarrow G_L t_h \ll n_{neq} \quad (3.26)$$

La evolución dinámica entre la situación de equilibrio y la estacionaria, a partir del instante en que se establece la iluminación de la muestra ($t = 0$) se obtiene como solución de la ecuación 3.24 con la condición inicial $p_n(0) = p_{neq}$ resultando en una evolución exponencial cuyo límite es el de la ecuación 3.25 al cabo de un tiempo, en teoría, infinitamente largo, aunque, a efectos prácticos, de tan sólo un múltiplo pequeño de t_h (teniendo en cuenta que el orden de magnitud típico de t_h es de μseg).

Si una vez alcanzado el régimen estacionario se apaga la fuente luminosa, es decir se hace $G_L = 0$ bruscamente, el exceso existente de portadores sobre la concentración de equilibrio no puede desaparecer en tiempo nulo, debiendo hacerlo por vía de la recombinación.

La ecuación 3.24 con $G_L = 0$ y la condición inicial $p_n(0) = p_{neq} + p'_L$ muestra que el retorno al equilibrio se realiza durante un tiempo del orden de t_h .

El parámetro fundamental que rige el proceso es la constante de tiempo t_h (t_e para un semiconductor de tipo p) que recibe el nombre de tiempo de vida de los portadores minoritarios.

3.1.1. DESCRIPCIÓN DEL *APPLET*

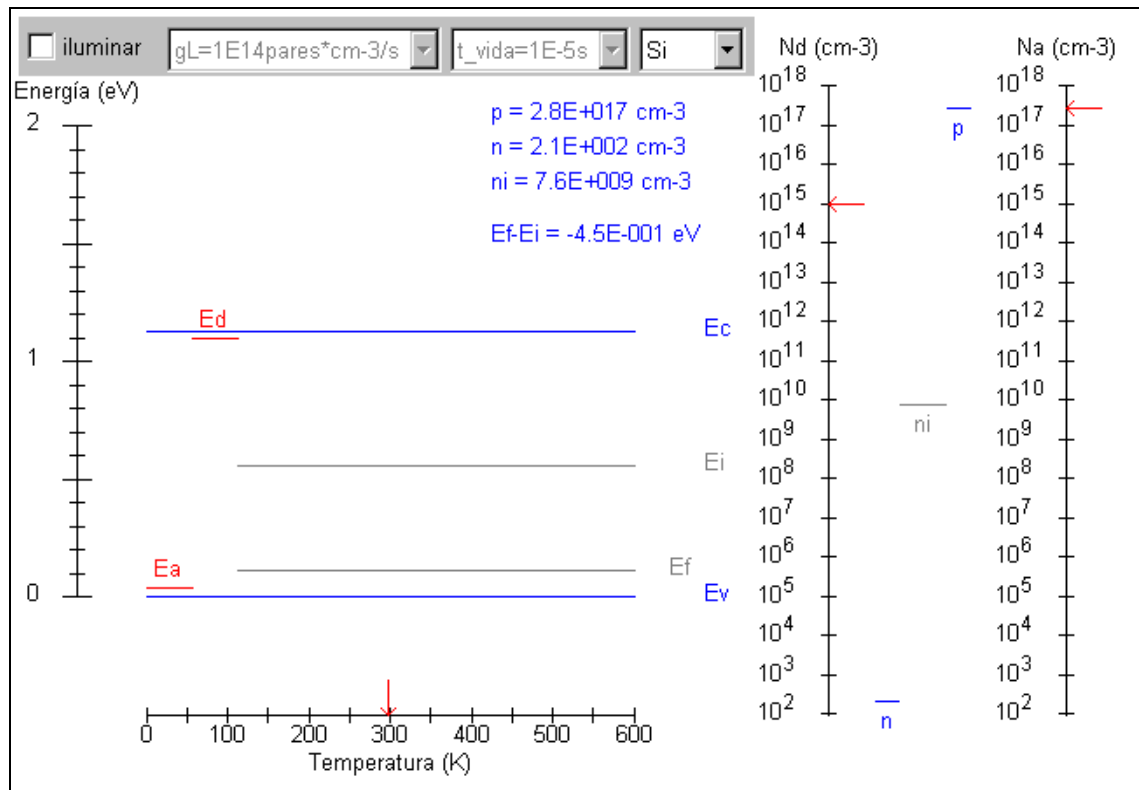


Figura 3.1. Applet de Evolución del nivel de Fermi

El *applet* ilustra la variación del nivel de Fermi, así como de la concentración intrínseca, de un material semiconductor según la temperatura, la posible presencia de luz, el tipo de material semiconductor (banda prohibida, número de estados equivalentes en banda de conducción y de valencia), el dopaje y el tipo de impurezas (que determinan la posición del nivel permitido dentro de la banda prohibida y, por tanto, el grado de ionización).

El objetivo de este *applet* es familiarizar al estudiante, tanto cualitativamente como cuantitativamente, con la relación entre la posición relativa del nivel de Fermi dentro del ancho de banda prohibido de un semiconductor y las concentraciones de huecos y electrones en las bandas. También es posible elegir entre diferentes materiales semiconductores.

La posición dentro de la banda prohibida juega un papel fundamental al controlar las propiedades electrónicas y las características del dispositivo semiconductor. No obstante, usando los medios tradicionales como las imágenes estáticas, ecuaciones y libros, es difícil proporcionar al alumno la habilidad de realizar una asociación inmediata de la posición del nivel de Fermi con las concentraciones de portadores y viceversa.

En la Figura 3.1 se observan algunas cantidades en rojo - **Na** (impurezas aceptoras), **Nd** (impurezas donadoras), **T** (temperatura), **Ea** (nivel de energía aceptor), y **Ed** (nivel de energía donador) – que pueden variarse arrastrando la marca correspondiente. Con esos cambios, el *applet* calcula n (concentración de electrones), p (concentración de huecos), n_i (concentración intrínseca), E_i (nivel de energía intrínseco) y E_f (nivel de Fermi).

En lugar de simular con precisión la dependencia con la temperatura, el programa muestra la dependencia que proporcionan los mecanismos estadísticos. Es decir, la temperatura depende solamente de la distribución de Fermi-Dirac. Se ignora la variación del ancho de banda prohibido y de otros parámetros del material. La temperatura no podrá ser muy baja (25°K). Esto se debe a que los términos como $\exp(E/kT)$ se hacen muy grandes o muy pequeños para valores bajos de T , haciendo el análisis numérico más difícil. El rango de temperatura de disponible es lo suficientemente grande como para ver los aspectos más importantes.

Siguiendo con la descripción del mismo se puede ver en la parte superior un panel de control donde se puede elegir el tipo de semiconductor entre Si, Ge y AsGa.

También está el cuadro de selección de iluminación que hace las veces de interruptor de la luz. Cuando se selecciona la opción de iluminar el semiconductor, automáticamente se hacen accesibles las listas de selección que permiten variar la tasa de generación y el tiempo de vida de los portadores minoritarios. Y ya habrán aparecido los pseudoniveles de Fermi en el diagrama de bandas del semiconductor. Por lo demás el funcionamiento no puede ser más sencillo e intuitivo y queda a disposición del usuario el jugar con las distintas posibilidades que ofrece.

A modo de resumen se van a incluir unas tablas que agrupan los controles y las variables del *applet*. En la tabla 3.1 están incluidos todos los controles:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
N_d	Flecha móvil	$10^2 \leftrightarrow 10^{18}$	cm^{-3}	Impurezas donadoras
N_a	Flecha móvil	$10^2 \leftrightarrow 10^{18}$	cm^{-3}	Impurezas aceptoras
T	Flecha móvil	$25 \leftrightarrow 600$	$^{\circ}\text{K}$	Temperatura
N_d	Flecha móvil	$10^2 \leftrightarrow 10^{18}$	cm^{-3}	Impurezas donadoras
E_d	Flecha móvil	Banda prohibida	eV	Nivel de energía donador
E_a	Flecha móvil	Banda prohibida	eV	Nivel de energía aceptor
<i>Semiconductor</i>	Botón de selección (<i>choice</i>)	Si-Ge-GaAs		Tipo de semiconductor
<i>iluminar</i>	Botón de comprobación (<i>checkbox</i>)			Presencia de luz
g_L	<i>Choice</i>	$10^{14} \leftrightarrow 10^{18}$	$\text{pares}\cdot\text{cm}^{-3}/\text{s}$	Tasa de generación
t_{vida}	<i>Choice</i>	$10^{-7} \leftrightarrow 10^{-5}$	s	Tiempo de vida

Tabla 3.1. Controles del *applet* de evolución del nivel de Fermi

La tabla 3.2 mostrará todas las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
n	Concentración de electrones	cm^{-3}
p	Concentración de huecos	cm^{-3}
n_i	Concentración intrínseca	cm^{-3}
E_i	Nivel de energía intrínseco	eV
E_f	Nivel de Fermi	eV
E_e	Pseudonivel de Fermi de electrones	eV
E_h	Pseudonivel de Fermi de huecos	eV

Tabla 3.2. Variables del *applet* de evolución del nivel de Fermi

3.2. SEMICONDUCTOR HOMOPOLAR EN EQUILIBRIO CON DISTRIBUCIÓN NO UNIFORME DE IMPUREZAS

Gran parte de los dispositivos semiconductores deben sus propiedades electrónicas a la distribución no homogénea de impurezas en su interior. Incluso puede decirse que la base de cualquier dispositivo electrónico radica en algún tipo de inhomogeneidad (de constitución o de comportamiento) puesto que el concepto de inhomogeneidad implica el de la posibilidad de su control. Más aún, cualquier fenómeno que lleve consigo variaciones, diferencias, inhomogeneidades, etc, desde un punto de vista eléctrico da o puede, en potencia, dar lugar a un dispositivo electrónico. Así, en semiconductores homogéneos, el efecto Hall (diferencia de comportamiento de impurezas ionizadas y portadores) o la variación de movilidad, por poner algún ejemplo derivado de fenómenos anteriormente estudiados, pueden ser el origen de dispositivos electrónicos.

Nos referimos ahora exclusivamente a las inhomogeneidades de dopaje de impurezas en cristales semiconductores. En el *applet* anterior han sido tratadas las propiedades estadísticas de los semiconductores extrínsecos homogéneos. Ahora se emprende el análisis de los casos de dopaje no homogéneo. El estudio de la situación de equilibrio térmico es a la vez el más simple y el punto de partida básico para el análisis de las pequeñas desviaciones respecto de él.

Aunque el concepto de inhomogeneidad es muy amplio, generalmente la noción de semiconductor no homogéneo suele restringirse a aquel que presenta una distribución de impurezas, donadoras o aceptadoras, variable punto a punto (en el sentido macroscópico).

A temperaturas muy bajas o muy altas la distribución inhomogénea de impurezas tiene escasa incidencia sobre el comportamiento eléctrico del semiconductor. En el primer caso porque la casi totalidad de las impurezas serán neutras (no ionizadas) y prácticamente existirá una neutralidad de carga local. En el segundo porque las concentraciones de electrones y huecos son prácticamente iguales a la intrínseca y muy superiores a la de impurezas y la densidad de carga será también nula en todo punto.

La situación es muy distinta en el rango de temperaturas medias. En este caso todas las impurezas se encuentran ionizadas y el semiconductor es extrínseco. Los portadores mayoritarios tenderán por un lado, debido a la atracción eléctrica, hacia una distribución inhomogénea semejante a la de las impurezas ionizadas correspondientes, de carga

opuesta. Por otro lado, a causa de la difusión (que a estas temperaturas afecta a los portadores pero no a las impurezas, que permanecen en posiciones fijas) los portadores tenderán hacia la homogeneidad de su concentración.

De la igualación, característica del equilibrio, de estas tendencias contrapuestas se deduce la necesaria presencia de campos eléctricos internos, variaciones de potencial y distribuciones de carga de alcance macroscópico en los semiconductores no homogéneos.

Por su propia definición el equilibrio se caracteriza por la nulidad del flujo neto de partículas en cualquier dirección. Por ello las dos tendencias anteriormente citadas deberán, en equilibrio, contrarrestarse mutuamente en todo punto. En semiconductores dopados homogéneamente el equilibrio es compatible tanto con la neutralidad de carga local (campo eléctrico nulo en todo punto) como con la homogeneidad en la distribución de concentraciones de portadores (flujo de difusión nulo en cada punto). La condición de nulidad de flujo neto se reduce entonces a la de nulidad de cada uno de los flujos componentes de arrastre y difusión.

Por el contrario, en semiconductores con dopaje no homogéneo el equilibrio no puede corresponder ni a una densidad de carga nula, que significaría una distribución de portadores mayoritarios semejante a la de impurezas, porque daría lugar a un campo eléctrico nulo (flujo de arrastre nulo) y a una difusión neta no nula; ni a una distribución homogénea de portadores con difusión nula y densidad de carga, campo eléctrico y flujo de arrastre no nulos. La distribución de portadores mayoritarios será pues intermedia entre la homogénea y la de impurezas de tal modo que en cada punto exista una compensación entre las tendencias de arrastre y difusión.

El equilibrio se caracterizará, por tanto, en el caso general por la presencia de una densidad de carga neta producida por descompensación, en cada volumen infinitesimal, entre la carga debida a impurezas ionizadas fijas y la debida a los portadores móviles. Esta densidad de carga dará lugar a una distribución (macroscópica) de campo y potencial eléctricos (superpuestas a sus distribuciones periódicas cuyos efectos se incluyen en el concepto de masa efectiva) y por tanto a un flujo de arrastre de portadores que, en cada punto, será compensado por un flujo igual y de sentido contrario producido por difusión.

En el caso de uniones homopolares la densidad de carga es debida fundamentalmente a la descompensación local entre las concentraciones de impurezas y de portadores mayoritarios siendo a estos efectos despreciable la concentración de portadores minoritarios en todo punto. Es de esperar por ello valores poco importantes tanto del

campo eléctrico como de la densidad de carga y desviaciones poco notables respecto de la condición de neutralidad de carga.

Utilizamos un modo práctico para calcular la distribución de campo eléctrico y densidad de carga en esta unión homopolar unidimensional que consiste en considerar que hay cuasi-neutralidad de carga ($n(x) \approx N_D(x)$) si la unión homopolar es del tipo $n^+ - n$.

El campo eléctrico se obtendría mediante:

$$E(x) = -\frac{D_e}{m_e} \frac{dn/dx}{n} \approx -\frac{D_e}{m_e} \frac{dN_D(x)/dx}{N_D(x)} \quad (3.27)$$

de cuya variación puede deducirse la distribución aproximada de $\rho(x)$ mediante

$$\frac{dE_x}{dx} = \frac{r}{e} \quad (3.28)$$

y mediante

$$r = e(N_d + p - n) \quad (3.29)$$

donde se ha designado por N_d la concentración neta de impurezas ionizadas:

$$N_d = N_D^+ - N_A^- \quad (3.30)$$

en combinación con la ley de acción de masas, una mejor aproximación para $n(x)$. Generalmente este proceso es suficiente ya que posteriores iteraciones siguiendo el mismo esquema de cálculo suponen muy ligeras modificaciones.

3.2.1. DESCRIPCIÓN DEL *APPLET*

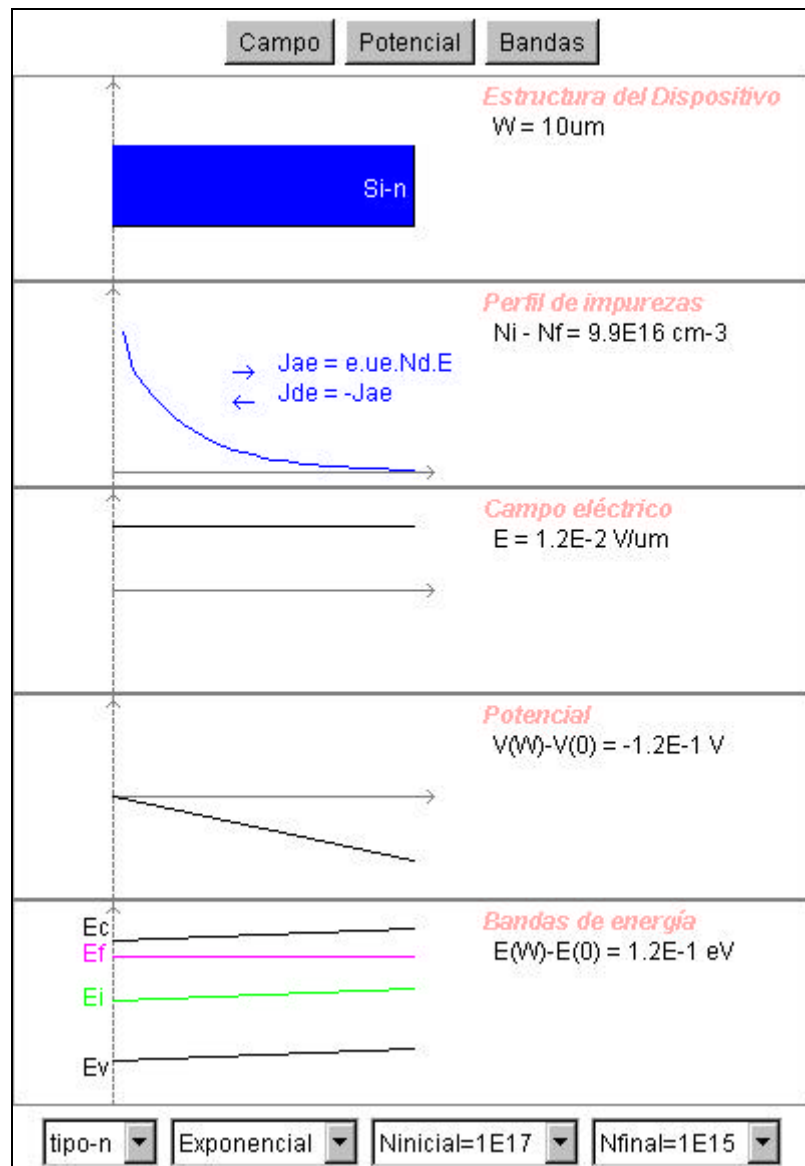


Figura 3.2. Distribución no uniforme de impurezas

En este *applet* que aparece representado en la Figura 3.2 podemos ver un conjunto de 5 diagramas:

En el primero de ellos tenemos el esquema del dispositivo donde se aprecia el tocho de semiconductor cuyo color dependerá del tipo de impurezas con las que está dopado. A continuación tenemos el perfil de impurezas con las que dopamos el silicio. Este perfil lo podemos variar mediante el panel de control que tenemos al final del *applet*. En el diagrama tenemos información sobre las características seleccionadas para el perfil de dopaje. Además podemos observar los procesos de difusión y arrastre mediante las flechas

que indican el sentido de las corrientes producidas por dichos mecanismos. El tercer diagrama tiene representado el campo eléctrico en el interior del semiconductor, así como la información correspondiente al valor máximo de campo. Justo debajo podemos apreciar el potencial eléctrico correspondiente al perfil de campo anterior y por último tenemos el diagrama de bandas de energía donde vemos en color verde en nivel intrínseco y en color violeta en nivel de Fermi.

Con los botones de la parte superior podemos ocultar o mostrar los esquemas correspondientes al campo, potencial eléctrico y al diagrama de bandas de energía. De este modo se puede prestar mayor atención a un determinado aspecto del dispositivo.

En el panel de control inferior podemos seleccionar el tipo de semiconductor y el perfil que tendrán las impurezas dentro del mismo. Podemos elegir entre un perfil exponencial o lineal. Los dos controles del final indican la concentración de impurezas en los bordes del semiconductor. El primero para la parte inicial y el segundo para el final. Con la característica de que el perfil siempre será decreciente debido a la simetría del dispositivo.

En la tabla 3.3 están incluidos todos los controles:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
<i>Campo</i>	Botón			Campo eléctrico
<i>Potencial</i>	Botón			Potencial eléctrico
<i>Bandas</i>	Botón			Bandas de energía
<i>Semiconductor</i>	<i>Choice</i>	Tipo n, tipo p		Tipo de semiconductor
<i>Perfil</i>	<i>Choice</i>	Lineal, exponencial		Perfil de dopaje
<i>Ninicial</i>	<i>Choice</i>	$5 \cdot 10^{14} \leftrightarrow 5 \cdot 10^{17}$	cm^{-3}	Concentración de impurezas inicial
<i>Nfinal</i>	<i>Choice</i>	$5 \cdot 10^{14} \leftrightarrow \text{Ninicial}$	cm^{-3}	Concentración de impurezas final

Tabla 3.3. Controles del *applet* de Distribución no uniforme de impurezas

La tabla 3.4 mostrará todas las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
E	Valor del campo eléctrico	$\text{V}/\mu\text{m}$
V	Valor del potencial eléctrico	V
E_h	Pseudonivel de Fermi de huecos	eV
E_i	Nivel de energía intrínseco	eV
E_f	Nivel de Fermi	eV
E_c	Nivel de energía de la banda de conducción	eV
E_v	Nivel de energía de la banda de valencia	eV

Tabla 3.4. Variables del *applet* de Distribución no uniforme de impurezas

3.3. UNIÓN PN EN EQUILIBRIO Y POLARIZADA

Este *applet* ha sido desarrollado para ayudar al alumno a adquirir soltura para dibujar el diagrama de bandas de una unión PN, lo que se hace posible al trabajar con el programa y observar los gráficos animados. Con este *applet* también se ayuda a familiarizarse con los principios subyacentes de la formación del diagrama de bandas, así como con el perfil de carga y campo eléctrico que se produce.

El diagrama de bandas es fundamental al explicar los principios físicos de la operación de los dispositivos semiconductores. El diagrama de bandas de un diodo de unión PN es el bloque fundamental o el punto de partida para los diagramas de bandas de otras estructuras de dispositivos más complicadas como el transistor bipolar. Usando imágenes fijas o matemáticas o los principios físicos solamente, los estudiantes tienen dificultad al aprender a dibujar por su cuenta el diagrama de banda en equilibrio.

En el *applet* se muestra el flujo de las corrientes de difusión, las corrientes inversas de arrastre, y la variación del diagrama de bandas cuando cambia la polarización. Sólo se tienen en cuenta las características del diodo ideal (esto es, no se consideran la generación o recombinación en la zona de depleción ni el efecto del alto nivel de inyección – cuando la concentración de portadores inyectados es comparable a la concentración de portadores mayoritarios por el dopaje, etc. No se muestra la difusión de los portadores inyectados en la región neutra que está coloreada.)

Existe una dependencia de la concentración de portadores con la posición. La densidad de portadores minoritarios disminuye al adentrarse en la muestra, aproximándose al logaritmo de la misma en equilibrio cero y se mantiene constante. Esta dependencia con la posición se puede expresar como:

$$n(x) = n_{peq} \exp[qV/kT] \exp[-x/L] \quad (3.31)$$

donde n_{peq} es la concentración de portadores en $x=0$ en equilibrio. Con cierta polarización, podemos obtener:

$$n(x, V_0) = n_{peq} \exp[qV_0/kT] \exp[-x/L] = n_0' \exp[-x/L] \quad (3.32)$$

Notar que n_0' es una constante para una determinada tensión. Notar también que:

$$n(x_0) = \int_0^\infty n(E_c, x_0) \exp\left(-\frac{E - E_c}{kT}\right) dE$$

L es la longitud de difusión. Es la distancia media que viaja un minoritario antes de recombinarse con un mayoritario, típicamente toma valores de entre micras a milímetros. Se puede obtener su valor como:

$$L = \sqrt{\frac{kT\mu\tau}{q}}$$
 Siendo μ la movilidad de los portadores, τ el tiempo de vida de los

mismos.

Ha habido muchos intentos para medir los tiempos de vida de los portadores, así como las movilidades y longitud de difusión. Para una concentración de impurezas mayor de 10^{19}cm^{-3} , los experimentos son complicados, ya que las concentraciones de portadores minoritarios son demasiado pequeñas, y los resultados obtenidos no son precisos. De modo que con un propósito de modelado de dispositivos se utilizan unas ecuaciones empíricas para los electrones y huecos minoritarios,

$$\mu_n = 232 + 1180 / (1 + (N_A / (8 * 10^{16}))^{0.9})$$

$$\mu_p = 130 + 370 / (1 + (N_D / (8 * 10^{17}))^{1.25})$$

$$1/\tau_n = 3.45 * 10^{-12} * N_A + 0.95 * 10^{-31} * N_A^2$$

$$1/\tau_p = 7.8 * 10^{-13} * N_D + 1.8 * 10^{-31} * N_D^2$$

Los datos utilizados en el *applet* se basan en dichas ecuaciones. Se pueden agrupar en una tabla los resultados que se obtienen para la longitud de difusión de los electrones y huecos minoritarios:

Dopaje, lado-p	$L_n(\mu m)$	Dopaje, lado-n	$L_p(\mu m)$
$N_A = 10^{15}\text{cm}^{-3}$	1023	$N_D = 10^{15}\text{cm}^{-3}$	1290
$N_A = 10^{16}\text{cm}^{-3}$	307	$N_D = 10^{16}\text{cm}^{-3}$	407
$N_A = 10^{17}\text{cm}^{-3}$	76	$N_D = 10^{17}\text{cm}^{-3}$	124
$N_A = 10^{18}\text{cm}^{-3}$	16	$N_D = 10^{18}\text{cm}^{-3}$	28
$N_A = 10^{19}\text{cm}^{-3}$	3.8183	$N_D = 10^{19}\text{cm}^{-3}$	3.8210

Se puede observar por los resultados obtenidos, que la longitud de difusión disminuye al aumentar el nivel de dopaje. Cuando se tiene el mismo nivel de dopaje en 10^{19}cm^{-3} , la longitud de difusión de los huecos minoritarios es mayor que la de los electrones.

Si la polarización es nula, el número de electrones que se mueven, por unidad de tiempo, desde el lado-p-al lado-n es exactamente igual al número moviéndose en la dirección opuesta (es decir, desde el lado-n-hasta el lado-p). Esto también es cierto para los huecos. Entonces la corriente total a través de la unión es cero para polarización nula (como debería ser).

3.3.1. DESCRIPCIÓN DEL APLET

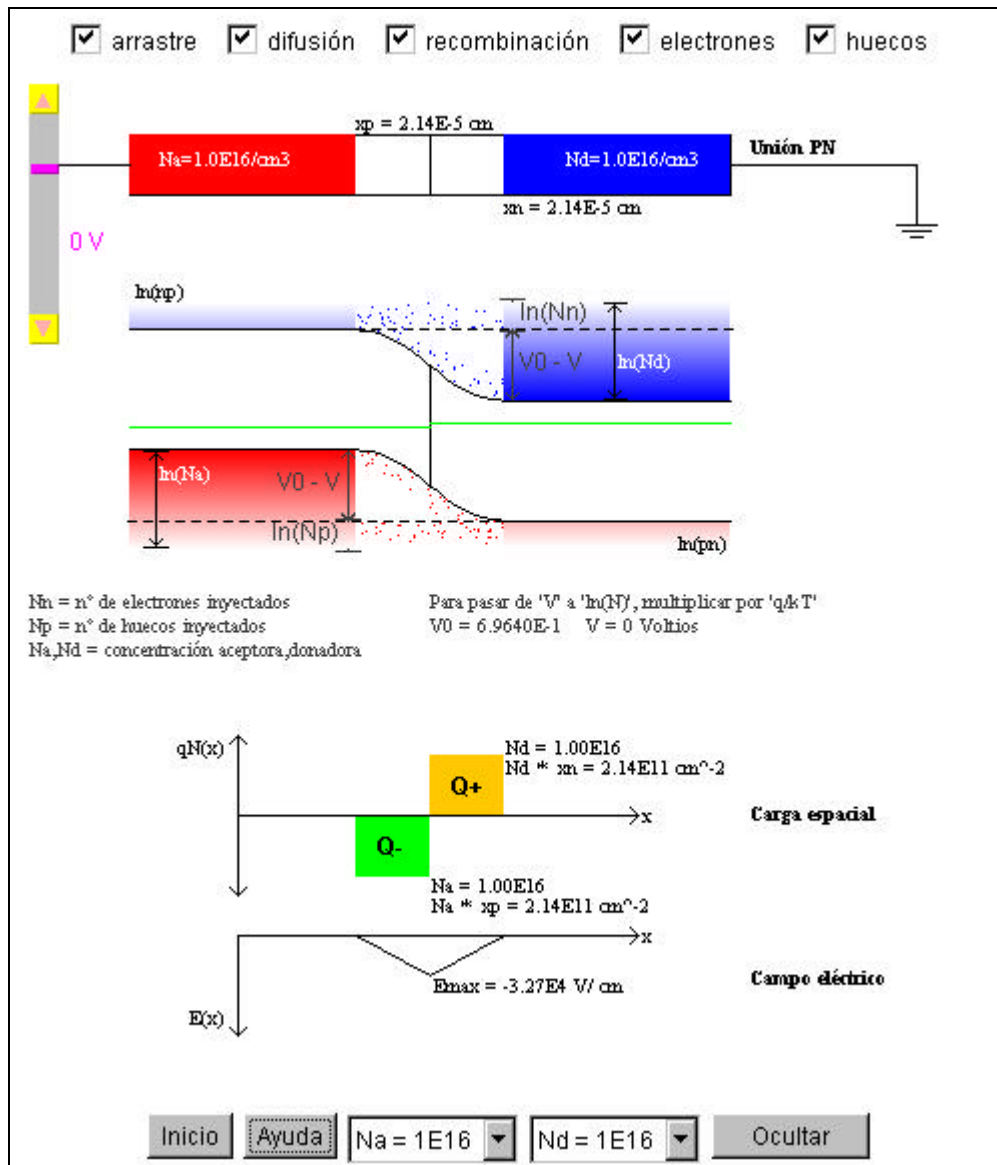


Figura 3.3. Unión PN en equilibrio

En la Figura 3.3, el esquema de la muestra (rectángulo rojo-blanco-azul en la parte superior) y el diagrama de bandas están dados para una polarización determinada que se ajustará utilizando la barra de desplazamiento de la esquina superior izquierda. El diagrama de bandas muestra tanto los portadores mayoritarios como los minoritarios (electrones = rectángulo azul; huecos = rectángulo rojo). Las alturas de los rectángulos en las bandas de conducción/valencia en el lado n/p respectivamente son proporcionales al logaritmo de la densidad de portadores mayoritarios. Las alturas de los rectángulos en las bandas de conducción/valencia en el lado p/n varían en función de la posición dentro de la

muestra. Justo en el borde de la zona de depleción en el lado p (n), la altura de los rectángulos es igual a la altura del logaritmo de la densidad de portadores mayoritarios por encima de la banda de conducción (valencia) en el lado n (p). La densidad de portadores va disminuyendo entonces conforme se adentra en la muestra, al final se aproximará al logaritmo de la densidad de portadores minoritarios en equilibrio y permanece constante. La longitud para la cual la densidad de portadores alcanza el valor constante es igual a la longitud para la que existe corriente de recombinación. La variación de la densidad de portadores en este proceso se produce según un descenso de tipo exponencial con un parámetro denominado longitud de difusión. La longitud de difusión es la distancia media que recorre un portador minoritario antes de recombinarse con un portador mayoritario, típicamente es de unas pocas micras o milímetros.

La escala vertical del diagrama de bandas de energía se utiliza tanto para energía como para potencial y la intensidad del color del rectángulo azul en la banda de conducción y del rectángulo rojo en la banda de valencia es proporcional al logaritmo de la concentración de portadores a esa energía (escala vertical en el diagrama de bandas).

La densidad de electrones en función de la energía, $n(E)$, es proporcional al término $\exp[-(E-E_c)/kT]$ en la banda de conducción, y la densidad de huecos $p(E)$ es proporcional a $\exp[(E-E_v)/kT]$ en la banda de valencia. Por lo tanto, $\ln[n(E)]$ disminuye linealmente con $(E-E_c)$ en la banda de conducción, y $\ln[p(E)]$ lo hace con (E_v-E) en la banda de valencia. Por eso en el *applet* el nivel de color en la banda varía linealmente con la separación de energía desde el borde de la banda, representando la variación de la densidad con la energía, $\ln[n(E)]$ frente a E .

El número de puntos azules (electrones) en E (es decir, coordenada- y), que se originan desde el rectángulo azul en la banda de conducción de la región neutra, es proporcional a $\ln[n(E)]$. En realidad, el número debería ser directamente proporcional a $n(E)$.

También se muestra el perfil de la densidad de carga especial y el campo eléctrico en la zona de depleción del diodo de unión-pn bajo la polarización aplicada.

Para aprender a utilizar el programa se empieza utilizando los controles de la parte superior. Con ellos el usuario puede familiarizarse con las cuatro componentes básicas de corriente, para lo cual se deben realizar las siguientes acciones:

- Corriente de arrastre de electrones: ésta es parte de la corriente de saturación inversa.
 1. Desconectar todos los recuadros excepto los de arrastre y electrones.

2. Observar los puntos azules que bajan por la barrera de potencial, son los electrones de arrastre.
- Corriente de arrastre de huecos: es la otra parte de la corriente de saturación inversa.
 1. Repetir lo mismo para arrastre y huecos,
 2. Observar los puntos rojos que suben la barrera de potencial, son los huecos de arrastre. Hay que hacer notar que en el diagrama de bandas de energía, la energía de los electrones aumenta conforme se sube. Pero la de los huecos lo hace al bajar en la banda. Por consiguiente, ambas corrientes de arrastre (electrones y huecos) fluyen en la dirección de descenso de la energía.
 - Corriente de difusión de electrones: una vez inyectados (dentro de la región-p, roja), los electrones se mueven por medio de la difusión (debido al gradiente de concentración en la región-p.)
 1. Para difusión y electrones,
 2. observar los puntos azules que se mueven horizontalmente del lado-n al lado-p , son los electrones inyectados.
 - Corriente de difusión de huecos: después de la inyección (dentro de la región-n neutra, azul), se mueven (hacia fuera de la unión) por difusión.
 1. Para difusión y huecos,
 2. Observar los puntos rojos que fluyen en la banda de valencia horizontalmente. Se mueven del lado-p al lado-n.
 - Corriente de recombinación de electrones: cuando la condición de equilibrio térmico de un sistema físico es distribuida, existen procesos para devolver el sistema al equilibrio. En la unión pn, el mecanismo existente es la recombinación banda a banda donde un par electrón-hueco se recombinan. Esta transición de un electrón desde la banda de conducción a la banda de valencia se hace posible por la emisión de un fotón o mediante transferencia de energía a otro electrón o hueco libre. Estos electrones forman la corriente de recombinación de electrones.
 1. Desconectar todos los recuadros excepto los de recombinación y electrones.
 2. Observar los puntos azules que fluyen desde la banda de conducción a la de valencia, verticalmente.
 3. Variar el nivel de dopaje de N_A , observar qué cambios se producen respecto a la altura inicial de la densidad de portadores minoritarios, la longitud de difusión, y

la longitud donde puede haber recombinación. Notar cómo disminuyen las dos últimas al aumentar el nivel de dopaje.

- Corriente de recombinación de huecos: es la otra parte de la corriente de recombinación.
 1. Desconectar todos los recuadros excepto los de recombinación y huecos.
 2. Observar los puntos rojos que fluyen desde la banda de valencia a la de conducción, verticalmente.
 3. Variar el nivel de dopaje de Nd, observar que cambios se producen respecto a la altura inicial de la densidad de portadores minoritarios, la longitud de difusión, y la longitud donde puede haber recombinación. Notar cómo disminuyen la longitud de difusión y la de la zona donde puede existir recombinación al aumentar el nivel de dopaje.

Haciendo clic en el botón de “Inicio” del panel de control inferior se anula la tensión de polarización aplicada, es decir tenemos 0 Voltios. En esta situación las magnitudes son tales que:

1. corriente de difusión de electrones = corriente de arrastre de electrones.
2. corriente de difusión de huecos = corriente de arrastre de huecos.
3. Con polarización nula, el número de electrones que se mueven desde el lado-p hasta el lado-n es exactamente igual al número moviéndose en sentido opuesto (es decir, desde el lado-n al lado-p). Esto también es válido para los huecos. Por lo tanto la corriente total que atraviesa la unión es cero para polarización nula (como debería ser).

Cuando se aplica un voltaje positivo en el lado-p respecto al lado-n (polarización directa), entonces más electrones en el lado-n están por encima de E_c del lado-p (esto es, el borde de la banda del otro lado de la unión) y estos electrones pueden atravesar la unión sin oposición de la barrera de potencial. Estos electrones constituyen una de las partes de la corriente de difusión, como ya se ha mencionado. (Por supuesto, incluso con polarización nula los electrones fluyen desde el lado-n al lado-p. Pero estos se cancelan por el flujo de electrones en el sentido opuesto.) Esos electrones en el lado-n que permanecen por debajo de E_c del lado opuesto no serán capaces de atravesar la unión y no contribuyen a ninguna corriente.

Cuando se aplica una polarización inversa (es decir, lado-p polarizado negativamente respecto al lado-n), el flujo de electrones desde el lado-n al lado-p disminuye y finalmente desaparece para una determinada tensión. Nota que, más allá de este valor de polarización inversa, la corriente es constante respecto a la tensión inversa aplicada.

Eso mismo es válido para la corriente de difusión de huecos bajo polarización directa o inversa.

Por otra parte, debajo de el diagrama de bandas del dispositivo se aprecia el esquema del diagrama de carga espacial. Se asume que tanto el lado-p como el lado-n están dopados uniformemente, con una aproximación de unión abrupta. Por lo tanto de la densidad de carga especial en la zona de depleción es uniforme: En color verde se muestra la densidad de impurezas aceptoras ionizadas (cargadas negativamente) y en naranja la densidad de impurezas donadoras polarizadas (positiva). En rojo tenemos la región-p neutra (el color rojo representa la presencia de huecos que compensan la carga de las impurezas aceptoras, siendo por tanto una región neutra), el azul representa la región neutra n (el azul representa la presencia de electrones móviles).

Al no ser nula la densidad de carga especial en la región de depleción tendremos un campo eléctrico. En la gráfica de la densidad de carga se observa que el campo eléctrico apunta desde la zona de carga positiva directo a la negativa, esto implica que apunta a la izquierda (valor negativo). El valor absoluto máximo del campo eléctrico corresponde al punto de la unión física (entre el lado-p y el lado-n).

Podemos variar la tensión de polarización y los niveles de dopaje para examinar los nuevos perfiles de carga especial y campo eléctrico.

Una vez familiarizados con el funcionamiento del *applet* se puede proceder a trabajar directamente con él. Se puede seguir un guión de ejemplo para comprobar las características del programa. Para ello se inicia por desconectar la corriente de huecos y la de difusión. Sólo quedará la corriente de arrastre de electrones. Cambiar la tensión de alimentación con la barra de desplazamiento situada a la izquierda a tal efecto. Se observa si permanece la corriente de arrastre de electrones constante.

La corriente de arrastre de electrones depende sólo de la concentración de electrones minoritarios en equilibrio del lado-p: $n_{peq} = ni^2 / Na$.

Desconectar todas las componentes de corriente excepto la corriente de difusión de electrones. Hay que hacer clic en el botón “Inicio” para poner a cero la polarización. Al variar la polarización se observan los cambios en el número de electrones inyectados. Para obtener unas útiles definiciones se pulsa el botón “Ayuda” que hay en el fondo del *applet*.

Bajo polarización directa, se puede obtener el número de electrones difundidos como sigue:

Usando la barra de la izquierda que permite fijar el valor de tensión, ponemos una polarización positiva, por ejemplo, 0.2 V.

El número de electrones disponibles para ser inyectados, N_n , es:

$$\ln(N_n) = \ln(N_d) - (q/kT) (V_0 - V) \quad (3.33)$$

Aquí N_d es el nivel de impurezas donadoras, y se multiplica por (q/kT) para pasar de tensión a concentración en la escala vertical del *applet*. (Notar que esto es cierto bajo ciertas suposiciones.)

Por tanto,

$$N_n = N_d \exp(-q(V_0 - V)/kT) \quad (3.34)$$

Para cero Voltios, $N_n = N_d \exp(-qV_0/kT)$. Esto debería ser igual a la concentración de electrones minoritarios del lado-p (se puede comprobar esto en el *applet* al pinchar en el botón “Inicio” y comparar $\ln(n_{peq})$ y $\ln(N_n)$). Es decir, $n_{peq} = N_d \exp(-qV_0/kT)$.

La corriente neta de electrones por lo tanto es proporcional a:

$$N_n - n_p = N_d \exp(-q(V_0 - V)/kT) - n_{peq} = N_d \exp(-qV_0/kT) (qV/kT - 1) \quad (3.35)$$

Ahora se desconectan todas las componentes de corriente excepto la de recombinación de electrones. Hay que hacer clic en el botón “Inicio” para poner a cero la polarización. Al variar la polarización se observan los cambios en el número de electrones recombinados. Cambiando el nivel de dopaje de Na, se observa lo que cambia la altura inicial de la densidad de portadores minoritarios, la longitud de difusión, y la longitud de la región donde puede haber recombinación. Hay que hacer notar que estas dos últimas disminuyen conforme aumenta el nivel de dopaje. Con el botón de “Ayuda” nuevamente se pueden ver algunas definiciones. Hacer lo mismo con las corrientes de huecos.

Al jugar con el valor de la tensión aplicada al diodo hay que destacar que el arrastre de portadores minoritarios se observa para todas las polarizaciones aplicadas, la corriente de difusión se incrementa bajo polarización directa (tensión positiva), y la corriente de recombinación también bajo polarización directa solamente.

Una vez analizado el completo funcionamiento del *applet* podemos pasar ahora a analizar una serie de cuestiones fundamentales en las que centrar la atención del alumno, y se verá cómo mediante este *applet*, se puede derivar la relación corriente-tensión para un diodo de unión pn:

1. Observar en qué dirección apunta el campo eléctrico en la región deplexión (zona blanca).
2. Con polarización nula (pinchando el botón “Inicio”), se deberían ver los electrones (puntos azules) atravesando la unión. ¿Es la corriente de arrastre de electrones (desde p a n) igual a la corriente de difusión de electrones (desde n a p) ? ¿Es cierto para los huecos también? Se observa cómo los electrones arrastrados bajan la barrera de potencial, mientras que los electrones difundidos se mueven horizontalmente. Notar el significado físico de estos flujos de electrones aparentemente diferentes.
3. Bajo polarización inversa (tensión negativa en el lado-p respecto al lado-n que está a masa), cambiando el nivel de dopaje del lado-p con el control del fondo se observa la corriente de arrastre de electrones. Explicar cualitativamente (basándose en la observación del *applet*) por qué la corriente de arrastre es constante, independientemente de la tensión aplicada. ¿Qué controla la magnitud de la corriente de arrastre de electrones?
4. Bajo polarización directa (lado-p más positivo que el lado-n), usando el botón “Ayuda” aparecen en pantalla la definición de los parámetros. Primero se hace clic en el botón “Inicio”. Después en los botones “Parámetros” y “Ayuda”. En la pantalla principal aparecen algunas definiciones:

n_{peq} = concentración de electrones minoritarios en equilibrio en el lado-p,

$Nn(0)$ = número total de electrones que pueden ser inyectados desde el lado-n al lado-p para $V = 0$ Voltios.

¿Es correcto $Nn(0)=n_{peq}$? ¿Es esta igualdad una condición necesaria para el equilibrio?

La cantidad de electrones difundidos es igual al número que puede ir al lado-p sin que se lo impida la barrera de potencial. A partir del diagrama de bandas, se puede encontrar el número de electrones inyectados, Nn , en términos del nivel de impurezas donadoras, Nd , y el potencial, V_0-V , donde V_0 es el potencial termodinámico de la unión y V es la tensión aplicada (Eso es, obtener esta ecuación: $Nn(V) = Nd \exp(-qV_0/kT) \exp(qV/kT)$).

Mostrar que

$$Nn(V) - Nn(0) = n_p(x_p) - n_{peq} = n_{peq} [\exp(qV/kT) - 1].$$

La igualdad, $Nn = n_p(x_p)-n_{peq}$, se mantendrá solo cuando los electrones difundidos (o inyectados) desde el lado-n ($x = x_n$) lleguen sin pérdidas al lado-p ($x = x_p$).

Hacer el mismo razonamiento para el número de huecos difundidos, N_p .

La condición de frontera para el número de electrones en el lado-p es $n_p(x_p)$ en el borde de la zona de carga espacial y n_{peq} en el interior del volumen.

Usando esta condición de frontera, resolver la ecuación de difusión estacionaria:

$$\frac{d^2 n_p(x)}{dx^2} = \frac{n_p(x) - n_{peq}}{L_n^2}$$

Aquí, $L_n = \sqrt{D_n \cdot t_n}$ = longitud de difusión de electrones en el lado-p, y t_n = tiempo de vida de los electrones en el lado-p.

La densidad de corriente de difusión resulta ser:

$$J_n(x) = qD_n \frac{dn_p(x)}{dx}.$$

Aquí D_n es el coeficiente de difusión de electrones en el lado-p. Encontrar la corriente de difusión de electrones en $x = x_p$.

$I = I_0(eqV/kT - 1)$. I_0 es función de D_n , L_n , n_{peq} , D_p , L_p , p_{neq} .

- Basándose en los cuatro puntos anteriores, se entenderá la ecuación **IV** para un diodo ideal con el significado físico de cada término claramente identificado.

Para resumir el funcionamiento se utilizan las tablas de controles y variables. Los controles aparecen en la tabla 3.5:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
<i>arrastre</i>	<i>Checkbox</i>			Corriente de arrastre
<i>difusión</i>	<i>Checkbox</i>			Corriente de difusión
<i>recombinación</i>	<i>Checkbox</i>			Corriente de recombinación
<i>electrones</i>	<i>Checkbox</i>			Flujo de electrones
<i>huecos</i>	<i>Checkbox</i>			Flujo de huecos
<i>polarización</i>	Barra de desplazamiento	-2 ↔ 0.72	Voltios	Tensión de polarización
<i>Inicio</i>	Botón			Polarización 0 V
<i>Ayuda</i>	Botón			Ayuda en pantalla
<i>Na</i>	<i>Choice</i>	$10^{15} \leftrightarrow 10^{19}$	cm ⁻³	Impurezas aceptoras
<i>Nd</i>	<i>Choice</i>	$10^{15} \leftrightarrow 10^{19}$	cm ⁻³	Impurezas donadoras
<i>Parámetros/Ocultar</i>	Botón			Presentación de parámetros

Tabla 3.5. Controles del *applet* de Unión PN en equilibrio

La tabla 3.6 mostrará todas las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
x_p	Anchura z.c.e. en el lado p	cm
x_n	Anchura z.c.e. en el lado n	cm
V_0	Potencial termodinámico	V
V	Polarización de la unión	V
E_{max}	Campo en la unión metalúrgica	V/cm

Tabla 3.6. Variables del *applet* de Unión PN en equilibrio

3.4. LA LEY DE SHOCKLEY

En este *applet* se explica la ecuación corriente-tensión de un diodo ideal PN (también llamada ecuación de Shockley). Se podrá investigar con la ecuación $I=I_0[\exp(qV/kT)-1]$ al trabajar con el programa.

Todas las componentes que componen la corriente total del diodo se muestran mediante una animación visual interactiva. Puede verse que la corriente de arrastre es independiente de la polarización aplicada. Para la corriente de inyección con polarización directa se puede observar su dependencia exponencial con la polarización directa aplicada.

Conociendo la distribución del exceso de portadores minoritarios en cátodo resulta fácil conocer la corriente transportada por ellos puesto que es debida fundamentalmente a difusión. De este modo

$$J_p(x) = -aD_p \frac{dp_n'}{dx} = \frac{eD_p p_{neq}}{L_p} \left[\exp \frac{eV}{kT} - 1 \right] \exp -\frac{x}{L_p} \quad (3.36)$$

Multiplicando y dividiendo por t_p se tiene

$$J_p(x) = \frac{eL_p p_{neq}}{t_p} \left[\exp \frac{eV}{kT} - 1 \right] \exp -\frac{x}{L_p} \quad (3.37)$$

Para ánodo se tiene

$$J_p(x) = \frac{eL_n p_{peq}}{t_n} \left[\exp \frac{eV}{kT} - 1 \right] \exp -\frac{x}{L_n} \quad (3.38)$$

Según estas fórmulas la corriente transportada por huecos disminuye según nos alejamos de la unión, de la misma forma que disminuye la concentración de huecos. Como la corriente total que atraviesa la unión debe ser constante, la disminución de $J_p(x)$ significa que la corriente de huecos se convierte en corriente de electrones conforme nos alejamos de la unión y penetramos en cátodo.

En el diagrama de corrientes que hay en la Figura 3.4 se puede observar que la corriente de electrones que entra por el contacto de cátodo se usa en parte para recombinar con los huecos inyectados en cátodo y el resto para inyectar electrones en ánodo.

Observando el diagrama se ve que entra una corriente I de electrones por un lado y por el otro sale una corriente I debida a huecos. Esto significa que en el diodo tiene lugar una transformación por medio de la recombinación, de la corriente de huecos en corriente de

electrones y viceversa. Por ello, la corriente que atraviesa el diodo será igual al número total de portadores recombinados por unidad de tiempo en todo el diodo puesto que cada recombinación requiere un electrón, que entra por cátodo y un hueco, que entra por ánodo.

Sin embargo en el modelo ideal la recombinación en la zona dipolar vale cero y resulta más cómodo determinar la corriente del diodo sumando las corrientes de electrones y huecos en un punto en el que ambas sean conocidas simultáneamente. Ese punto es el origen y así

$$J = J_n(0) + J_p(0) \quad (3.39)$$

y por tanto

$$J = \left[\frac{eL_p p_{neq}}{t_p} + \frac{eL_n n_{peq}}{t_n} \right] \left(\exp \frac{eV}{kT} - 1 \right) \quad (3.40)$$

que a menudo se escribe de la forma

$$J = J_{sat} \left(\exp \frac{eV}{kT} - 1 \right) \quad (3.41)$$

siendo

$$J_{sat} = \frac{eL_p p_{neq}}{t_p} + \frac{eL_n n_{peq}}{t_n} \quad (3.42)$$

Esta es la célebre ecuación de Shockley que constituye la característica tensión-corriente del diodo ideal.

La Fórmula 3.45 da lugar a una corriente muy disimétrica, según se polarice el diodo directa o inversamente, pues en polarización directa se establecen enseguida corrientes muy elevadas a poca diferencia de potencial que se aplique al diodo, mientras que en polarización inversa, por muy alta que sea la tensión nunca se obtiene una corriente mayor que J_{sat} , que viene a ser del orden de 100 pA/cm². Según esto el diodo se comporta casi como un cortocircuito en polarización directa y casi como un circuito abierto en polarización inversa.

3.4.1. DESCRIPCIÓN DEL APLET

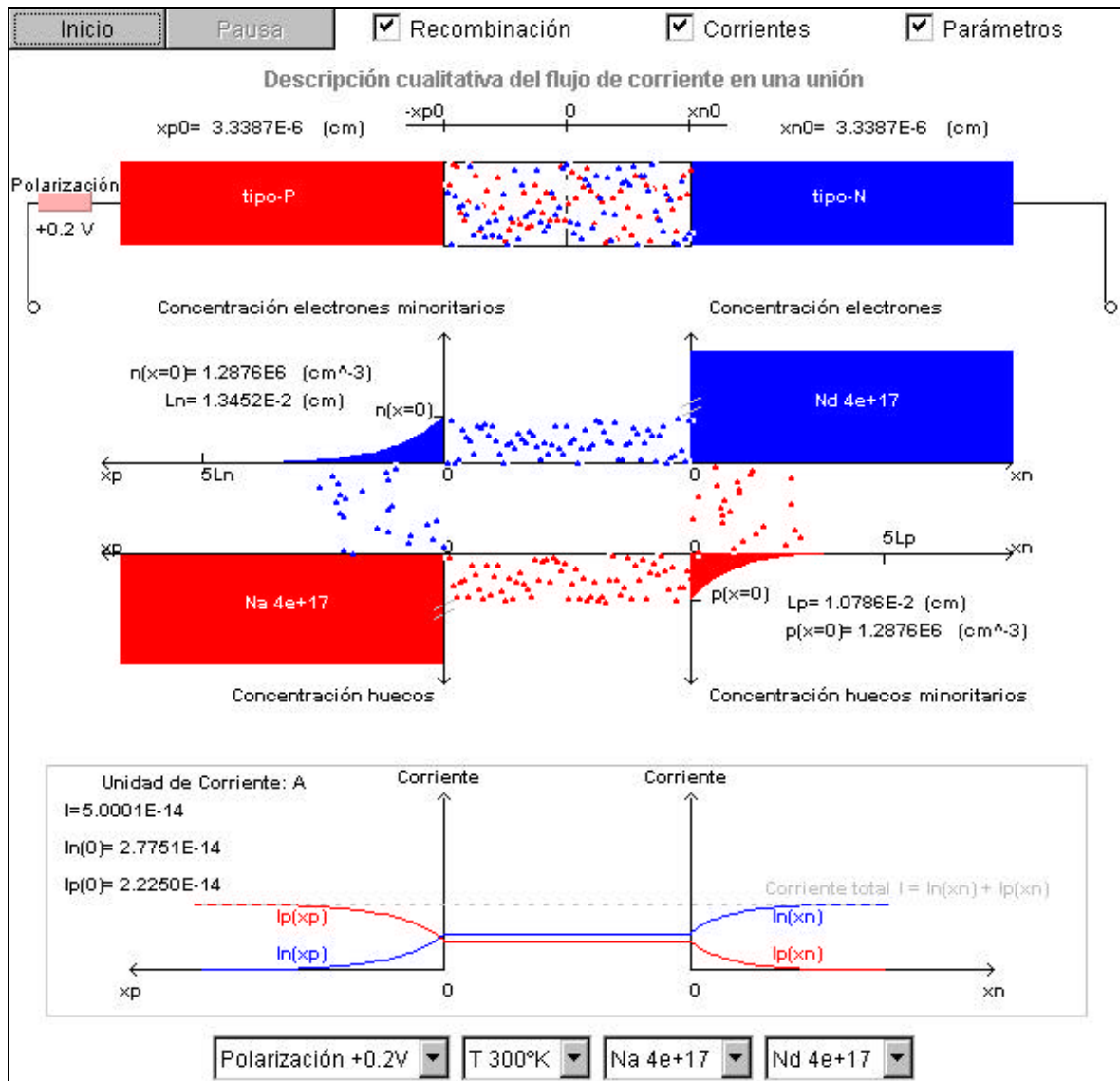


Figura 3.4. La ley de Shockley

El convenio que se ha utilizado en el programa es el siguiente:

Puntos móviles: rojo = huecos, azul = electrones.

Regiones coloreadas: rojo = tipo-p, azul = tipo-n.

$p(x=0)$ = concentración de huecos en exceso en la frontera de depleción

$n(x=0)$ = concentración de electrones en exceso en la frontera de depleción

L_p = longitud de difusión de huecos.

L_n = longitud de difusión de electrones.

I = corriente total

$I_n(x)$ = corriente de electrones (curva azul)

$I_p(x)$ = corriente de huecos (curva roja)

En este *applet* se puede observar el flujo de corriente en un diodo de unión-pn. Bajo polarización directa (lado-p es más positivo que el lado-n, o el lado-n más negativo que el lado-p), se inyectan huecos desde el lado-p (rojo), atravesando la región de depleción (alrededor de la unión, región despoblada de portadores móviles) hasta el lado-n (azul). Los huecos inyectados dentro del lado-n son portadores minoritarios. Cerca de la frontera de la zona de depleción, la concentración de huecos minoritarios está por encima de la concentración en equilibrio térmico debido a esos huecos inyectados eléctricamente. En la región neutra hay dos procesos activos para los huecos minoritarios inyectados: (1) difusión y (2) recombinación.

(1) Difusión: Hay más huecos cerca de la frontera de depleción que en el interior de la región-n. Por lo tanto se produce una difusión térmica, resultando en un flujo neto de huecos (rojo) alejándose de esa frontera.

(2) Recombinación: La concentración de huecos tenderá a recuperar su valor en equilibrio térmico, y entonces intentará librarse del exceso de huecos ($p(x=0)$). Este exceso de huecos sufre una recombinación con los electrones que son los portadores mayoritarios que aniquilan los huecos y electrones formando pares (flujo vertical de puntos rojos). Algunos electrones perdidos (portadores mayoritarios) son rápidamente repuestos a través del contacto óhmico y el cable metálico del lado más lejano de la región tipo-n. Este mismo argumento se aplica a los electrones minoritarios inyectados.

El diagrama de en medio puede ser visto como un diagrama de bandas, pero la banda a través de la región de depleción no se muestra aquí.

Los procesos de difusión y recombinación conjuntamente producen un perfil de concentración exponencial para los huecos minoritarios, como puede verse en el la segunda gráfica (parte inferior derecha en rojo). Para una determinada polarización directa constante, el número de huecos inyectados a través de la unión por unidad de tiempo iguala al número de huecos perdidos por recombinación, estableciéndose entonces una situación estacionaria (es decir, constante en el tiempo).

La componente de huecos de la corriente total se muestra en la curva en rojo de la tercera figura. Desde la izquierda a la derecha, la corriente de huecos en la región tipo-p es una corriente de arrastre (debida al pequeño campo eléctrico en la región neutra tipo-p, en rojo), corriente de arrastre en la región de depleción (debida al campo eléctrico en su

interior), y corriente de difusión en el lado-n (azul) debido al gradiente de concentración. La concentración de huecos minoritarios en exceso llega a cero en un determinado punto.

Los electrones (azul) son inyectados desde el lado-n (azul) a el lado-p (rojo) atravesando la unión (o región de deplexión, gris). Siguen el mismo proceso que los huecos.

Cuando el *applet* aparece por primera vez en la pantalla, comenzará automáticamente la animación. En el panel de control superior tenemos una serie de botones que sirven para controlar esta animación.

Si se quiere cambiar algún parámetro hay que pulsar el botón “Parar”, de este modo se detiene la animación y se activan tanto los controles del panel superior como los inferiores. En las casillas de verificación superiores se puede seleccionar los diagramas que se deseen que aparezcan, así como los parámetros característicos dentro de los mismos. Una vez realizada la selección hay que volver a activar la animación con el botón “Inicio” para poder observar los cambios realizados.

También se dispone del botón “Pausa” para poder detener momentáneamente la animación. Pero en este caso no se pueden modificar los parámetros. Hay que continuar la ejecución pinchando de nuevo en el mismo botón que habrá cambiado su etiqueta anterior y ahora reflejará “Continuar”.

En el panel de control inferior hay cuatro listas desplegables. La primera de ellas es para cambiar la polarización, donde podremos seleccionar tanto valores positivos como negativos. De este modo se puede estudiar el comportamiento del dispositivo con polarización directa e inversa. En segundo lugar tenemos el control de temperaturas que servirá para observar los cambios relativos respecto a variaciones de 5° en torno a la temperatura ambiente que se utiliza por defecto. Los dos últimos controles son de sobra conocidos ya que no es la primera vez que aparecen en el tutorial. Obviamente seleccionarán el nivel de las impurezas aceptoras y donadoras en ánodo y cátodo respectivamente.

Todos los controles aparecen resumidos en la tabla 3.7:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
<i>Inicio/Parar</i>	Botón			Detener y reiniciar
<i>Pausa/Continuar</i>	Botón			Resume la simulación
<i>Recombinación</i>	<i>Checkbox</i>			Diagrama de bandas
<i>Corrientes</i>	<i>Checkbox</i>			Diagrama de corrientes
<i>Polarización</i>	<i>Choice</i>	$-4 \leftrightarrow 0.4$	V	Polarización de la unión
<i>T</i>	<i>Choice</i>	$295 \leftrightarrow 310$	°K	Temperatura
<i>Na</i>	<i>Choice</i>	$4 \cdot 10^{17} \leftrightarrow 4 \cdot 10^{18}$	cm ⁻³	Impurezas aceptoras
<i>Nd</i>	<i>Choice</i>	$4 \cdot 10^{17} \leftrightarrow 4 \cdot 10^{18}$	cm ⁻³	Impurezas donadoras

 Tabla 3.7. Controles del *applet* de La ley de Shockley

La tabla 3.8 muestra las variables:

VARIABLE	DESCRIPCIÓN	UNIDADES
x_{p0}	Anchura z.c.e. en el lado p	cm
x_{n0}	Anchura z.c.e. en el lado n	cm
$n(x=0)$	Exceso de electrones minoritarios	cm ⁻³
$p(x=0)$	Exceso de huecos minoritarios	cm ⁻³
L_n	Longitud de difusión de electrones	cm
L_p	Longitud de difusión de huecos	cm
I	Corriente total	A
$I_n(0)$	Corriente de electrones	A
$I_p(0)$	Corriente de huecos	A

 Tabla 3.8. Variables del *applet* de La ley de Shockley

3.5. EL DIODO VARACTOR

Este apartado se centra en otro aspecto fundamental de la unión PN. Se estudiará la capacidad en función de la tensión en un diodo de unión PN.

El término varactor proviene de las palabras *variable reactor* (reactancia variable) y define un dispositivo cuya reactancia puede variarse de forma controlada mediante una tensión de polarización.

Las relaciones capacidad-voltaje de una unión PN polarizada en inversa siguen la fórmula general

$$C = Ae_s/l \quad (3.43)$$

Donde C es la capacidad de la unión, A el área de la misma, l la anchura de la zona dipolar y ϵ_s la constante dieléctrica del material.

La dependencia capacidad-tensión se establece a través de la anchura de la zona dipolar l , cuya ley de variación $l(V)$, es función de la distribución de impurezas.

El propósito de este *applet* es aprender cuantitativamente el valor relativo de la capacidad de la unión y la de difusión en un diodo de unión PN. Para ello habrá que observar la carga en función de la tensión aplicada en la unión PN. Nuevamente habrá que realizar un estudio por separado para los casos de polarización inversa y polarización directa.

Bajo polarización inversa ($V < 0$), las cargas se almacenan en la zona de depleción del diodo, donde se encuentran ionizadas las impurezas donadoras/aceptoras cuya carga no es compensada por la carga de los electrones/huecos. Cuando la tensión aplicada varía dV , el ancho de la zona de depleción varía dx y la carga espacial total asociada a las impurezas donadoras cambia dQ mientras que la asociada a las impurezas aceptoras lo hace por $-dQ$. La capacidad de la unión será entonces $C_j = dQ/dV$.

Si la polarización es directa ($V > 0$), las cargas no sólo se almacenan en la zona de depleción de la unión, sino también en los portadores minoritarios inyectados. Ahora habrá que sumar dos componentes de la capacidad: la capacidad de la unión y la de difusión. Aunque más allá de cierta polarización directa (aprox., 0.5 V), la capacidad de difusión es la predominante.

3.5.1. DESCRIPCIÓN DEL *APPLET*

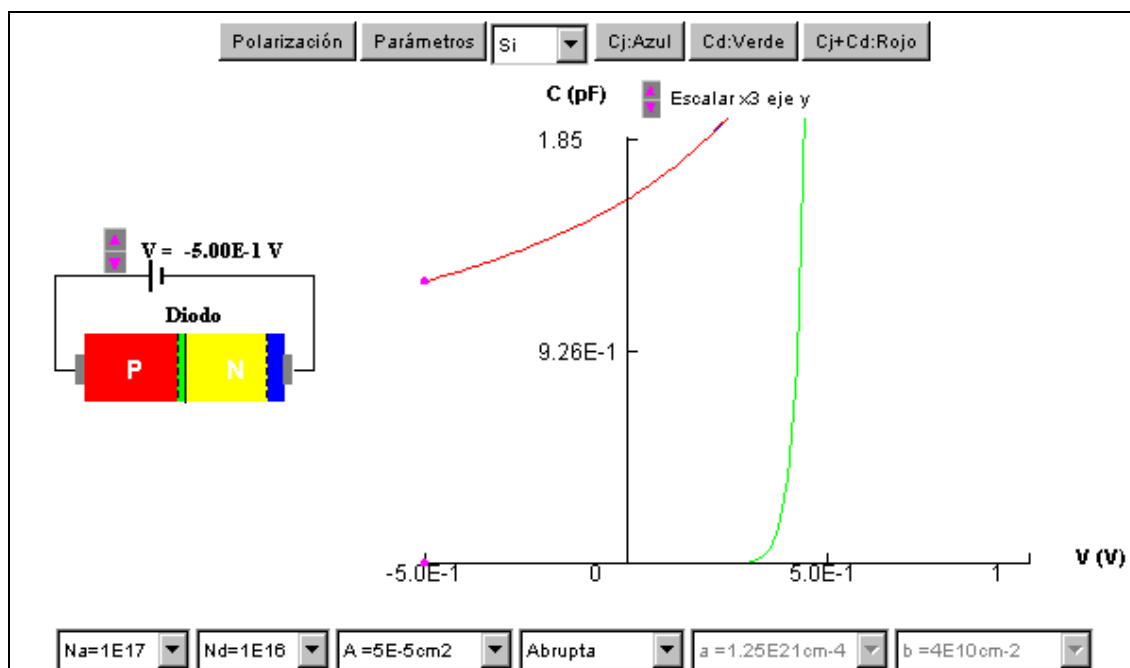


Figura 3.5. El diodo varactor

En el *applet*, la carga neta (sin compensar) se dibuja para las cargas de las impurezas ionizadas (dorado y verde). En la curva C-V, la capacidad de la unión (C_j) se pinta en color azul y la de difusión (C_d) en verde, siendo la capacidad total ($C_j + C_d$) dibujada en rojo.

Como muestra el gráfico, más allá de cierta polarización directa (sobre 0.5 V), la capacidad de difusión es la predominante.

A estas alturas del tutorial el funcionamiento del *applet* resulta de extrema facilidad. En la barra-superior: los botones “Polarización” y “Parámetros” despliegan una ventana flotante. Estas ventanas permiten cambiar el valor de los parámetros.

El valor de polarización que aparece junto a la batería en el esquema del dispositivo tiene a su lado unas flechas que permiten variar su valor. Conforme cambia el valor de V se puede ver cómo nos desplazamos por las curvas de la gráfica con unos puntos que hacen la función de marcadores. La tensión de polarización también se puede cambiar con uno de los botones mencionados anteriormente. Al pulsar el botón de “Polarización” aparecerá la ventana de la Figura 3.6, ahí se puede poner un valor más exacto de la tensión de polarización que utilizando las flechas correspondientes.



Figura 3.6. Ventana de polarización

Con el botón de “Parámetros” aparecerá la ventana de la Figura 3.7.

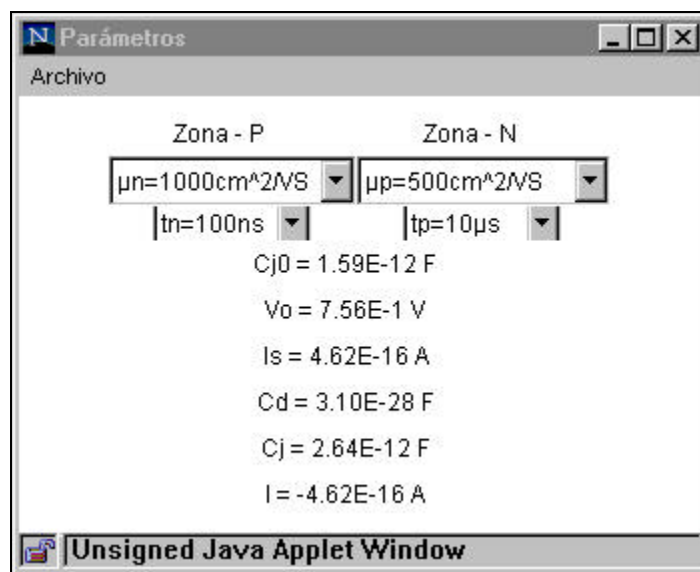


Figura 3.7. Ventana de parámetros

Aquí se puede modificar la movilidad y el tiempo de vida de ambos tipos de portadores. Además tenemos disponible una serie de información complementaria muy interesante. C_{j0} es el valor de la capacidad cuando la polarización es nula, V_0 es el potencial termodinámico, I_s la corriente inversa de saturación, C_d y C_j los valores de las capacidades de difusión y de la unión respectivamente, I es la corriente que atraviesa el diodo. Por los valores que aparecen en el ejemplo está claro que la unión está polarizada en inversa.

A continuación de los botones del panel de control superior descritos hay una selección entre “Si” y “GaAs” para el material del diodo semiconductor. Y los tres últimos botones (C_j , C_d , y C_d+C_j) muestran u ocultan las curvas G-V correspondientes cada vez que se pulsan.

En la barra de controles que componen el panel inferior se podrá variar el tipo de perfil y el área de la unión del diodo. Se puede escoger entre un perfil de la unión lineal, abrupta o hiperabrupta. En función de la selección realizada en el control correspondiente se

activarán los controles adecuados que definen el perfil de la unión. Para una unión lineal podemos variar el parámetro a que controla la pendiente del perfil, y para una unión hiperabrupta tenemos el parámetro b que controla la caída exponencial de la distribución de impurezas del diodo.

Los restantes controles sirven para cambiar el área del dispositivo y el nivel de impurezas de ambas regiones.

En la pantalla central del *applet* se puede ver el esquema del diodo con la polarización, y la gráfica Capacidad-Tensión (C-V) con las respectivas curvas.

Un ejemplo de funcionamiento del programa trabajaría por separado en las dos zonas de polarización:

- Polarización Inversa

1. Cambiar V entre $-0.5V < V < 0$. Observar el valor de la capacidad indicado por el punto en la curva.
2. Pulsar los botones "Cd:Verde" and "Cd+Cj:Rojo" para ocultar las curvas correspondientes. Notar que la capacidad de la unión C_j es la única componente en polarización inversa.
3. Observar también que la zona de depleción (colores verde y dorado) se hace más estrecha conforme la polarización es menos negativa.

- Polarización Directa

1. Pulsar la flecha superior en el botón junto a V , haciendo la polarización positiva.
2. Si la curva C_d (verde) no es visible, utilizar el botón que la hace visible de nuevo.
3. Al incrementar V , el valor de C_d crece de repente exponencialmente. Por encima de cierta polarización $V > 0.5 V$ más o menos, la capacidad de difusión (C_d) es dominante.

Para resumir el funcionamiento se utilizan las tablas de controles y variables. Los controles aparecen en la tabla 3.9:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
<i>Polarización</i>	Botón			Ventana flotante
<i>Parámetros</i>	Botón			Ventana flotante
<i>Semiconductor</i>	<i>Choice</i>	Si, GaAs		Tipo de semiconductor
<i>Cj:Azul</i>	Botón			Curva Cj
<i>Cd:Verde</i>	Botón			Curva Cd
<i>Cj+Cd:Rojo</i>	Botón			Curva Cj+Cd
V	Flechas de control	-0.5 $\leftrightarrow$ 0.74	Voltios	Polarización del varactor
<i>Escalar</i>	Flechas de control	0.8 $\leftrightarrow$ 1.9 $\cdot 10^9$	pF	Escala el eje de capacidad
<i>Na</i>	<i>Choice</i>	10 ¹⁵ $\leftrightarrow$ 10 ¹⁹	cm ⁻³	Impurezas aceptoras
<i>Nd</i>	<i>Choice</i>	10 ¹⁵ $\leftrightarrow$ 10 ¹⁹	cm ⁻³	Impurezas donadoras
<i>A</i>	<i>Choice</i>	10 ⁻⁵ $\leftrightarrow$ 10 ⁻⁴	cm ²	Área de la unión
<i>Unión</i>	<i>Choice</i>	Lineal, abrupta, hiperabrupta		Tipo de unión
<i>a</i>	<i>Choice</i>	4 $\cdot 10^{19}$ $\leftrightarrow$ 4 $\cdot 10^{25}$	cm ⁻⁴	Perfil unión lineal
<i>b</i>	<i>Choice</i>	1.25 $\cdot 10^{10}$ $\leftrightarrow$ 1.25 $\cdot 10^{12}$	cm ⁻²	Perfil unión hiperabrupta
<i>m_n</i>	<i>Choice</i>	10 ³ $\leftrightarrow$ 10 ⁴	cm ² /(V·s)	Movilidad de electrones
<i>m_p</i>	<i>Choice</i>	10 ² $\leftrightarrow$ 10 ³	cm ² /(V·s)	Movilidad de huecos
<i>t_n</i>	<i>Choice</i>	10 ⁻¹⁰ $\leftrightarrow$ 10 ⁻⁵	s	Tiempo de vida electrones
<i>t_p</i>	<i>Choice</i>	10 ⁻¹⁰ $\leftrightarrow$ 10 ⁻⁵	s	Tiempo de vida huecos

Tabla 3.9. Controles del *applet* de El diodo varactor

La tabla 3.10 mostrará todas las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
C_{j0}	Capacidad con polarización nula	F
V_0	Potencial termodinámico	V
I_s	Corriente inversa de saturación	A
C_d	Capacidad de difusión	F
C_j	Capacidad de la unión	F
I	Corriente en el diodo	A

Tabla 3.10. Variables del *applet* de El diodo varactor

3.6. PROBLEMA DEL DIODO VARACTOR

Una de las principales aplicaciones de los diodos varactores es la sintonización de circuitos. Cuando se utiliza en un circuito resonante, el varactor actúa como una capacidad variable permitiendo ajustar la frecuencia de resonancia mediante un nivel de tensión variable.

En el *applet* de la Figura 3.8 el diodo varactor y el inductor forman un circuito resonante paralelo. Los valores de las resistencias de polarización están seleccionados para que no se cargue al circuito en alterna. Igualmente C_1 , C_2 , C_3 y C_4 son capacidades de desacoplo para prevenir que el filtro cargue al circuito de polarización. Estas capacidades no tienen efecto en la respuesta en frecuencia del filtro porque sus reactancias son despreciables a la frecuencia de resonancia. C_1 previene un camino de continua entre el contacto móvil del potenciómetro y el generador de alterna a la entrada a través del inductor y R_1 . C_2 previene del camino de continua desde el cátodo al ánodo del varactor a través del inductor. C_3 evita el camino desde la toma media del potenciómetro a una carga en la salida por el inductor. Y C_4 corta la componente continua de la toma del potenciómetro a masa.

Las resistencias R_2 , R_3 , R_5 y el potenciómetro R_4 forman un divisor de tensión continua que permite alimentar al varactor. La tensión inversa de polarización se puede variar con el potenciómetro.

Hay que tener en cuenta que la frecuencia de resonancia del circuito paralelo es

$$f_r \cong \frac{1}{2\pi\sqrt{LC}} \quad (3.44)$$

Primero se determina el rango de tensiones de polarización inversa del circuito. La tensión de polarización del cátodo está fija a

$$V_K = \left(\frac{R_3 + R_4 + R_5}{R_2 + R_3 + R_4 + R_5} \right) \cdot V_{BIAS} = \left(\frac{2.44M\Omega}{4.64M\Omega} \right) \cdot 60 = 31.6(V) \quad (3.45)$$

La polarización del ánodo se puede variar desde un valor mínimo a uno máximo con el potenciómetro R_4 .

$$V_{A(min)} = \left(\frac{R_5}{R_2 + R_3 + R_4 + R_5} \right) \cdot V_{BIAS} = \left(\frac{220K\Omega}{4.64M\Omega} \right) \cdot 60 = 2.85(V) \quad (3.46)$$

$$V_{A(max)} = \left(\frac{R_4 + R_5}{R_2 + R_3 + R_4 + R_5} \right) \cdot V_{BIAS} = \left(\frac{2.2M\Omega}{4.64M\Omega} \right) \cdot 60 = 28.7(V) \quad (3.47)$$

De modo que los valores mínimo y máximo de polarización inversa del varactor V_R vendrán dados por:

$$V_{R(min)} = V_K - V_{A(max)} = 31.6 \text{ V} - 28.7 \text{ V} = 2.9 \text{ V}$$

$$V_{R(max)} = V_K - V_{A(min)} = 31.6 \text{ V} - 2.85 \text{ V} = 28.7 \text{ V}$$

Utilizando la curva característica del varactor se obtienen unos valores de capacidad de 54.8 pF y 19.4 pF respectivamente para esos valores de polarización. La mínima frecuencia de resonancia del filtro será:

$$f_{r(min)} \cong \frac{1}{2p\sqrt{LC}} = \frac{1}{2p\sqrt{(1mH)(54.8nF)}} = 680kHz \quad (3.48)$$

La máxima frecuencia de resonancia es:

$$f_{r(max)} \cong \frac{1}{2p\sqrt{LC}} = \frac{1}{2p\sqrt{(1mH)(19.4nF)}} = 1.14MHz \quad (3.49)$$

3.6.1. DESCRIPCIÓN DEL APPLET

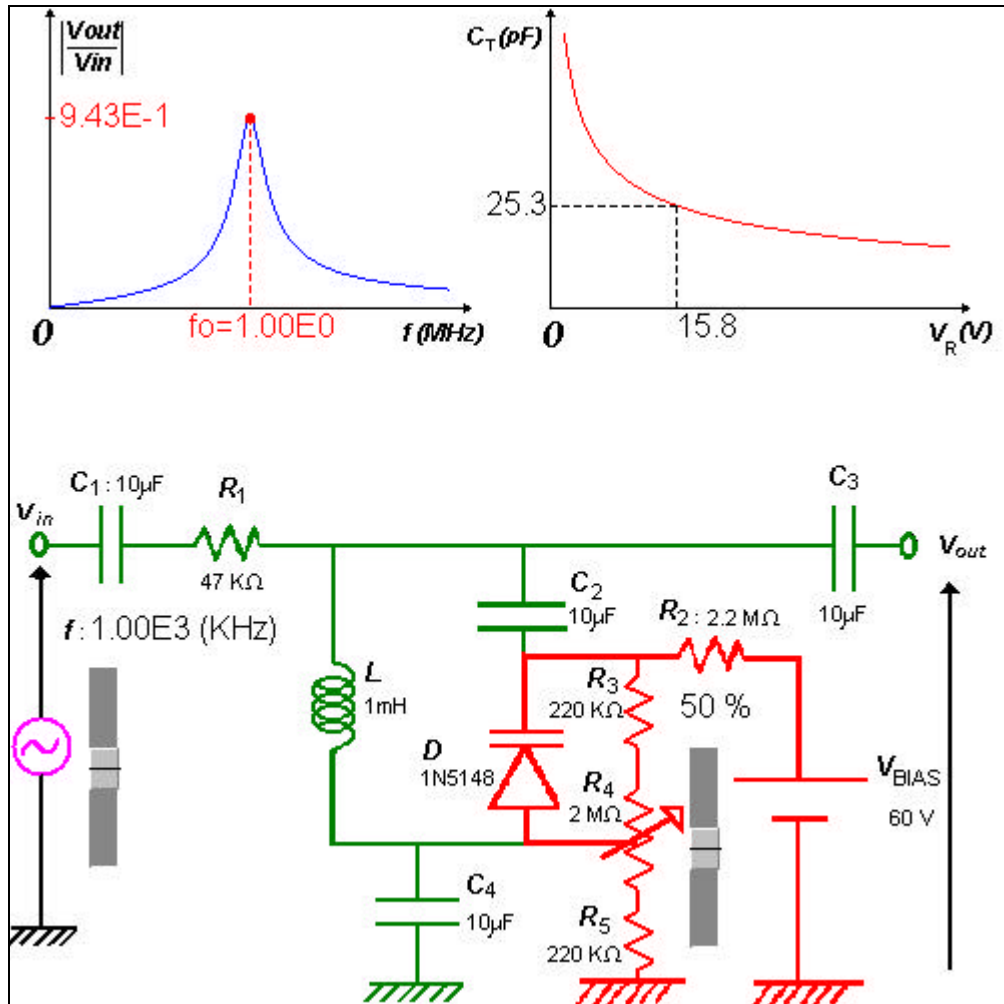


Figura 3.8. Circuito de sintonización con varactor

En este *applet* los controles se encuentran integrados dentro de la estructura del circuito. Se dispone de dos barras de desplazamiento. Una de ellas actúa como potenciómetro de valor 2MΩ. Habrá que arrastrar el control de la barra para variar el valor de su resistencia. La posición actual viene expresada en tanto por ciento justo encima de la barra. Lógicamente con este porcentaje sobre el valor nominal del potenciómetro podemos conocer la resistencia que existe entre la toma intermedia y los extremos.

Al variar el potenciómetro estamos cambiando la polarización del varactor. Esto se verá reflejado en la curva característica del mismo. En esa curva se puede apreciar la tensión inversa de polarización que se le está aplicando, así como el valor de la capacidad obtenida en picroFaradios.

La variación del valor de la capacidad hace que la función de transferencia del circuito resonante paralelo se vea modificada.

En este punto es donde resulta útil el segundo control. Con esta otra barra de desplazamiento se controla la frecuencia del generador de entrada. De este modo se puede examinar el aspecto de la respuesta en frecuencia del circuito ya que se pueden observar los siguientes parámetros:

La frecuencia de resonancia f_0 aparece justo debajo de la campana correspondiente en la gráfica. Esta frecuencia ya se ha comentado que depende del valor de la inductancia y de la capacidad, por lo que cambiará al mover el potenciómetro.

El valor de la frecuencia del generador de entrada está situado encima de la barra de desplazamiento con la etiqueta “f”. Esa frecuencia se corresponde al valor de abscisas del punto de marcación que se mueve por la respuesta en frecuencia.

El valor correspondiente en el eje de ordenadas para la frecuencia del generador se puede leer directamente en el eje de la gráfica. Con este dato se puede estudiar el aspecto de la respuesta en frecuencia del filtro para poder buscar algunos parámetros interesantes como puede ser el ancho de banda de 3dB.

Como se ha visto el *applet* tiene pocos controles, lo que facilita su uso. En la tabla 3.11 se observan los dos controles:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
$R4$	Barra de desplazamiento	$0 \leftrightarrow 100$	%	Potenciómetro de 2 $M\Omega$
f	Barra de desplazamiento	$0 \leftrightarrow 2$	MHz	Frecuencia del generador de entrada

Tabla 3.11. Controles del *applet* de El problema del diodo varactor

La tabla 3.12 enumera las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
V_R	Polarización inversa del varactor	V
C_T	Capacidad total del varactor	pF
f_0	Frecuencia de resonancia del filtro	MHz
V_{out}/V_{in}	Función de transferencia	V/V

Tabla 3.12. Variables del *applet* de El problema del diodo varactor

3.7. DISTRIBUCIÓN DE CORRIENTES Y PORTADORES EN EL BJT

Este *applet* estudia el funcionamiento de un transistor bipolar polarizado en función de la tensión aplicada. Permite observar el flujo de portadores y obtener la distribución de corrientes. La naturaleza del semiconductor es programable, pudiendo trabajar con transistores NPN y PNP.

El circuito que se utiliza para el estudio es un amplificador en Emisor-Común donde el transistor bipolar opera en la región activa.

En la Figura 3.9 se representa el esquema del *applet*. El diagrama de la parte superior izquierda muestra el perfil de la concentración de portadores en la región de la base. Los puntos azules son los electrones que entran y salen de la región de la base, se muestran en la banda de conducción (por encima de E_c , el borde de la banda de conducción). Los puntos rojos son huecos, entran y salen de la base también. El número de puntos es aproximadamente proporcional al logaritmo del valor de la corriente.

Debajo de ese diagrama está el circuito que contiene el transistor bipolar, resistencias $R_b=10\text{ k}\Omega$ y $R_c=5\text{ k}\Omega$, la alimentación suministrada V_{cc} y una señal de tensión. Esta señal de tensión se puede controlar mediante las flechas de control en el fondo del *applet* junto a la etiqueta “V”.

En la parte superior derecha está representada la característica de salida (corriente de Colector I_c frente a la tensión Colector-Emisor V_{ce}) del transistor. La línea verde es la recta de carga de continua:

$$V_{ce} = V_{cc} - I_c \times R_c \quad (3.50)$$

Si las curvas en azul representan las propiedades internas del transistor (característica de salida), entonces la línea verde es la característica del circuito, externa al transistor (recta de carga ya que depende de la resistencia de carga R_c).

En la zona inferior derecha hay listados varios valores de corrientes. El primer subíndice es la zona del transistor, (e=emisor, b=base, c=colector), y el segundo es el tipo de portadores (n=electron, p=hueco). Por ejemplo, I_{en} es la corriente de emisor debida a los electrones.

En el modo activo de funcionamiento la polarización directa de la unión de emisor mantiene un exceso de portadores minoritarios en ambos bordes de la zona dipolar. Este

exceso se va reduciendo al alejarse de la unión debido a la recombinación, tanto en el emisor como en la base. En los bordes de la zona dipolar de la unión de colector la polarización inversa de esta unión produce un efecto de portadores tal que, en los casos normales de una tensión $|V_{CB}| \gg kT/e$, la concentración de minoritarios es nula en esos puntos.

En el emisor, el exceso de concentración disminuye, por recombinación, hasta anularse para puntos muy alejados de la unión. En el colector, por el contrario, el defecto de portadores da lugar a una generación neta de portadores que produce el aumento de la concentración de minoritarios recuperándose la concentración de equilibrio en los puntos alejados de la unión.

3.7.1. DESCRIPCIÓN DEL APPLET

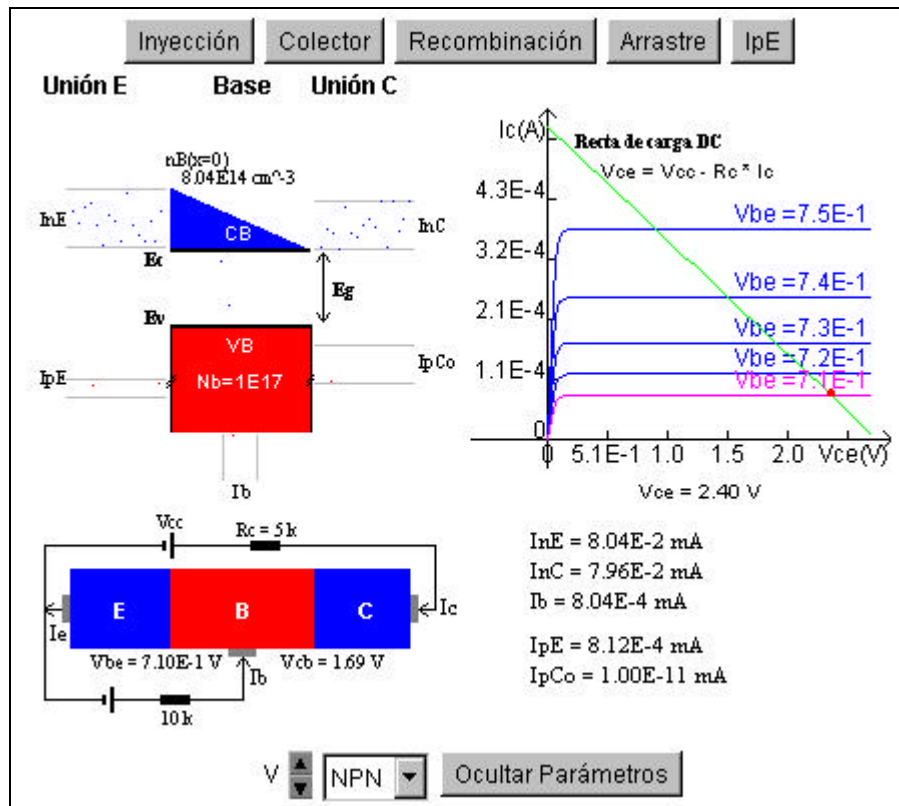


Figura 3.9. El transistor bipolar

Lo primero que hay que hacer es familiarizarse con las componentes de corriente pinchando los botones en la parte superior del *applet*. Los botones ocultan o muestran los correspondientes flujos de portadores.

Segundo, se puede utilizar las flechas de control en la parte inferior del *applet* para cambiar la tensión de entrada. Al variar la tensión estamos observando sólo la región activa directa (unión-Emisor = polarizada en directa; unión-Colector = polarizada en inversa) del BJT. Conforme va variando la tensión de entrada, el punto de trabajo de continua (punto de color rojo) se mueve por la curva característica del BJT dentro de la región activa.

Al variar la tensión de entrada, mientras el transistor se mueve entre el límite superior de este *applet* (más cerca de la región de saturación) y el límite inferior (próximo al corte del transistor), se observa cuánto varía la tensión de la unión de Emisor V_{be} . Lo hace entre 0.66 V y 0.75 V para el transistor mientras el punto de trabajo recorre prácticamente toda la región activa. Por eso casi siempre se asume una $V_{be} = 0.7 \text{ V}$.

En el ejemplo de la Figura 3.9, para silicio con $n_i = 10^{10} \text{ cm}^{-3}$, y dopado con impurezas aceptoras $N_a = 10^{17} \text{ cm}^{-3}$, la concentración de electrones minoritarios en equilibrio es

$$n_{B0} = n_i^2/N_a = 10^{20}/10^{17} = 1000 \text{ electrones/cm}^3.$$

Si se aplica una polarización directa de $V=0.71$ Voltios a un diodo de unión PN de silicio donde el nivel de dopaje tipo-p es de 10^{17}cm^{-3} , la concentración de electrones minoritarios inyectados en el borde de la unión es:

$$n_B(0) = n_{B0} \exp(V/V_T) = 1000 \cdot \exp(0.71/0.0259) = 1000 \cdot 8.042 \cdot 10^{11} = 8.042 \cdot 10^{14} \text{cm}^{-3}.$$

Cambia la tensión de entrada hasta una polarización de la unión Base-Emisor de $V_{be}=0.72$ Voltios. El factor de transporte de Base se define como la fracción de la corriente inyectada por el emisor en la base que alcanza el colector sin recombinarse:

$$\alpha_T \circ I_{nC}/I_{nE} = 1.17\text{E-}1/1.18\text{E-}1 = 0.9915$$

La eficiencia de emisor (g) o eficiencia de inyección se define como la relación entre la corriente de portadores inyectados por el emisor en la base y la corriente total que atraviesa la unión de emisor:

$$g = I_{nE} / (I_{nE} + I_{pE}) = 1.18\text{E-}1 / (1.18\text{E-}1 + 1.20\text{E-}3) = 0.9899$$

La beta del transistor es la ganancia de corriente. Para el caso habitual en que la corriente de base, y , por lo tanto, también la de colector, es mucho mayor que I_{C0} , la expresión de β puede aproximarse por

$$\beta @ I_c / I_b = I_{cn} / I_b = 1.17\text{E-}1 / 1.18\text{E-}3 = 99.15$$

El resumen de controles del *applet* se muestra en la tabla 3.13:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
<i>Inyección</i>	Botón			Flujo de portadores inyectados en la base
<i>Colector</i>	Botón			Ventana flotante
<i>Recombinación</i>	Botón			Curva Cj
<i>Arrastre</i>	Botón			Curva Cd
I_{pE}/I_{nE}	Botón			Curva Cj+Cd
V	Flechas de control		Voltios	Tensión de entrada
<i>Transistor</i>	<i>Choice</i>	NPN, PNP		Tipo de transistor
<i>Ocultar/Mostrar Parámetros</i>	Botón			Parámetros del circuito

Tabla 3.13. Controles del *applet* de El transistor bipolar

La tabla 3.14 mostrará todas las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
$N_B(x=0)$	Perfil de minoritarios en la base	cm^{-3}
V_{be}	Tensión base-emisor	V
V_{cb}	Tensión colector-base	V
V_{ce}	Tensión colector-emisor	V
I_{nE}	Electrones inyectados en la base	mA
I_{nC}	Electrones unión de colector	mA
I_b	Corriente de base	mA
I_{pE}	Huecos unión de emisor	mA
I_{pC0}	Huecos unión de colector	mA

Tabla 3.14. Variables del *applet* de El transistor bipolar

3.8. CONMUTACIÓN DEL TRANSISTOR

En este *applet* se simula la operación dinámica del transistor bipolar, mostrándose en particular que los tiempos de conmutación están relacionados con la carga de los minoritarios almacenada en la base y que la lenta respuesta al apagarse se debe a la saturación del BJT. Se tratan los tiempos de conmutación y los procesos físicos asociados del BJT. Esto es uno de los primeros temas de estudio de los circuitos digitales con transistores bipolares.

Dentro de los circuitos digitales con transistores bipolares, los tiempos de retardo de puerta mejoran desde la Lógica Transistor Resistor (RTL), a la Lógica de Transistor Diodo (DTL), la Lógica de Transistor Transistor (TTL) y la Lógica de Emisor Acoplado (ECL). El retardo de puerta de un moderno TTL puede ser menor de 1.5 ns, y para ECL menor de 1 ns en SSI (Escala de Integración pequeña) y MSI (Escala de Integración Media) e incluso más pequeño en VLSI (Escala de Integración Muy Grande). ECL encuentra su aplicación en circuitos de comunicaciones digitales y otros circuitos de alta velocidad. Los tiempos de subida y bajada que se ven en el *applet* son mucho mayores de 1 ns.

En la actualidad se sigue utilizando TTL aunque en muchas de sus aplicaciones ha sido superada por CMOS. Para bajos retardos de puerta (o velocidades altas de operación), es importante mantener el BJT fuera de la saturación. Tanto TTL como ECL evitan la saturación. Otras tecnologías digitales en proceso emergente son GaAs y BiCMOS. Esta última combina las ventajas de CMOS y de los circuitos bipolares y proporciona los medios para realizar circuitos integrados muy densos, de bajo consumo y de alta velocidad.

Como propósitos principales del programa se busca poder asociar los aspectos de la conmutación dinámica del BJT a los procesos físicos del transistor, en particular, asociar los tiempos de conmutación a la carga en la base, y experimentar la lenta respuesta del BJT saturado.

Se empleará un circuito inversor. La acción de invertir se lleve a cabo por la sencilla ecuación:

$$V_o = V_{cc} - R_c \times I_c$$

Cuando V_i es alta, I_b es alta y entonces I_c también. Por consiguiente V_o es baja. Cuando V_i es baja, V_o es alta.

El proceso de conmutación del BJT depende de la condición inicial de la carga en la Base y de la señal de tensión de entrada aplicada. Se utilizan las siguientes definiciones:

Q_{bo} = cantidad de exceso de carga en la base en el momento de la conmutación

Q_{eos} = carga en la Base al borde de la saturación, es decir, en la frontera entre la región Activa y la de Saturación.

V_1 = tensión de entrada mínima.

V_2 = tensión de entrada máxima.

I_{sat} = corriente de colector en saturación.

t_s = constante de tiempo de saturación.

b = ganancia de corriente.

3.8.1. DESCRIPCIÓN DEL *APPLET*

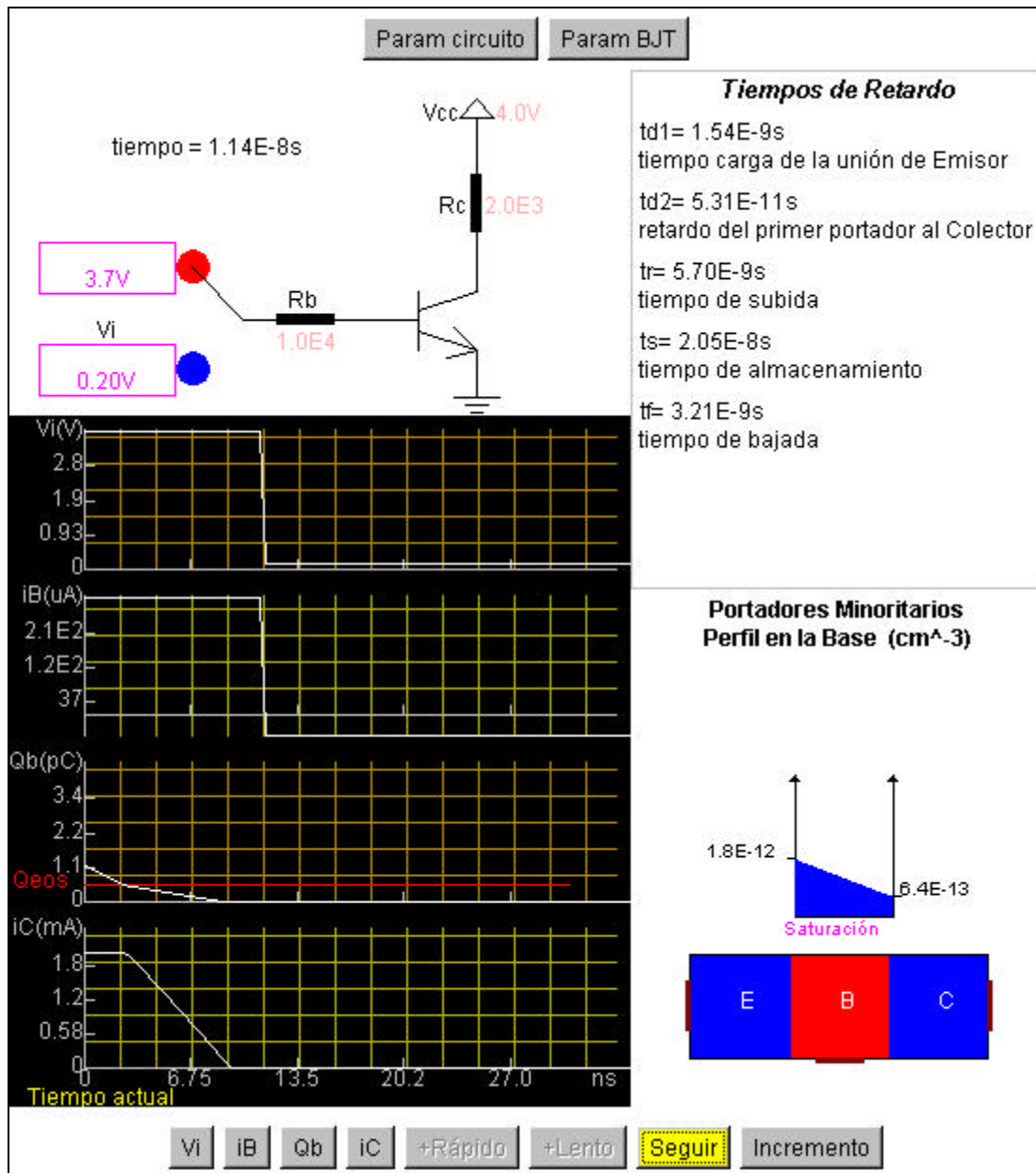


Figura 3.10. El transistor como conmutador

El *applet* se compone de varias secciones:

- Arriba a la izquierda está el circuito básico RTL donde la conmutación se realiza cuando el usuario pincha en las zonas de la tensión de entrada.
- Debajo del circuito hay cuatro gráficas, V_i (tensión de entrada), i_B (corriente de base), Q_b (carga del exceso de portadores minoritarios en la base), e i_C (corriente de colector).

- iii) En la parte derecha de la pantalla aparece el perfil del exceso de portadores minoritarios en la base, y
- iv) en la parte inferior derecha está el esquema del dispositivo.

En la barra superior hay dos botones, el primero para ver los parámetros del circuito (que son variables), y el otro para los parámetros del transistor (también son modificables).



Figura 3.11. Parámetros del circuito

En la Figura 3.11 se representa la ventana flotante que aparece cuando se pulsa el botón “Param circuito”. Como puede apreciarse se pueden modificar los valores máximo y mínimo de la señal de entrada. También son modificables otros parámetros más relacionados con las características del propio circuito como el valor de la alimentación y de la resistencias de base y colector.

En la Figura 3.12 se aprecia la ventana de parámetros del transistor que se obtiene con el botón “Param BJT”.

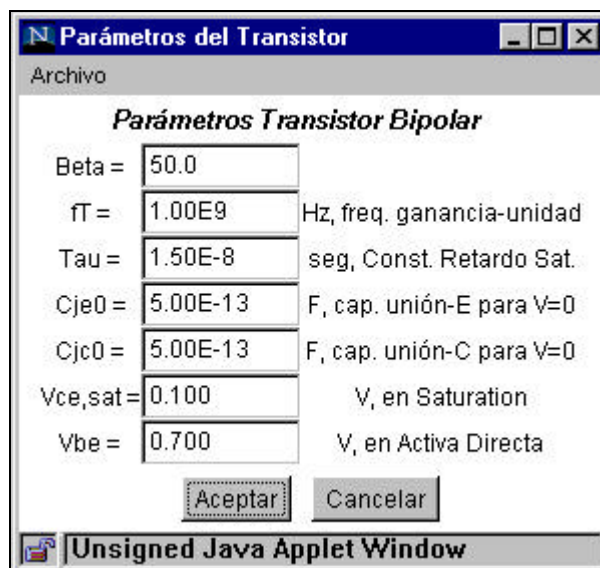


Figura 3.12. Parámetros del transistor

En este caso son los parámetros característicos del transistor los que pueden ser modificados. Por orden de aparición se tienen la ganancia de corriente (β), la frecuencia de transición (f_T), el tiempo de vida medio efectivo de la región de saturación (t_s), las capacidades de la unión de emisor y colector con polarización nula (C_{je0} , C_{jc0}), la tensión colector-emisor de saturación (V_{CEsat}), y la tensión base-emisor en directa (V_{BE}).

La barra inferior sirve para controlar la animación. Los cuatro primeros botones sirven para ocultar sus respectivas gráficas. A continuación tenemos dos botones para acelerar o frenar la animación (la velocidad cambia cada vez que pulsamos los botones), el siguiente botón es para detener o continuar con la animación, y el último botón de 'Incremento', es para moverse hacia delante con pasos discretos de tiempo. Éste último botón sólo está disponible cuando la animación está en modo pausa (botón anterior).

Las cuatro gráficas se dibujan en función del tiempo. La separación entre las marcas horizontales de las gráficas es de aproximadamente 5 ns. Las ondas viajan de izquierda a derecha y los ejes de abscisas son iguales en las cuatro gráficas.

La etiqueta "tiempo = " que se muestra en el diagrama del circuito en la parte superior izquierda de la pantalla, junto al circuito RTL del *applet*, representa el tiempo transcurrido desde la última conmutación en la tensión de entrada, V_i .

En cuanto a los tiempos de retardo del BJT, en la parte superior derecha del *applet* hay cinco tipos diferentes de tiempos de retardo presentes. Estos retardos se explican a continuación: El retardo de encendido consta de t_{d1} , t_{d2} y t_r . Y el retardo de apagado consiste en t_s y t_f .

- Retardo de encendido: $t_{d1} + t_{d2} + t_r$

t_{d1} = tiempo de carga de la unión de emisor:

1. Hacer clic en la zona de V_i (tensión de entrada) para ajustar V_i al la tensión inferior (0.2 V). Se deja funcionar la animación hasta que el BJT llega al corte (observar el perfil de portadores minoritarios en la base: en corte, el exceso de portadores minoritarios en la Base es cero – el perfil de carga en la Base está vacío).
2. Cambiar la tensión de entrada haciendo clic en el conmutador de V_i para pasar al valor superior (3.7 V). El ancho de la zona de carga especial de la unión Base-Emisor disminuye (desde $x_n+x_p=100\text{nm}$ correspondientes a $V_{be}=0.2$ V hasta $x_n+x_p=59\text{nm}$ para $V_{be}=0.7$ V.) Este es el proceso de carga de la unión Base-Emisor, desde el ancho de depleción en corte (próximo a la polarización inversa de la unión PN) hasta la región activa directa.
3. El tiempo transcurrido en este proceso es $td1$ (1.54 ns con los parámetros por defecto.)

t_{d2} = Tiempo que tarda el primer minoritario en llegar a la unión de Colector. En los BJT actuales, este tiempo es despreciable. Se puede omitir este retardo y pasar directamente al tiempo de subida t_r , o bien realizar los siguientes pasos para comprobar lo pequeño que es este retardo para un dispositivo con $f_T = 1\text{GHz}$.

1. Una vez que la unión Base-Emisor está cargada con una polarización directa de 0.7 V aproximadamente, el tiempo restante para que cambie del estado de Corte al de Activa Directa es t_{d2} . En ese punto, el primer portador minoritario ha llegado ya a la unión de Colector. En la zona de tiempos de retardo podemos comprobar que es de $5.31\text{E}-11$ s = 0.053 ns, por lo que es despreciable.
2. Aproximadamente, este valor es igual a $1/(6\cdot\pi\cdot f_T)$. Cuanto más alta sea la frecuencia de corte f_T del BJT, más rápido llegará ese minoritario a la unión de Colector. El valor por defecto es $f_T = 1$ GHz en el *applet*. Los transistores actuales tienen mayor frecuencia de corte (ganancia de corriente unidad) f_T , y el t_{d2} puede ser ignorado.

t_r = tiempo de subida de la corriente de Colector.

1. Con el *applet* en modo de “Pausa”. La Base (o el BJT) acaba de entrar en el modo Activa Directa.
2. Haciendo dos veces clic en V_i ponemos a cero el “tiempo”.
3. Hacer clic en el botón “Incremento” hasta que el BJT entre en modo Saturación por primera vez. Se lee el valor de “tiempo”. Este valor debería ser un poco mayor que el que el tiempo de subida t_r que aparece en el recuadro de tiempos de

retardo ya que t_r se calcula para una subida de i_C desde el 10% al 90% de I_{sat} , en lugar de 0 a I_{sat} .

4. Notar que el perfil de la concentración de portadores minoritarios en la Base ha alcanzado la máxima pendiente. Esto también significa que la corriente de Colector tiene su valor máximo ya que i_C es proporcional a esta pendiente (corriente de difusión).
5. El tiempo de subida t_r está relacionado con cuánto se incrementa la carga en la base desde cero hasta el valor en el borde de la saturación, Q_{eos} .
6. Al continuar en la Saturación no se incrementa la pendiente del perfil de concentración. La carga de la Base continua incrementándose, pero la i_C ha alcanzado su máximo, I_{sat} , como muestran las gráficas de Q_b e i_C . Esto se puede ver si se continúa pulsando el botón “Incrementar” o si se pulsa “Seguir” para que funcione la animación.

- Retardo de apagado: $t_s + t_f$

t_s = tiempo de almacenamiento o retardo de saturación.

1. Cuando V_i se conmuta al valor inferior (0.2V), i_C no responde inmediatamente sino que permanece constante un tiempo t_s . Este es el tiempo requerido para eliminar la carga de saturación de la Base y Q_b alcanza Q_{eos} desde el máximo $t_s \times I_{B2}$ donde $I_{B2} = (V_2 - V_{be})/R_B = (3.7 \text{ V} - 0.7 \text{ V})/10 \text{ k}\Omega$.
2. Durante este tiempo t_s , la unión EB permanece polarizada en Directa $V_{be} = 0.7 \text{ V}$ y $V_i = 0.2 \text{ V}$. Una corriente inversa de base $I_{B1} = (V_i - V_{be})/R_B = (0.2 \text{ V} - 0.7 \text{ V})/10 \text{ k}\Omega$ fluye y ayuda a descargar la Base. Si esta corriente I_{B1} no está presente, la carga de saturación será eliminada por completo por el proceso de recombinación, lo que llevará mucho más tiempo en eliminar Q_b .
3. El tiempo de almacenamiento, t_s , está dado por $t_s = t_s \times \ln[(I_{B2} - I_{B1})/(I_{Bsat} - I_{B1})]$ donde $I_{Bsat} = I_{sat}/\beta$.
4. Para observar t_s en el *applet* hay que seguir los siguientes pasos.
 - a) Poner $V_i = 3.7 \text{ V}$ al valor máximo, y dejar la animación hasta que Q_b no aumente más.
 - b) Hacer clic en el botón “Pausa para suspender la animación. Conmutar V_i al valor inferior $V_i = 0.2 \text{ V}$ y poner a cero la etiqueta de tiempo.

- c) Con el botón “Seguir” se continúa la animación. En el momento en que se pase del estado de Saturación al de activa Directa, se hace clic en “Pausa” para detener la animación.
- d) En este punto, Q_b debería haber alcanzado el valor en el borde de la saturación, Q_{eos} ; i_C todavía permanece en su valor máximo, mientras que el perfil de carga en la Base muestra la misma pendiente.
- e) Se debería leer un tiempo transcurrido algo superior a $2.0E-8 \text{ s} = 20 \text{ ns}$. Este es el tiempo de almacenamiento que produce un tiempo de apagado muy lento para la lógica de BJT saturado.

t_f = tiempo de bajada.

1. Conforme Q_b baja desde Q_{eos} hasta cero, i_C cae exponencialmente con una constante de tiempo que depende de las capacidades de las uniones, C_{je} y C_{jc} . [Por el teorema de Miller, la capacidad de la unión de Colector C_{jc} se ve aumentada por la ganancia de tensión A_v . Pero en este *applet* no se considera el efecto Miller y por tanto el tiempo de bajada puede ser algo mayor que el valor real.]
2. La pendiente del perfil de carga de los minoritarios en la Base disminuye hasta cero, y por lo tanto i_C decrece desde su valor máximo, I_{sat} , hacia cero.
3. Finalmente, la capacidad de la unión Base-Emisor se carga hasta el valor correspondiente a la polarización inversa (o el valor del estado apagado, que es $V_I=0.2 \text{ V}$ en este *applet*). No obstante, esta carga de la union EB no contribuye al retardo de apagado porque i_C ya ha llegado a cero antes de que esta unión comience a cargarse.

Sumario:

Retardo de encendido: $t_{d1} + t_{d2} + t_r = 1.54 + 0.05 + 5.70 = 7.3 \text{ ns}$.

Retardo de apagado: $t_s + t_f = 20.5 + 3.21 = 23.7 \text{ ns}$.

Estos números representan una estimación grosera. Sobre todo si se comparan estos números con los retardos de puerta inferiores a 1.5 ns de la actual TTL e inferiores a 1 ns para ECL incluso en escalas de integración SSI o MSI.

El manual de utilización puede resumirse en la 3.15:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
<i>Param circuito</i>	Botón			Ventana flotante
<i>Param BJT</i>	Botón			Ventana flotante
V_i	Zona de selección	$V_{\min}, V_{\max}$	V	Conmutador de señal de entrada
V_i	Botón			Gráfica tensión de entrada
i_B	Botón			Gráfica corriente de base
Q_b	Botón			Gráfica carga en la base
i_C	Botón			Gráfica corriente de colector
+Rápido	Botón			Aumenta la velocidad
+Lento	Botón			Disminuye la velocidad
<i>Pausa / Seguir</i>	Botón			Parar y reanudar
<i>Incremento</i>	Botón			Simulación paso a paso

 Tabla 3.15. Controles del *applet* de Conmutación del transistor

La tabla 3.16 mostrará todas las variables del *applet*, en este caso algunas son de entrada y otras de salida ya que cuando aparecen las ventanas flotantes hay unos campos dónde se pueden modificar estos parámetros.

VARIABLE	DESCRIPCIÓN	RANGO	UNIDADES	TIPO DE VARIABLE
$V_{\min}$	Tensión mínima de entrada	$-3 \leftrightarrow 0.5$	V	Entrada
$V_{\max}$	Tensión máxima de entrada	$2.5 \leftrightarrow 6$	V	Entrada
V_{cc}	Tensión de alimentación	$2.5 \leftrightarrow 6$	V	Entrada
R_B	Resistencia de base	$0.5 \leftrightarrow 50$	k Ω	Entrada
R_C	Resistencia de colector	$0.1 \leftrightarrow 10$	k Ω	Entrada
β	Ganancia de corriente	$10 \leftrightarrow 10^3$		Entrada
f_T	Frecuencia de transición	$10^7 \leftrightarrow 10^{12}$	Hz	Entrada
t_s	Constante de tiempo en saturación	$10^{-12} \leftrightarrow 10^{-7}$	s	Entrada
C_{JE0}	Capacidad unión de emisor	$5 \cdot 10^{-14} - 5 \cdot 10^{-10}$	F	Entrada
C_{JC0}	Capacidad unión de colector	$5 \cdot 10^{-14} - 5 \cdot 10^{-10}$	F	Entrada
$V_{ce,sat}$	Tensión colector-emisor en saturación	$0.03 \leftrightarrow 0.25$	V	Entrada
V_{be}	Tensión base-emisor	$0.55 \leftrightarrow 0.9$	V	Entrada
<i>tiempo</i>	Tiempo transcurrido		s	Salida
t_{d1}	Tiempo carga unión de emisor		s	Salida
t_{d2}	Retardo primer portador en colector		s	Salida
t_r	Tiempo de subida		s	Salida
t_s	Tiempo de almacenamiento		s	Salida
t_f	Tiempo de bajada		s	Salida

 Tabla 3.16. Variables del *applet* de Conmutación del transistor

3.9. CANAL EN UNA ESTRUCTURA JFET

El transistor de efecto de campo de unión (JFET) es uno de los dispositivos basados en la modulación de la conductancia de un canal a través del cual circula una corriente de portadores mayoritarios. La modulación de la conductancia se logra en el JFET variando la sección transversal del canal mediante la variación de la anchura de las zonas de carga espacial de las uniones inversamente polarizadas.

La corriente en un JFET es transportada por un solo tipo de portadores, los mayoritarios, que circulan fundamentalmente por arrastre. Se dice que es un dispositivo unipolar, en contraste con el transistor de unión que es un dispositivo bipolar.

En la Figura 3.13 se puede apreciar el corte típico de un JFET de canal tipo n. Se toman tres contactos, como se indica en la figura, para los terminales de fuente, puerta y drenador. Aunque para ser estrictos se ven dos contactos de puerta a ambos lados del canal. La corriente, circulará entre drenador y fuente, que son intercambiables, a través del canal formado entre las dos zonas de carga espacial; canal cuya anchura efectiva puede variarse mediante una tensión inversa aplicada a la unión de puerta.

3.9.1. DESCRIPCIÓN DEL *APPLET*

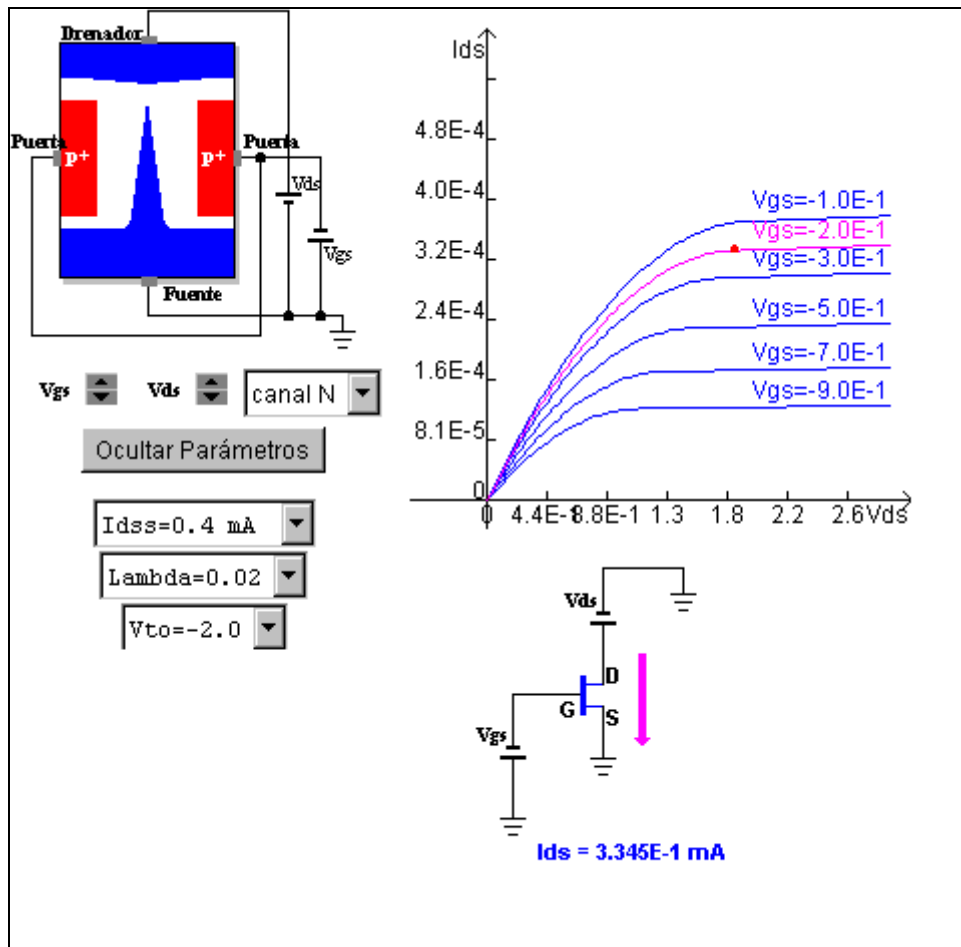


Figura 3.13. Canal en una estructura JFET

En el *applet* se pueden distinguir cuatro zonas distintas:

(1) Arriba a la izquierda se ve el esquema del dispositivo donde se puede observar que el terminal de fuente está conectado a masa, también se ve la polarización de la puerta y del drenador respecto a la citada fuente. El canal del JFET se aprecia entre los dos terminales de puerta, el color dependerá del tipo de transistor que estemos estudiando. Para el JFET de canal n será de color azul y para canal p, rojo. En blanco se pueden apreciar las zonas de carga espacial que son las que modulan la conductancia del canal.

(2) Debajo del esquema del dispositivo se encuentra la zona de control del *applet*. Aquí tenemos todo el control de la simulación. Con las flechas de polarización se puede mover el punto de trabajo del dispositivo al seleccionar valores para la tensión Puerta-Fuente, V_{gs} , y para la tensión Drenador-Fuente, V_{ds} . También podemos cambiar el tipo de

transistor con el que trabajamos y pinchando en el botón de 'Parámetros Físicos' podremos jugar con los valores de beta y lambda del JFET y también con la tensión umbral V_{t0} .

(3) En la parte superior derecha se pueden observar las curvas características del JFET. Aquí es donde representamos el valor de la corriente de drenador en función de la tensión Drenador-Fuente. Las distintas curvas representadas corresponden a los distintos valores de la polarización de la puerta que regula la anchura del canal. Al variar con los controles de polarización podemos observar como se mueve el punto de trabajo por las curvas características.

(4) Por último, la representación simbólica del JFET con su polarización correspondiente. Cuando aumentamos la tensión de drenador se incrementa la corriente que atraviesa el canal, esto se reflejará en este diagrama mediante una flecha que representa la corriente de drenador I_d , se indica la dirección de la misma y además su tamaño irá variando en función del valor de la corriente.

A continuación se presentan la tabla 3.17 con el sumario de controles del *applet*:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
V_{gs}	Flechas de control		Voltios	Tensión puerta-fuente
V_{ds}	Flechas de control		Voltios	Tensión drenador-fuente
<i>canal</i>	<i>Choice</i>	Canal N, canal P		Tipo de JFET
<i>Parámetros Físicos/Ocultar</i>	Botón			Parámetros del JFET
I_{dss}	<i>Choice</i>	0.32 $\leftrightarrow$ 0.48	mA	Corriente de saturación
l	<i>Choice</i>	0 $\leftrightarrow$ 0.08		Factor de modulación del canal
V_{t0}	<i>Choice</i>	-1.8 $\leftrightarrow$ -2.2	V	Tensión umbral

Tabla 3.17. Controles del *applet* de La estructura del JFET

La tabla 3.18 mostrará las variables de salida:

VARIABLE	DESCRIPCIÓN	UNIDADES
V_{gs}	Tensión puerta-fuente	V
V_{ds}	Tensión drenador-fuente	V
I_{ds}	Corriente de drenador	mA

Tabla 3.18. Variables del *applet* de La estructura del JFET

3.10. LA CAPACIDAD MOS

Una estructura MOS ideal se compone de un semiconductor cuyas propiedades son idénticas en todo el volumen; un óxido homogéneo de espesor X_o , aislante perfecto, y que no contiene ninguna carga eléctrica; y un metal que tiene la misma energía de extracción (función trabajo) que el semiconductor, y por tanto niveles de Fermi iguales.

El óxido separa el electrodo metálico (puerta), del semiconductor que está conectado a un potencial de referencia fijo mediante un contacto óhmico situado en la cara opuesta a la del óxido. La diferencia de potencial entre la puerta y el contacto óhmico se llama tensión de puerta (V_G).

Como el óxido es un aislante perfecto, no existe flujo de corriente continua entre el metal y el semiconductor, y por tanto el nivel de Fermi es constante. El semiconductor estará en equilibrio térmico para cualquier tensión de puerta V_G cumpliéndose, en todo punto del mismo, que $pn = n_i^2$. Es decir, para $V_G = 0$ las bandas del M.O.S. ideal son planas.

En la Figura 3.14 puede verse el diagrama de bandas del MOS ideal tipo p , en equilibrio térmico, para $V_G = 0$.

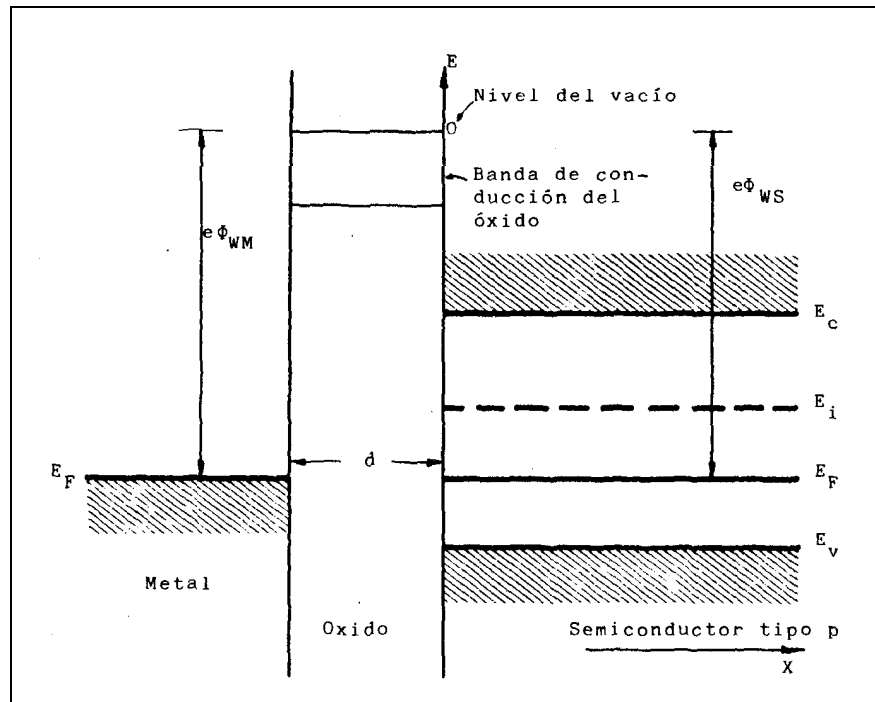


Figura 3.14. Diagrama de bandas de energía del MOS ideal

En un MOS pueden darse tres casos en la superficie del semiconductor:

- Acumulación.

Al aplicar una tensión negativa a la puerta metálica, se atraen los huecos hacia la superficie y se produce una acumulación junto a la interfase óxido-semiconductor. La carga espacial estará formada por dos capas muy estrechas y de alta concentración en las interfases con el óxido.

El campo eléctrico se calcula con la ecuación de Poisson y resulta ser uniforme y de módulo

$$E_{ox} = Q_G / e_{ox} \quad (3.51)$$

donde Q_G es la carga neta por unidad de superficie en el metal, y e_{ox} es la constante dieléctrica del óxido.

Integrando esa expresión se obtiene el potencial:

$$f(x) = V_G (1 - x/x_0) \quad (3.52)$$

que obviamente es una recta.

La caída de tensión en el óxido (V_o) coincide con la tensión de puerta en acumulación:

$$V_G = V_o = E_{ox} x_o = Q_G x_o / e_{ox} \quad (3.53)$$

- Deplexión o vaciamiento.

Con una tensión positiva pequeña en la puerta, los huecos se alejan de las proximidades de la interfase dejando allí una zona con los átomos aceptores ionizados sin compensar. Las bandas de energía del semiconductor se doblan hacia abajo de forma que la banda de valencia se aleja del nivel de Fermi.

La carga por unidad de área contenida en el semiconductor, Q_s , suponiendo vaciamiento total, viene dada por

$$Q_s = - e N_A x_d \quad (3.54)$$

donde

N_A concentración de impurezas aceptoras.

x_d anchura de la región vacía superficial.

La carga del semiconductor se extiende sobre una longitud apreciable, x_d ya que los iones de impureza están fijos.

Una parte de la tensión aplicada, que denominaremos potencial en superficie, f_s se desarrolla en el semiconductor, de modo que la tensión total aplicada V_G se reparte así

$$V_G = V_o + f_s \quad (3.55)$$

Esto es así debido a que en el MOS ideal la energía de extracción (llamada también función trabajo) del metal y del semiconductor son iguales.

De esa distribución de cargas se obtiene, mediante la ecuación de Poisson, el campo eléctrico, cuyo valor máximo vale

$$E_{s,max} = E_s(x_o) = e N_A x_d / \epsilon_s \quad (3.56)$$

siendo ϵ_s la cte. dieléctrica del semiconductor.

El potencial se obtiene integrando el campo eléctrico. Tiene un tramo lineal en el óxido y uno parabólico en el semiconductor. La expresión analítica del potencial en el semiconductor viene dada por

$$f(x) = f_s (1 - x/x_d)^2 \quad (3.57)$$

siendo

$$f_s = E_{s,max} x_d / 2 = e N_A x_d^2 / (2 \epsilon_s) \quad (3.58)$$

como puede comprobarse en la Fig. 3.15.

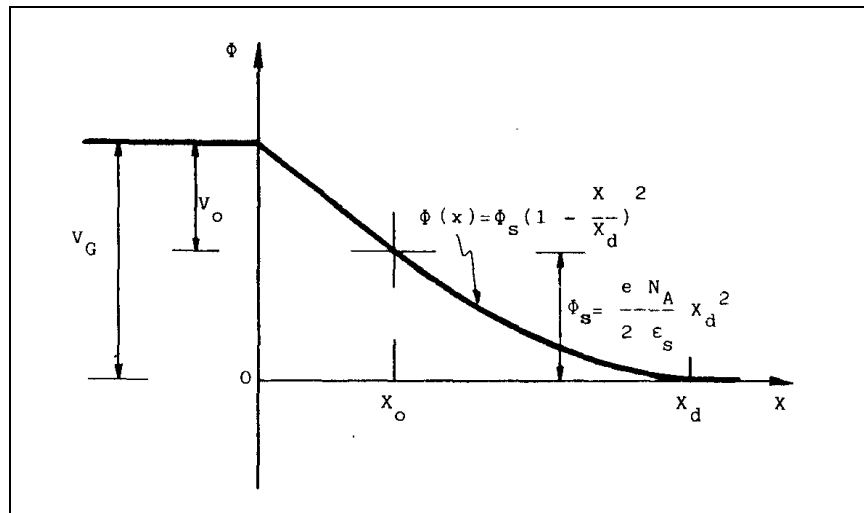


Fig. 3.15. Distribución del potencial en una estructura MOS en deplexión.

No existiendo cargas en la interfase óxido-semiconductor, el teorema de Gauss establece la continuidad del vector desplazamiento eléctrico (Figura 3.16).

$$D_{ox}(x_o) = D_s(x_o) \quad (3.59)$$

y por tanto la discontinuidad del vector E en dicha interfase

$$e_{ox}E_{ox}(x_o) = e_sE_s(x_o) \quad (3.60)$$

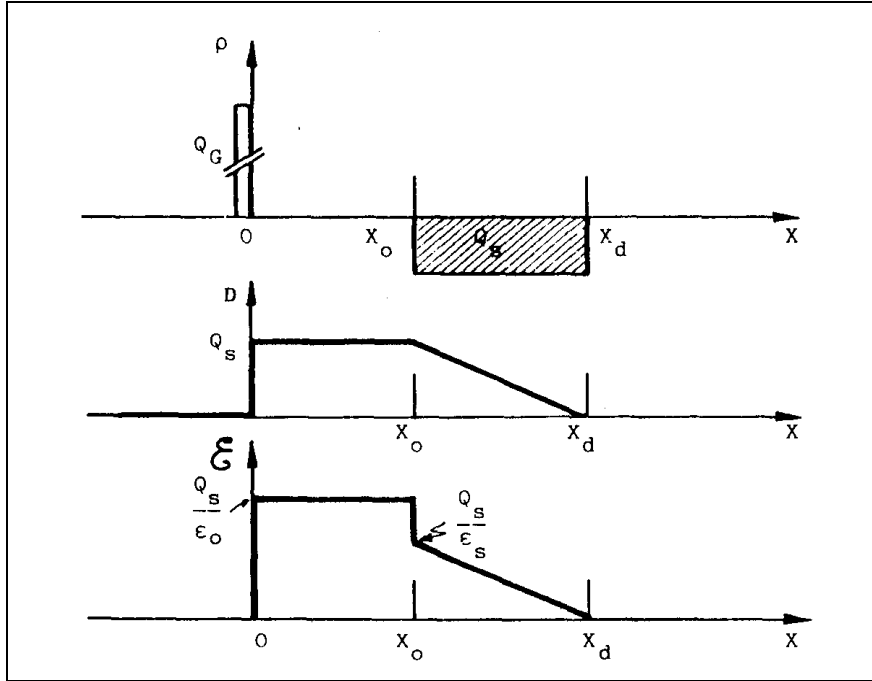


Figura 3. 16. Los vectores D y E en función de la carga espacial .

Por el mismo teorema aplicado al semiconductor se puede escribir, según se ve en la Figura 3.17.

$$e_sE_s(x_o) = -Q_s \quad (3.61)$$

donde Q_s es la carga total en el semiconductor (con su signo correspondiente). Por último, mediante la ecuación de Poisson se establece que, al no haber cargas en el interior del óxido, el campo eléctrico en él debe ser uniforme y de valor

$$E_{ox} = V_o/x_o \quad (3.62)$$

Con las tres últimas ecuaciones se obtiene el valor de la tensión a través del óxido

$$V_o = -x_o Q_s/e_{ox} = -Q_s/C_o \quad (3.63)$$

donde $C_o = e_{ox}/x_o$ es la capacidad de la capa de óxido por unidad de superficie.

Con las Ecuaciones 3.55 y 3.63 puede escribirse

$$V_G = -Q_s/C_o + f_s \quad (3.64)$$

que pone de manifiesto la relación entre la tensión de puerta y la carga espacial en superficie.

Introduciendo en esta última expresión los valores de Q_s y f_s de las Ecuaciones 3.54 y 3.58 se obtiene una ecuación en x_d cuya solución es

$$x_d = \frac{e_s}{e_o} x_o \left(\sqrt{1 + \frac{2e_o^2}{e_s x_o^2 e N_A} V_G} - 1 \right) \quad (3.65)$$

Si la tensión de polarización V_G es cero o negativa, no existe zona de depleción y el semiconductor actúa simplemente como una resistencia en serie con la capacidad del óxido.

En depleción la capacidad total varía porque lo hace C_s a través de la anchura de la zona dipolar x_d .

- Inversión.

Si la tensión de puerta se hace mayor, la anchura de la zona de depleción x_d seguirá creciendo y de acuerdo con la Ecuación (3.58) crecerá también el potencial electrostático de superficie f_s y las bandas de energía se doblarán más hacia abajo como se observa en la Figura 3.17. Pero al irse doblando las bandas, el fondo de la banda de conducción llegará a situarse muy cerca del nivel de Fermi. Cuando esto suceda, la concentración de electrones cerca de la superficie se incrementará muy rápidamente, formándose en superficie una capa en la que los electrones serán mayoritarios. Dicha capa es llamada “capa de inversión”. Visto de otro modo la superficie va haciéndose cada vez más tipo n

Todo aumento posterior de la carga negativa inducida en el semiconductor aparecerá como una carga Q_n debida a electrones situados en una capa de inversión en superficie, sumamente estrecha.

Cuando la capa de inversión se ha formado, la anchura de la zona de depleción alcanza un máximo. Se produce una inversión fuerte y el más pequeño incremento en el doblamiento de las bandas, al que corresponde un pequeñísimo incremento en x_d , produce un gran aumento de carga en la capa de inversión.

Se dice que existe inversión fuerte cuando la concentración de electrones (por unidad de volumen) cerca de la superficie iguale o supere la concentración de iones de impureza del sustrato. Es decir, cuando

$$n_s = N_A \quad (3.66)$$

siendo n_s la concentración de electrones de capa de inversión (cm^{-3})

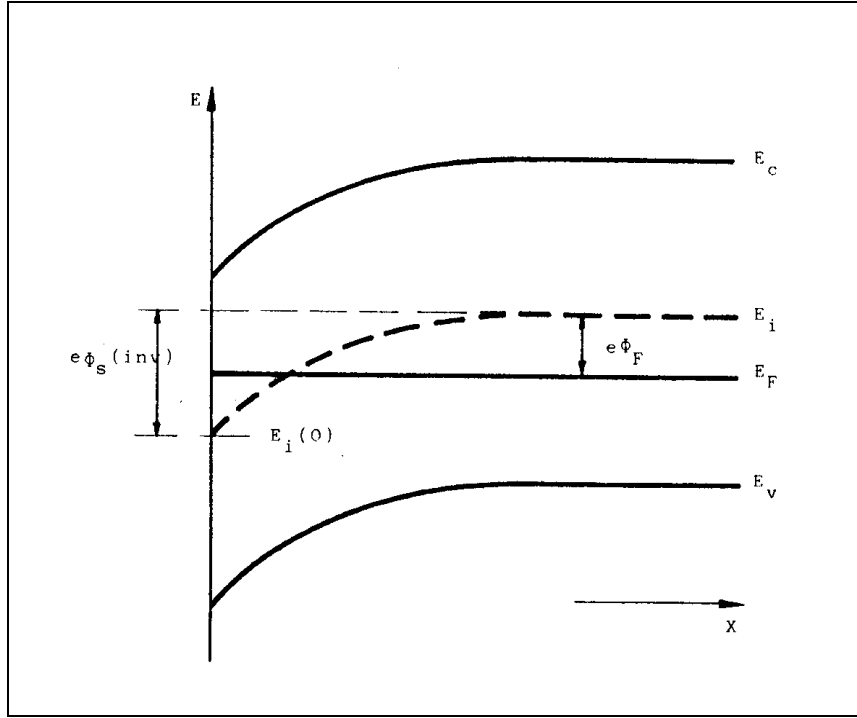


Figura 3.17. Bandas de energía en inversión fuerte.

En estas condiciones y a la vista de la Figura 3.17 se tiene:

en el volumen

$$N_A = n_i \exp (E_i - E_F) / kT \quad (3.67)$$

en la superficie

$$n_s = n_i \exp (E_F - E_i(0)) / kT \quad (3.68)$$

Aplicando la condición de inversión fuerte (Ecuación 3.66) se obtiene:

$$e f_s(inv) = E_i - E_F + E_F - E_i(0) \quad (3.69)$$

de donde se deduce que el potencial en superficie para inversión fuerte vale:

$$f_s(inv) = 2 (E_i - E_F) / e = -2 f_F \quad (3.70)$$

donde

$$f_F = (E_F - E_i) / e \quad (3.71)$$

es el potencial de Fermi.

Una vez alcanzada la inversión fuerte, la tensión f_s se satura en $f_s(inv)$ con lo cual la anchura x_d tampoco crece más. Particularizando la Ecuación 3.57 y despejando se obtiene la anchura máxima de la zona despoblada

$$x_{d,max} = \sqrt{\frac{2e_s f_s(inv)}{eN_A}} \quad (3.72)$$

Si llamamos V_T a la tensión de puerta que hace posible la inversión fuerte podremos escribir

$$V_T = -Q_B/C_o + f_s(inv) \quad (3.73)$$

Esta tensión V_T recibe el nombre de tensión de puesta en conducción (*Turn-on*).

En condiciones de inversión fuerte la carga inducida en el semiconductor por unidad de área viene dada por

$$Q_s = Q_n - e N_A x_{d,max} = Q_n + Q_B \quad (3.74a)$$

$$Q_B = - e N_A x_{d,max} \quad (3.74b)$$

donde Q_B es la carga de la zona de depleción para la anchura máxima.

En la representación el campo eléctrico a partir de la distribución de cargas. Se observará un mayor salto brusco en la interfase óxido-semiconductor debido a la carga Q_n de la capa de inversión.

El potencial responde a las mismas expresiones obtenidas para el caso de depleción con sólo sustituir x_d por $x_{d,max}$.

Hasta el momento la estructura MOS era ideal, pero a partir de ahora se tratará el efecto de la diferencia de energías de extracción (función trabajo), así como la posible presencia de cargas en el óxido.

Las energías de los niveles de Fermi del semiconductor y del metal de una estructura MOS son, en general, diferentes.

La diferencia entre esos niveles de Fermi se suele expresar como diferencia entre las funciones trabajo del metal y del semiconductor.

Cuando el metal de una estructura MOS es puesto en cortocircuito con el semiconductor, aparece un flujo de electrones del metal al semiconductor, o viceversa, hasta que se establezca una diferencia de potencial entre ambos que compense exactamente la diferencia de energías de extracción. Ver Figura 3.18.

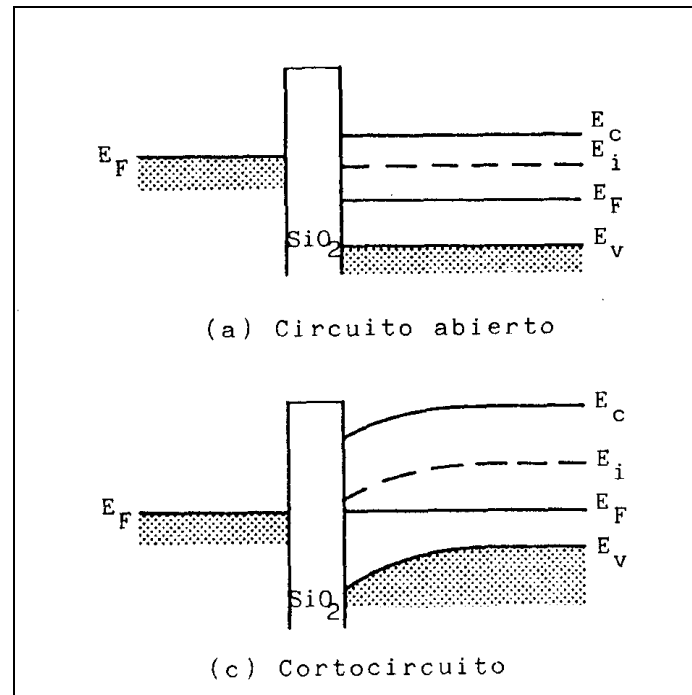


Figura 3.18. Efecto de la diferencia de energías de extracción entre metal y semiconductor.
a) Bandas de energía en circuito abierto. c) Bandas en cortocircuito.

Alcanzado el equilibrio los dos niveles de Fermi están nivelados y por tanto existirá una diferencia de potencial electrostático entre una región y otra como muestra la Figura 3.18c.

Para cuantificar el efecto de la diferencia de los niveles de Fermi, lo más sencillo consiste en determinar el valor de la tensión de puerta necesario para contrarrestar el efecto de la diferencia de niveles de Fermi logrando así una condición de banda plana en el semiconductor (Figura 3.19).

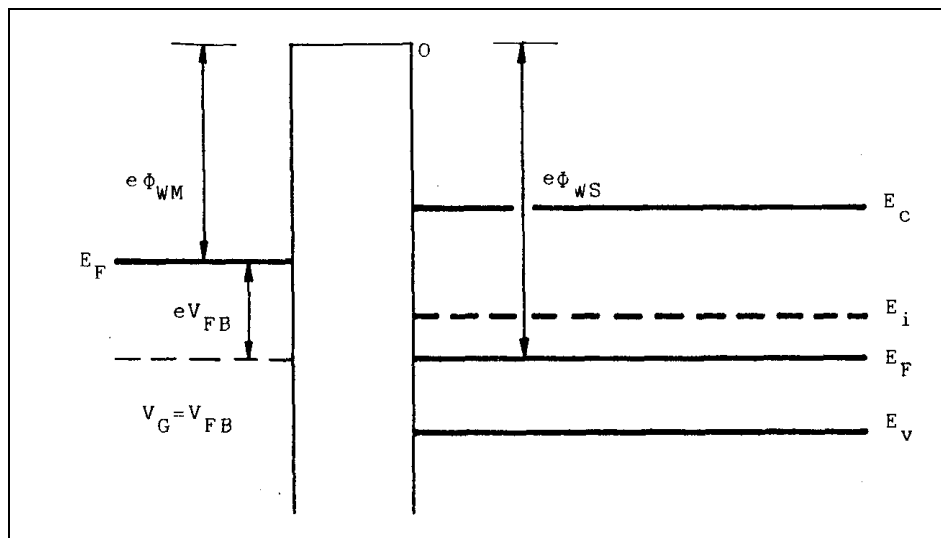


Figura 3.19. Diagrama de bandas de energía del M.O.S. en la condición de banda plana.

Dicha tensión recibe el nombre de tensión de banda plana y se representa por V_{FB} . A la vista de la Figura 3.19:

$$V_{FB} = f_{WM} - f_{WS} = f_{WMS} \quad (3.75)$$

Ahora se considerara que existe una carga pelicular de valor Q por unidad de área en el óxido de una estructura MOS ideal.

En la condición $V_G=0$, esta película de carga inducirá una carga imagen, parte de ella en el metal y parte en el semiconductor.

Para conseguir una condición de banda plana, es decir, que no haya cargas en el semiconductor, es preciso aplicar una tensión negativa al metal. Aumentando esta tensión negativa, ponemos más carga negativa en el metal y por tanto desplazamos el campo eléctrico hacia la puerta hasta que el campo eléctrico en la superficie del semiconductor sea cero. En estas condiciones el área bajo la curva de campo eléctrico es la tensión de banda plana V_{FB} la cual según la ecuación de Poisson vendrá dada por

$$V_{FB} = -x E = -x Q/e_{ox} = -x Q/(x_o C_o) \quad (3.76)$$

Siendo x la distancia de la carga al metal.

La tensión de banda plana no depende sólo de la magnitud de la hoja de carga sino también de su posición dentro del óxido. Así cuando la carga está junto al metal, en $x=0$, no induce carga alguna en el semiconductor y por tanto no tiene efecto en la superficie del mismo. En el otro extremo, cuando la carga está junto al semiconductor ($x=x_o$) ejerce su máxima influencia y da lugar a una tensión de banda plana de

$$V_{FB} = -x_o Q/e_o = -Q/C_o \quad (3.77)$$

La carga en el óxido y las diferencias de energía de extracción entre metal y semiconductor dan lugar, ambas, a una traslación del punto de banda plana desde $V_G=0$ a lo largo del eje de tensiones.

Así pues los efectos descritos determinan un desplazamiento de la curva C-V en el valor

$$V_{FB} = f_{WMS} - Q_{ss}/C_o \quad (3.78)$$

siendo Q_{ss} la densidad de carga en la interfase óxido-silicio.

La Ecuación 3.64 se convierte en

$$V_G - V_{FB} = -Q_s/C_o + f_s \quad (3.79)$$

3.10.1. DESCRIPCIÓN DEL *APPLET*

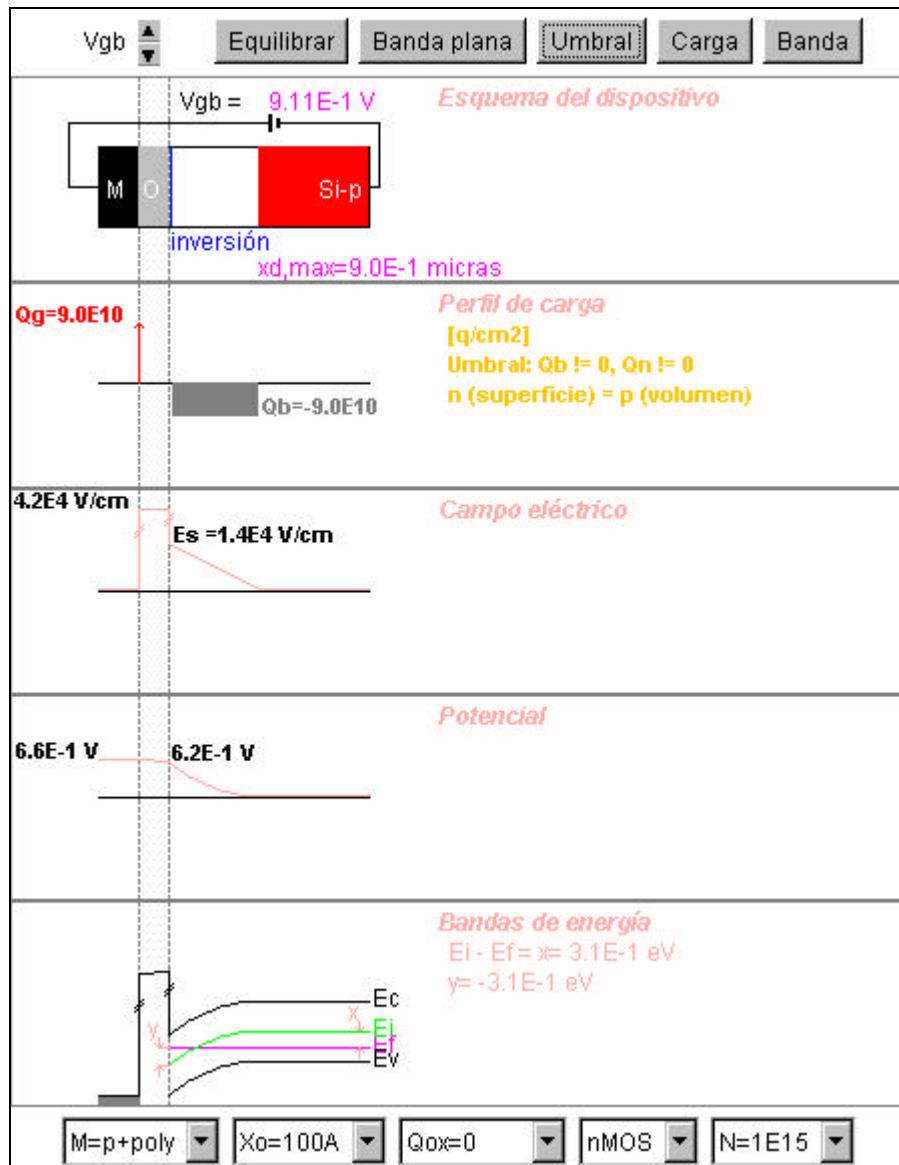


Figura 3.14. La capacidad MOS

La física de la capacidad MOS no es sencilla. Este *applet* proporciona una herramienta dinámica de aprendizaje para esta importante estructura.

En este *applet* se estudiarán las variaciones y distribución espacial de carga, campo y potencial, así como de las bandas de energía, de una estructura MOS en función de la polarización.

Este *applet* ilustra el comportamiento físico de la capacidad MOS bajo una tensión puerta-sustrato aplicada. Se ilustran también tres condiciones importantes: condición de equilibrio, condición de banda plana, y la condición de umbral.

Utilizando las flechas de polarización podemos seleccionar o variar el valor de la tensión de polarización. Se puede observar el perfil de carga, el de campo eléctrico, el de potencial, o el diagrama de bandas de energía. El campo eléctrico y el potencial dentro del semiconductor se obtienen resolviendo la ecuación de Gauss con el perfil de carga dado en el segundo diagrama.

En el modo de depleción, el perfil de campo y potencial en el Si se entiende que se basa en la densidad de carga espacial en la zona de depleción. Para los modos de inversión y acumulación, la caída de campo y potencial en el Si es constante, independiente de la tensión aplicada. Esto se debe a que toda la carga adicional se concentra en la superficie del Si (la carga inducida aumenta exponencialmente con el doblamiento de las bandas que varía con la polarización). Esto es como un efecto de escudo de la carga en la superficie del Si.

Si se aumenta la polarización más allá del valor umbral, el aumento de la inversión de carga bajo el óxido de puerta se ilustra en el espesor de la carga del canal en el diagrama superior.

Se pueden probar diversos metales de puerta, espesores de óxido, el tipo y nivel de dopaje del Si.

El control de la polarización es más útil para observar lo que sucede más allá del umbral. Tres puntos adicionales son:

- Condición de equilibrio: Se aplican cero Voltios. Alineamiento entre el nivel de Fermi del metal (encima de la banda de energía en el lado metálico) y el nivel de Fermi del semiconductor (línea horizontal de color morado).
- Condición de banda plana: La banda de energía del semiconductor es plana. Esto significa (1) no depleción (causa principal para que las bandas se doblen) y (2) carga nula en el canal (la carga en el canal causa una caída de potencial en el semiconductor pequeña, aunque no nula). El requisito de neutralidad de carga implica que una carga $-Q_{ox}$ estará presente en el lado de la puerta.
- Condición umbral: Es la condición donde la densidad volumétrica de carga invertida es igual a la densidad de portadores en el sustrato. Esto tiene lugar cuando $(E_i - E_f)$ en la superficie del Si es igual a $(E_f - E_i)$ en el sustrato. Se debe notar que el nivel de Fermi es constante en el Si (condición de equilibrio térmico) porque no fluye corriente a través de la capa de óxido.

Se pueden ocultar el perfil de carga y el diagrama de bandas si queremos tener una representación más clara. Para eso se utilizan los botones “Carga” y “Banda” que están en el panel de control superior.

Para ver la condición de equilibrio de la capacidad MOS hay que pinchar en el botón “Equilibrar”. Este proceso de equilibrado comienza con las cargas alrededor de la interfaz O-Si: $(Q_{ox}) + (Q_{ch} \text{ o } Q_b) = 0$. Esto es una situación arbitraria que se elige como la condición más natural de inicio. El proceso de equilibrado finaliza con la condición:

E_{FM} en el metal de puerta = E_{FS} en al Si.

Una serie de aspectos a observar puede ser:

E_{FM} (el máximo de energía del metal) se mueve directo a E_{FS} (la línea de color violeta en el Si).

Para equilibrar el MOS se necesita un cable de conexión ya que por la capa de óxido no puede pasar corriente (aislante):

En el Si sólo cae potencial para la condición de deplexión.

La clase de metal afecta al nivel de E_{FM} .

El tipo y nivel de dopaje del Si afectan al nivel de E_{FS} .

El espesor del óxido afecta a la cantidad de carga en el canal, Q_n , necesaria para una correspondiente caída de tensión en el óxido, V_o .

Los resumen de controles aparece en la tabla 3.19:

NOMBRE	TIPO	RANGO	UNIDADES	DEFINICIÓN
V_{gb}	Flechas de control	$-3 \leftrightarrow 3$	Voltios	Tensión puerta-sustrato
<i>Equilibrar</i>	Botón			Condición de equilibrio
<i>Banda plana</i>	Botón			Condición de banda plana
<i>Umbral</i>	Botón			Condición umbral
<i>Carga</i>	Botón			Perfil de carga
<i>Banda</i>	Botón			Bandas de energía
M	<i>Choice</i>	Polisilicio p+, polisilicio n+, Al		Tipo de puerta
X_o	<i>Choice</i>	$50 \leftrightarrow 10^4$	Angstrom	Espesor del óxido
Q_{ox}	<i>Choice</i>	$0 \leftrightarrow 10^{12}$	e/cm^{-2}	Carga en el óxido
<i>Tipo</i>	<i>Choice</i>	NMOS, pMOS		Tipo de MOS
N	<i>Choice</i>	$5 \cdot 10^{14} \leftrightarrow 5 \cdot 10^{17}$	cm^{-3}	Nivel de Impurezas

Tabla 3.19. Controles del *applet* de La capacidad MOS

La tabla 3.20 mostrará todas las variables del *applet*.

VARIABLE	DESCRIPCIÓN	UNIDADES
x_d	Ancho de la zona de deplexión	micras
Q_g	Carga en la puerta	e/cm^{-2}
Q_b	Carga en el sustrato	e/cm^{-2}
Q_n, Q_p	Carga en el canal	e/cm^{-2}
E_s	Campo eléctrico en el sustrato	V/cm
x, y	Diferencia de energía entre el nivel de Fermi y el nive intrínseco	eV

Tabla 3.20. Variables del *applet* de La capacidad MOS

CAPÍTULO 4: CONCLUSIONES Y LÍNEAS FUTURAS

En este capítulo final se resumen las principales ventajas e inconvenientes de los *applets* Java. Estas conclusiones servirán de guía para el futuro desarrollo del presente tutorial, así como de otros nuevos que se puedan desarrollar.

A favor de Java se puede apuntar que es un lenguaje de programación natural para cursos basados en Internet. Por otro lado también proporciona un mayor grado de interactividad y más potencia de cálculo que otras herramientas para crear tutoriales.

En la ardua tarea de demostrar procesos y conceptos físicos complicados, resulta de vital importancia la facilidad que nos ofrece Java para crear gráficos animados. Por lo tanto se puede asegurar sin lugar a duda que es un lenguaje de programación ideal para integrar temas de ciencia e ingeniería en un mismo elemento.

Otra de las grandes ventajas de los *applets* es que pueden estar universalmente disponibles (clase, Universidad, casa, WWW) y en cualquier sistema (PC, MAC, UNIX, en un navegador). Y eso es un aspecto fundamental que los hace adecuados para la educación a distancia.

Los *applets* se integran naturalmente en documentos con hiperenlaces que son esenciales para proporcionar el contexto a los materiales docentes. De este modo publicar un documento en la red que contenga *applets* es tan rápido como puede ser pegar una imagen.

Si se une esto a la gran difusión que ha tenido Internet, se puede afirmar que es el medio ideal para alcanzar el objetivo de este tutorial.

Dentro de los inconvenientes que se pueden encontrar habría que mencionar el gran gasto de tiempo que hay que emplear para desarrollar un elemento del tutorial bien acabado. Más tiempo que con otra herramienta. Por eso es de vital importancia una elección adecuada del tema a tratar en el *applet*. Si todos los pasos se llevan a cabo correctamente, al final merecerá la pena el esfuerzo realizado. Pero existe también la posibilidad de llegar a una aplicación que sea totalmente inútil después de la inversión realizada.

En línea con este último punto hay que comentar que únicamente con el *applet* será difícil introducir un principio o un concepto por primera vez al alumno. Es decir, un

tutorial por muy bien elaborado que esté no podrá sustituir la labor del profesor y del método tradicional de aprendizaje. Esa no es la misión de los *applets*, su función es la de actuar como material de apoyo. Y en esa actividad si que son claramente superiores a cualquier otro método.

Cuando un alumno se enfrenta por primera vez a unos conceptos difíciles como los de la física de los dispositivos de estado sólido, necesita ir paso a paso en el proceso de aprendizaje. En este proceso no es dónde son más adecuados los *applets*. Sin embargo cuando hay que asimilar esos conceptos individuales en su integración conjunta, será el momento adecuado para la visualización interactiva que permitirá el asentamiento de todo el proceso de aprendizaje anterior.

Otro de los factores fundamentales a la hora de decantarse por un tutorial basado en Java es el reciclado que se puede llevar a cabo con el material construido. No será interesante, después del esfuerzo realizado, tener un código tan específico que sólo sirva para un *applet*. Y en ese aspecto Java también marca diferencias.

Se pueden desarrollar varias capas de objetos o componentes de software reutilizable, capas como por ejemplo: Un API (*Application Programmer's Interface*) de estado sólido para programadores de Java; o los propios *applets* de Electrónica de Dispositivos son también objetos reutilizables para profesores y educadores en general.

Java, como lenguaje de programación, es completamente orientado a objetos. Es beneficioso para la productividad programando que el diseñador siga una técnica de diseño orientada a objeto. Al seguir este principio de diseño se irá creando una librería de clases relacionadas con la teoría de los dispositivos electrónicos que ayudará en el proceso de desarrollo. Esta librería de estado sólido se desarrolla en paralelo con la programación de los *applets*, y las clases de la librería irán mejorando con el tiempo.

Los programas desarrollados formarán parte finalmente de una librería de *applets* de estado sólido que será distribuida por Internet. Como ya es conocido, esta librería proporcionará las herramientas de simulación visuales (es decir, los *applets*) como objetos reutilizables para que sean utilizados por el personal docente que no necesitará conocimientos de programación.

Realmente Internet proporciona una fuente inagotable de recursos útiles para los propósitos del tutorial. No sería adecuado lanzarse a la programación de un *applet* si resulta que en la red hay alguno que incluso puede mejorar las expectativas.

Los *applets* son adecuados para diversas disciplinas en el campo de la enseñanza. En Internet se pueden encontrar verdaderas maravillas en el campo de la Física en general.

Particularmente adecuados son los *applets* de electromagnetismo y mecánica. También hay *applets* que abarcan más aspectos relacionados con la asignatura de Electrónica de Dispositivos. Entre estos últimos se pueden encontrar algunos que tratan la estructura de los semiconductores, su diagrama de bandas, circuitos con diodos, operación dinámica de transistores, simulación de los procesos de fabricación, circuitos digitales, memorias, funcionamiento de un CD-ROM, etc.

Por eso al pensar en el futuro que tendrán los *applets* de estado sólido hay que encaminarse a completar la librería de clases y la de *applets*. Habrá que estar atento a las novedades que puedan producirse en Internet y poder seguir aportando algunos aspectos interesantes de los que no haya mucho material.

Sería interesante tratar algo más el tema de la cinética del desequilibrio, como por ejemplo los conceptos del tiempo de vida de los portadores minoritarios, y la variación con el tiempo de su concentración. También se podrían incluir algún *applet* más que utilice los dispositivos semiconductores en alguna aplicación práctica. Un ejemplo de esto último sería un puente rectificador de diodos tan utilizado en los transformadores de corriente, o bien se podría incluir una fuente de corriente.

También sería importante avanzar en el desarrollo de ejercicios y prácticas con los *applets*. De este modo se evitaría jugar sin un propósito concreto con los programas y sería fácil perder la atención de lo que realmente interesa. Esa también sería una línea interesante para mejorar el tutorial ya que si se acompañan los programas de hojas de trabajo y ejercicios que puedan publicarse con sus posteriores resoluciones, se incrementará la productividad del aprendizaje del alumno.

BIBLIOGRAFÍA

- [1] A. Froufe, *JAVA 2 Manual de usuario y tutorial*. Madrid: Ra-Ma Editorial, 2000.
- [2] G. Araujo, G. Sala y J. Ruiz, *Física de los Dispositivos Electrónicos Volumen I*. Madrid: Servicio de Publicaciones E.T.S.I. de Telecomunicación de Madrid, 1986.
- [3] G. Araujo, G. Sala y J. Ruiz, *Física de los Dispositivos Electrónicos Volumen II*. Madrid: Servicio de Publicaciones E.T.S.I. de Telecomunicación de Madrid, 1986.
- [4] G. Sala, *Transistores de efecto de campo*. Madrid: Servicio de Publicaciones E.T.S.I. de Telecomunicación de Madrid, 1986.
- [5] C. Wie, “Educational Java Applet Service”, 2001, Documento en formato HTML accessible por internet en la dirección: <http://jas2.eng.buffalo.edu/jas2.eng.buffalo.html>.
- [6] C. Wie, “Application of the Java Applet Tecnology in a Semiconductor Course”, 1996, Documento en formato HTML accessible por internet en la dirección: http://jas.eng.buffalo.edu/papers/mrs96_jme/paper.htm.
- [7] C. Wie, “Development of Java Applet Resources for Solid State Materials”, 1997, Documento en formato HTML accessible por internet en la dirección: <http://jas2.eng.buffalo.edu/talks/MRS97/mrs97paper.html>.
- [8] C. Wie, “Educational Java Applets in Solid State Materials”, 1998, Documento en formato HTML accessible por internet en la dirección: <http://jas2.eng.buffalo.edu/papers/ieee/ieee2.html>.