

INSTITUTO SUPERIOR POLITECNICO JOSE ANTONIO ECHEVERRIA

SEDE MUNICIPAL PLAZA DE LA REVOLUCION

T tulo: Reconocimiento de Caracteres Manuscritos

Grupo: 2

Curso: 2008 – 2009

Fecha de Realizaci n: 24/05/2009

RESUMEN

A trav s de los a os se ha vuelto imprescindible la conservaci n de documentos y al estos hacerse viejos el papel pierde calidad, entre otros accidentes que puede sucederle. Debido a esto la ciencia ha trabajado arduamente en la soluci n de este problema inventando as - diversos dispositivos electr nicos que realizan internamente m todos totalmente matem ticos que digitalizan todo aquello que se procese, ejemplo de esto tenemos el esc ner, las c maras digitales, entre otros.

Todo este proceso de digitalizaci n se conoce como el Reconocimiento de Caracteres Manuscritos que es el conjunto de t cnicas inform ticas con el objetivo de reconstituir los caracteres de un documento a partir de su propia imagen. En la actualidad esta disciplina cient fica no s lo engloba la reconstrucci n de caracteres, sino la estructuraci n de los documentos, y otras t cnicas de digitalizaci n.

Los sistemas de reconocimiento de caracteres comprenden diversas etapas como son: la adquisici n de la imagen, el pretratamiento, la segmentaci n, el reconocimiento de caracteres, el reconocimiento de fuentes, el proceso de vectorizaci n, el reconocimiento de gr ficos, el reconocimiento estructural y la clasificaci n de documentos.

El reconocimiento de caracteres sigue siendo un problema complejo que tropieza con dificultades a n no resueltas y que son actualmente un objeto de numerosas investigaciones.

SUMMARY

Over the years it has become essential to the preservation of these documents and to take the old paper loses quality, among other accidents that can happen. Because of this science has worked hard on solving this problem by inventing and various electronic devices that perform all mathematical methods internally digitized everything that is processed, we have the example of this scanner, digital cameras, among others.

This whole digitization process known as character recognition Manuscripts is the set of techniques with the aim of reconstituting the character of a document from his own image. At present this field of science not only includes the reconstruction of characters, but the structuring of documents and other digital technologies.

The character recognition system comprising different stages such as: the image adquisici n, pretreatment, segmentation, character recognition, recognition of sources, the vectorization process, the recognition of graphs, structural recognition and classification documents.

The character recognition is still complex problems that is facing difficulties not yet resolved and are currently a subject of numerous investigations.

INDICE

INTRODUCCION.5

DESARROLLO.....6

CONCLUSIONES.....12

RECOMENDACIONES..13

BIBLIOGRAFIA...14

ANEXOS..15

GLOSARIO..16

INTRODUCCION

El problema de la mayoría de las personas e instituciones, es la conservación de los documentos en papel, especialmente cuando estos tienen un siglo de antigüedad. Es por esto que se ha intentado conservar tanto el documento, como la información que este contiene. Un avance en la tecnología a finales del siglo XX, ha dado origen al desarrollo de las bibliotecas digitales, las cuales ponen a disposición de un mayor número de personas la información. Por medio de un escáner se puede digitalizar libros, revistas, documentos, cartas, telegramas, entre otros; pero falta un paso muy importante, interpretar el contenido e información de estas imágenes digitales.

Existen en el mercado una variedad de software de reconocimiento de caracteres, sin embargo todos poseen un margen de error muy alto cuando la entrada es un documento de caracteres manuscritos. Es por esto que se da comienzo a este trabajo.

Una buena segmentación y reconocimiento del carácter, depende de la calidad del documento original y de la calidad de la imagen digitalizada. El principal problema en el reconocimiento de caracteres manuscritos es lo ilegible que son estos caracteres aún para el ser humano.

Durante las últimas dos décadas se han propuesto una gran cantidad de sistemas para el Reconocimiento de Caracteres Manuscritos. Sin embargo aun existen varias limitaciones relativas a su funcionamiento y porcentaje de reconocimiento sobre todo cuando los caracteres manuscritos son de tipo cursivo

El objetivo principal que se persigue con este trabajo es conocer en que consiste el método de Reconocimiento de Caracteres Manuscritos, y aunque este se lleva a cabo en la actualidad con distintos métodos que han ido surgiendo al pasar de los años, aquí se describe en que consiste el Reconocimiento de Caracteres, se explica el método principal y luego de varios de los métodos investigados se selecciona uno de ellos que en este caso es El Sistema Geométrico de Reconocimiento Óptico de Caracteres Manuscritos y se desarrolla, del cual podemos decir que está basado en el uso de un dispositivo de exploración óptica que puede reconocer la letra impresa.

En la actualidad el Reconocimiento de Caracteres Manuscritos no sólo encierra la reconstrucción de caracteres, también su estructuración. Esta técnica comenzó aplicándose en documentos para los cuales ninguna forma electrónica estaba disponible y a medida que la tecnología ha evolucionado, sus aplicaciones han ido en aumento como por ejemplo las cámaras digitales cuando procesan esa foto que tomamos, las computadoras de bolsillo que vienen con pantallas portátiles al igual que las nuevas tecnologías de teléfonos móviles, Ipod, etc. El reconocimiento de caracteres sigue siendo un problema

complejo que tropieza con dificultades aún no resueltas y que son actualmente un objeto de numerosas investigaciones.

DESARROLLO

Introducción al reconocimiento de caracteres manuscritos:

Definición: Es el conjunto de técnicas informáticas cuyo objetivo es reconstituir los caracteres de un documento a partir de su propia imagen. En la actualidad esta disciplina científica no sólo engloba la reconstrucción de caracteres, sino la estructuración de los documentos (títulos, subtítulos, bloques de texto, etc...)

Comenzó aplicándose en documentos para los cuales ninguna forma electrónica estaba disponible. A medida que evoluciona la tecnología, sus aplicaciones han ido en aumento. Los resultados obtenidos hasta ahora distan mucho de ser perfectos. El reconocimiento de caracteres sigue siendo un problema complejo que tropieza con dificultades aún no resueltas y que son actualmente un objeto de numerosas investigaciones.

Varios factores son la causa de estas dificultades:

- ◆ Ausencia de un objetivo universal. Los resultados dependen mucho de la aplicación.
- ◆ Son técnicas por lo general costosas.
- ◆ Muchas son las causas que pueden provocar que los resultados no sean los correctos.

Por ejemplo:

- resolución insuficiente de la imagen.
- introducción óptica de mala calidad.
- documento deteriorado.

En general, los sistemas de reconocimiento de documentos y, por lo tanto, de caracteres comprende las siguientes etapas:

1– Adquisición de la imagen: mediante escáneres y cámaras.

2– Pretratamiento (Binarización, Filtrado, Rectificación).

– Binarización: es el proceso que toma una imagen en escala de grises y se define un umbral: los píxeles con un valor por debajo del umbral se vuelven negros, y el resto se hacen blancos. Esta es una operación útil en tanto que la mayoría de los problemas buscan distinguir dos objetos diferentes en la imagen: el objeto de interés y el fondo. El valor del umbral debe ser escogido de acuerdo a un criterio particular no sólo para cada problema sino para cada imagen, debido a las variaciones en la tinción o en la iluminación de la muestra; en este aspecto es bastante útil la información que da el histograma de la imagen. Los métodos de binarización pueden dividirse en seis categorías diferentes como son:

Métodos basados en la forma del histograma; sus máximos y mínimos, la curvatura, etc.

Métodos basados en clustering; se busca clasificar los distintos niveles de

intensidad en dos grupos diferentes.

Métodos basados en entropía* (Ver Glosario, Pág. 16). Generalmente son métodos que usan la entropía del frente y el fondo de la imagen, o la entropía cruzada entre la imagen original y la binarizada.

Métodos basados en atributos de objetos; se apoyan en medidas como la

coincidencia de bordes, similitud de formas, etc.

Métodos espaciales, basados en distribuciones de probabilidad o en la correlación entre los píxeles de la imagen.

Métodos locales, que adaptan el umbral de acuerdo a las características

de la región analizada.

– Filtrado: en el área de procesamiento de imágenes hay una gran cantidad de operaciones o filtros con efectos muy específicos sobre la imagen aplicada; actualmente son la parte mecánica de los trabajos en el área, pero no por ello menos importante. Son la solución más directa para retocar una imagen antes de intentar cualquier proceso de segmentación.

– Rectificación: es el método que se encarga de rectificar cada uno de estos métodos antes descritos.

3– Segmentación: Mediante la segmentación se divide la imagen en las partes u objetos que la forman. El nivel al que se realiza esta subdivisión depende de la aplicación en particular, es decir, la segmentación termina cuando se hayan detectado todos los objetos de interés para la aplicación. En general, la segmentación automática es una de las tareas más complicadas dentro del procesamiento de imagen. La segmentación va a dar lugar en última instancia al éxito o fallo del proceso de análisis. En la mayor parte de los casos, una buena segmentación dará lugar a una solución correcta, por lo que, se debe poner todo el esfuerzo posible en la etapa de segmentación. Los algoritmos de segmentación de imagen generalmente se basan en dos propiedades básicas de los niveles de gris de la imagen: discontinuidad y similitud. Dentro de la primera categoría se intenta dividir la imagen basándonos en los cambios bruscos en el nivel de gris. Las áreas de interés en esta categoría son la detección de puntos, de líneas y de bordes en la imagen. Las áreas dentro de la segunda categoría están basadas en las técnicas de umbrales, crecimiento de regiones, y técnicas de división y fusión.

Según el grado de asociación entre las operaciones de segmentación y las de reconocimiento, se distinguen tres tipos principales de métodos de segmentación:

– Los métodos explícitos o segmentación en unidades físicas: Estos métodos, intervienen avanzando el proceso de reconocimiento. Las partes segmentadas se dividen prácticamente en letras, tanto que la segmentación se considera una parte del proceso de reconocimiento.

– Los métodos de segmentación implícitos o segmentación en unidades lógicas: Los métodos implícitos, consisten generalmente en una segmentación más fina y así conseguir los puntos de corte correctos. Las partes segmentadas son llamadas grafemas. Estos se usarán más adelante, durante el proceso de reconocimiento/ clasificación. Los grafemas estarán compuestos por fragmentos de caracteres, caracteres o grupos de caracteres.

– Los métodos de segmentación implícitos y exhaustivos: En este caso, es el reconocimiento quien guía la segmentación, así que el sistema de evaluación que se aplica aquí implica un reconocimiento por cálculo de las posiciones sucesivas de la imagen y escoger las posiciones de

segmentación que se correspondan con las responsables de las partes más significativas.

4- Reconocimiento de caracteres: Esta es sin duda la etapa de mayor dedicación ya que el reconocimiento de caracteres es en resumen el proceso donde se aplica métodos de comparación de configuraciones a las formas de los caracteres leídos para determinar qué caracteres alfanuméricos o signos de puntuación representan las formas. Debido a que la diversidad de tipos y formatos de letras (por ejemplo, negrita y cursiva) puede traducirse en grandes diferencias en la forma de los caracteres, el reconocimiento de caracteres puede dar errores.

5- Reconocimiento de fuentes.

6- Vectorización: este proceso transforma las características de la imagen en una línea poligonal o curvilínea, en resumen consiste en convertir imágenes que están formadas por píxeles en imágenes formadas por vectores. Esto se logra dibujando todos los contornos y rellenos de la imagen mediante curvas B-splines. Las curvas B-splines son ampliamente utilizadas en computación gráfica debido a que requieren poco espacio de almacenamiento y son independientes de la resolución de salida que se utilice. Su uso actual se extiende desde la representación de tipografías hasta el modelado de objetos tridimensionales. Los dibujos obtenidos mediante la vectorización son imágenes de contornos perfectamente definidos, que pueden ampliarse o reducirse a cualquier tamaño sin que se modifique su alta calidad. (Ver Anexos, pág. 15: Anexo No 1).

7- Reconocimiento de gráficos (si es que los hay).

8- Reconocimiento estructural: determina la organización lógica de las entidades elementales o compuestas, se asume una representación estructurada de los objetos, estando usualmente los subobjetos representados de manera similar a la utilizada en los métodos globales. Las formas se representan mediante una serie de "reglas de composición" que deben cumplir los subobjetos para pertenecer al conjunto, existiendo un amplio abanico de posibilidades (árboles de decisión, expresiones lógicas, redes, modelos de Markov, reglas, entre otros).

9- Clasificación de documentos: se distingue el tipo de documento reconocido.

En particular, en la etapa del reconocimiento de caracteres se divide en dos sub-etapas:

1- Extracción de características:

- Permite conocer medidas (tamaño, perímetro, centro de gravedad, momentos...).
- Características topológicas (orientación de segmentos, número de agujeros, número de extremidades, etc...).

2- Etapa de decisión: Tres técnicas destacan sobre las demás:

- Redes neuronales (capacidad de aprendizaje).
- Cadenas Ocultas de Markov. Estudios y algoritmos probabilísticos.
- Voto mayoritario. Combinación de diferentes estrategias. Se escoge la clase con mayor número de clasificaciones. Es la técnica que mejores resultados ofrece.

En la mayoría de los tratamientos se requiere para su buen funcionamiento una contribución de información del contexto, dependiente del tipo del documento analizado. Esta información se proporciona

por los llamados modelos de documentos. Varias etapas del reconocimiento requieren esta clase de conocimientos: un reconocedor de caracteres utilizará; por ejemplo una base de datos de caracteres de referencia o diccionarios lingüísticos; el reconocimiento de fuentes necesitará; una base de conocimiento de las características de las fuentes en cuestión.

El reconocimiento de caracteres tiene como objeto el asociar a una imagen la identidad correspondiente de entre los símbolos de un determinado alfabeto. Existen varias modalidades entre las cuales se pueden distinguir reconocimiento on line y reconocimiento off line, es decir teniendo o no en cuenta la información temporal asociada a los trazos que componen el carácter. Sólo se tratará; aquí- la segunda opción. Si la información del carácter se obtiene a través de medios ópticos, suele hablarse de Reconocimiento Óptico de Caracteres (OCR). En este caso, el problema del reconocimiento de los caracteres es sólo una parte de un problema mayor que se conoce como Análisis o Comprensión de Documentos. Se trata, evidentemente, de obtener una representación simbólica lo más fiel y completa posible a partir de la imagen digitalizada de un documento escrito. Las etapas de que consta un sistema de análisis de documentos son: adquisición de la imagen y preproceso, segmentación de ésta en bloques de gráficos, bloques de texto, líneas de texto y caracteres, y ordenación de estos elementos, reconocimiento de los caracteres impresos independientemente del tipo y tamaño de letra y, finalmente, recuperación de errores en el texto por corrección ortográfica y aplicación de un modelo de lenguaje.

La complejidad del sistema varía en función del tipo de caracteres que aparecen en el documento. En este sentido se distinguen las siguientes variedades de menor a mayor dificultad: reconocimiento de uno o varios tipos de letra impresa (fixed-font OCR y multifont OCR), de cualquier tipo de letra impresa (omnifont OCR), de caracteres aislados manuscritos (handwritten OCR) y, finalmente, de escritura manual cursiva (script recognition). En cuanto al método de reconocimiento propiamente dicho, dos aproximaciones distintas que coinciden con las que se distinguen en Reconocimiento de Formas en general son: el reconocimiento geométrico (estadístico o basado en la teoría de la decisión) y el reconocimiento estructural. Cada uno de ellos tiene sus ventajas y sus inconvenientes. La información estructural presente en los patrones a reconocer es difícilmente aprovechada, en general, por los métodos geométricos. Esta información, sin embargo, puede tener extraordinaria importancia en muchos casos. Por otro lado, los métodos de extracción de primitivas (paso necesario previo a la aplicación de un método estructural y, en particular, sintáctico) son generalmente más costosos y pueden causar problemas en presencia de fuerte ruido. Los propios métodos sintácticos añaden a menudo más coste que los geométricos y pueden ser también más sensibles al ruido si no incorporan una formulación estocástica. La mayor parte de los trabajos que abordan el reconocimiento de caracteres manuscritos hace uso de métodos híbridos sintáctico-geométricos, en este trabajo se presenta un método de parametrización, en vectores de talla fija, de imágenes digitalizadas de caracteres manuscritos. La representación de los objetos en forma de vectores es una característica fundamental de un sistema de reconocimiento geométrico de formas. El espacio vectorial definido con esta representación es el marco de aplicación de gran cantidad de técnicas tanto paramétricas como no paramétricas. Entre las primeras pueden citarse las basadas en la estimación de probabilidad asumiendo distribuciones normales, mezclas de gaussianas, etc. Entre las segundas, las funciones discriminantes, la búsqueda de vecinos y las ventanas de Parzen asumiendo una métrica, etc. Tanto unas como otras se apoyan en bases teóricas procedentes de la estadística y la teoría de la decisión.

A continuación se presenta de forma resumida el método de Parametrización:

Parametrización: El objetivo es encontrar un conjunto de parámetros que definan la forma del carácter bajo análisis de forma precisa y única. Además, la continuidad de la representación es imprescindible para asegurar el mayor grado de inmunidad al ruido y la mayor capacidad de generalización posibles. Esto significa que objetos similares deben dar lugar a representaciones similares. En conclusión, tres características que definen un método de parametrización adecuado son: precisión, unicidad y continuidad. Por otro lado la concisión del método de parametrización, dependiente del número de

parámetros que se generan (dimensionalidad de la salida) y en menor medida del rango y la cuantización de éstos, es también un factor clave por varios motivos. Uno de ellos es la complejidad temporal y espacial de muchos algoritmos de clasificación, que crece muy deprisa cuando aumenta la dimensionalidad de los vectores a clasificar. En algunos casos este crecimiento puede llegar a ser exponencial. Otro argumento en contra de una alta dimensionalidad de la representación es el incremento exponencial de la dispersión de un conjunto de puntos en espacios de dimensión creciente. Esto significa que la obtención de estimaciones estadísticamente fiables en estos espacios requiere un número de muestras que crece exponencialmente con la dimensión, en realidad es el aumento de la dimensionalidad intrínseca o topológica la que determina este hecho, pero ésta crece generalmente con la dimensionalidad algebraica. Otros aspectos también relevantes a la hora de diseñar o escoger un método de parametrización dependen directamente de las particularidades del problema. El factor quizás más importante a este respecto es la invarianza de la representación. La parametrización debe ser invariante a distorsiones o alteraciones que no varíen la naturaleza del objeto representado, esto es, que no hagan variar la clase en la que éste se va a clasificar finalmente. Desgraciadamente, esta importante propiedad es difícil de obtener a través de métodos analíticos o teóricos y se aborda generalmente mediante heurísticos y procedimientos hasta cierto punto intuitivos.

Clasificación geométrica de patrones: Se han empleado diversos métodos de clasificación para evaluar las prestaciones del método de representación LLF (Local Line Fitting). Una red neuronal (Perceptrón Multicapa), una variante del LVQ (Learning Vector Quantization) llamada DSM (Decision Surface Mapping), la regla de clasificación del vecino más próximo con la (distancia euclídea)* como medida de distancia y esta misma regla con un conjunto de prototipos editado y condensado han sido los paradigmas utilizados. Las redes neuronales son estructuras formadas por elementos de proceso muy simples. La topología de un perceptrón multicapa consiste en cierto número de capas de unidades conectadas en cascada. A la primera capa (capa de entrada) se le presentan las muestras a clasificar (vectores de dimensión fija). Las capas posteriores (capas ocultas) alimentan la capa de salida para que ésta proporcione una salida (otro vector de dimensión fija) en función de los pesos de los enlaces entre unidades. Son precisamente estos pesos los que se van adaptando en la fase de aprendizaje por medio de algoritmos que comparan el resultado proporcionado por la red con el deseado y modifican el valor de los pesos de modo que la discrepancia sea mínima. En este caso se utiliza el algoritmo Backward error propagation sobre un perceptrón con tres capas de pesos.

La regla del vecino más próximo es una de las técnicas de clasificación más conocidas. Sea $\{x_1, x_2, \dots, x_n\}$ un conjunto de n muestras etiquetadas. Sea x_i la muestra más cercana a x según una cierta medida de distancia. La aplicación de la regla del vecino más próximo supone asignar a x la clase a que pertenece x_i . El análisis para $n \rightarrow \infty$ demuestra que el error cometido nunca es mayor que el doble del error de Bayes para un conjunto infinito de muestras. La regla de los k -vecinos más próximos es una extensión de la anterior en la que x se clasifica en la clase que más veces aparezca en el conjunto de los prototipos $x_1, x_2, \dots, x_K$ más cercanos a x . El número óptimo de vecinos depende fuertemente (en relación inversa) de n y de la dimensión del espacio en que se definen los patrones. Aunque la dimensión aplicable es la topológica, una estimación probablemente optimista (asumiendo incluso una dimensión topológica un orden de magnitud menor que la algebraica) para nuestro problema revela una k óptima menor que 1. Este hecho se ha comprobado experimentalmente en el caso de los dígitos con LLF donde los mejores resultados se obtienen con $k=1$. Las técnicas de edición y condensado son ya clásicas en el área de reconocimiento de formas. El objetivo es obtener un subconjunto del conjunto de muestras original que preserve en la medida de lo posible las fronteras de decisión de éste. Un menor número de prototipos implica un reconocimiento más eficiente y, en algunos casos, mejores prestaciones debido a la presencia masiva de outliers en el conjunto de entrenamiento. Otro tipo de técnicas son las de modificación adaptativa de los prototipos. Las más conocidas son, posiblemente, las técnicas LVQ. En este caso, se da libertad a los prototipos para moverse libremente en el espacio de características. Es evidente que esta libertad no puede existir si el espacio en el que se definen las muestras no es vectorial (bien sea métrico, pseudométrico, etc.). La aproximación adoptada ha sido, concretamente, llamada también DSM (Decision Surface Mapping). La característica principal de esta técnica es la modificación del prototipo más

cercano de la misma clase que la muestra acercándolo a ésta y la modificación del más cercano de otra clase alejándolo de ella. Las muestras sólo provocan la alteración de prototipos si son clasificadas incorrectamente con el conjunto de prototipos actual. Aunque las superficies de decisión obtenidas no respetan (en la misma medida que las técnicas LVQ) las fronteras de decisión de Bayes, los resultados experimentales parecen superar ampliamente los obtenidos con LVQ. Se utiliza una variante original de la aproximación DSM en la que se fija el error máximo de resubstitución y no el número de prototipos. De esta forma, la especificación del modelo es mucho más natural. El coste temporal es mayor, sin embargo, ya que el método empieza por asignar un prototipo a cada clase y va añadiendo nuevos prototipos a aquellas clases que no superen el error de resubstitución prefijado.

CONCLUSIONES

Debido al problema que ocurre en la actualidad en cuanto a la conservación de los documentos en soporte papel ante todo cuando tiene hasta un siglo de antigüedad, al igual que la información en las bibliotecas, que se está trabajando para que sea digital debido a que ya existen muchos de sus libros deteriorados, además de que por una causa lógica debemos ir a favor del desarrollo y por supuesto es una forma más moderna, menos costosa, más fácil y más eficiente para el trabajo, la ciencia se ha esmerado por trabajar dedicadamente en diseñar métodos para la digitalización de los mismos. Todos conocemos que por medio de un escáner se puede digitalizar libros, revistas, documentos, cartas, telegramas, entre otros; como también ocurre el proceso de digitalización a través de las cámaras digitales, las minicomputadoras portátiles, la nueva telefonía móvil, la tecnología de Ipod reciente, pero falta un paso muy importante que consiste en interpretar el contenido e información de estas imágenes digitales. En la actualidad existe gran variedad de software para el reconocimiento de caracteres pero sin embargo poseen una margen de error alto cuando la entrada es un documento de caracteres manuscritos. Debido a esto se realizó este trabajo en aras de conocer todo el procedimiento principal que utiliza el Reconocimiento de Caracteres Manuscritos mientras digitaliza la imagen, además de que ya existen diversos métodos y sistemas para este proceso aunque aun no se ha logrado demasiada perfección en la práctica de los mismos. También se desarrolla uno de estos métodos que fue en este caso el que se escogió para explicar, pues fue el que despertó mayor interés y se conoce como El Sistema Automático de Caracteres Manuscritos que está basado en el uso de un dispositivo de exploración óptica que puede reconocer la letra impresa.

La técnica del Reconocimiento de Caracteres Manuscritos comenzó aplicándose en documentos para los cuales ninguna forma electrónica estaba disponible y a medida que la tecnología ha avanzado, sus aplicaciones han evolucionado. El reconocimiento de caracteres sigue siendo un problema complejo que tropieza con dificultades aún no resueltas y que son actualmente un objeto de numerosas investigaciones.

RECOMENDACIONES

Después de haberse investigado e indagado en el tema del Reconocimiento de Caracteres Manuscritos se puede decir que:

- Sería bueno que la ciencia trabajara más profundo aún, en el tema, debido a que las técnicas del Reconocimiento de Caracteres Manuscritos son por lo general muy costosas y no estaría de más proponerse un menor gasto de recursos.
- Seguir trabajando en el perfeccionamiento de las técnicas y métodos utilizados en el Reconocimiento de Caracteres Manuscrito ya que distan mucho de ser perfectos y aún en la actualidad poseen un margen de error muy grande.
- Que cada día la ciencia innove más dispositivos en los cuales se pone de manifiesto el Reconocimiento de Caracteres Manuscritos como son la nueva telefonía móvil, las computadoras de bolsillo que son con pantalla táctil, la última tecnología de Ipod, las cámaras digitales, escaners, fotocopadoras es decir todo aquello que cada día nos facilita más el trabajo, la comunicación, como también el despeje.

BIBLIOGRAFIA

Blasco L pez, Antonio y F lez Esteban, Francisco: Reconocimiento de Caracteres Manuscritos. sl, se, sf.

P rez, Juan Carlos y Vidal, Enrique y S nchez, Lourdes: Sistema Geom trico de Reconocimiento  ptico de Caracteres. Valencia, se, sf.

B ez Rojas, J.J y Guerrero, M.L y Conde Acevedo, J y Padilla Vivanco, A y Urcid Serrano, G: Segmentaci n de Im genes. M xico, se, 2004.

Martin, Marcos: T cnicas Cl sicas de Segmentaci n de Im genes. sl, se, 2002.

Rojas, Javier: Segmentaci n y Reconocimiento de Patrones. sl, se, 2005.

Toscano Medina, Karina y Nakano Miyatake, Mariko y S nchez P rez, Gabriel y P rez Meana, H ctor M y Yasuhara, Makoto: Reconocimiento de Caracteres Manuscritos, en Instituto Polit cnico Nacional, Cient fica, M xico, 9(003), 2005, p. 143–154.

ANEXOS

Anexo No 1: [1]

Aqu  van algunos ejemplos de las curvas B zier:

Curvas lineales de B zier

Dado los puntos P_0 y P_1 , una curva lineal de B zier es una l nea recta entre los dos puntos. La curva viene dada por la expresi n:

$$B(t) = P_0 + (P_1 - P_0)t = (1 - t)P_0 + tP_1, t \in [0,1].$$

Curvas Cuadr ticas de Bezier

Una curva cuadr tica de B zier es el camino trazado por la funci n $B(t)$, dados los puntos: P_0 , P_1 y P_2 ,

$$B(t) = (1 - t)^2 P_0 + 2t(1 - t)P_1 + t^2 P_2, t \in [0,1].$$

Curvas C bicas de Bezier

Cuatro puntos del plano o del espacio tridimensional, P_0 , P_1 , P_2 y P_3 definen una curva cubica de B zier. La curva comienza en el punto P_0 y se dirige hacia P_1 y llega a P_3 viniendo de la direcci n del punto P_2 . Usualmente no pasara por P_1 ni por P_2 . Estos puntos solo est n ah  para proporcionar informaci n direccional. La distancia entre P_0 y P_1 determina que longitud tiene la curva cuando se mueve hacia la direcci n de P_2 antes de dirigirse hacia P_3 .

La forma parametrica de la curva es:

$$B(t) = P_0(1 - t)^3 + 3P_1t(1 - t)^2 + 3P_2t^2(1 - t) + P_3t^3, t \in [0,1].$$

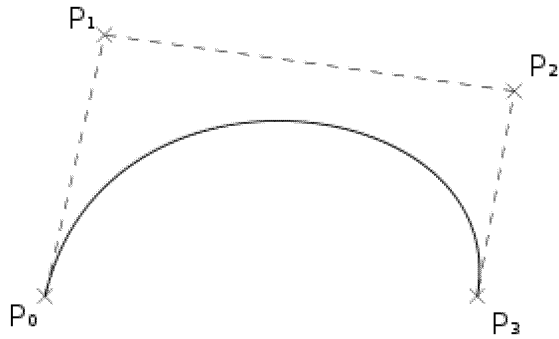


Fig. No 1: Curva Cúbica de Bézier donde se aprecian los puntos o nodos* de anclaje P1 y P2.

[1] Paul Bourke: Bézier curves, <http://local.wasp.uwa.edu.au/~pbourke/>

GLOSARIO

En esta sección va la definición de toda aquella palabra que fue marcada con * porque se entendió que podrían causar dificultades al conocer su significado.

- Entropía: Magnitud termodinámica que indica el grado de desorden molecular de un sistema.
- Distancia Euclídea: es la distancia normal que todos conocemos.
- Nodos: es un punto de intersección o unión de varios elementos que confluyen en el mismo lugar