

White paper – executive briefing

**Construya relaciones
rentables con sus clientes
usando Data Mining**

Herb Edelstein, President
Two Crows Corporation



Usted ha recopilado información acerca de sus clientes - ¿está haciendo buen uso de esta información?

Customer Relationship Management (CRM) ayuda a las compañías a mejorar la rentabilidad de las interacciones con sus clientes, mientras que a su vez, hace que estas interacciones sean más amigables gracias a la individualización. Para lograr el éxito implementando un proyecto de CRM, las compañías necesitan que sus productos y campañas correspondan a lo que sus prospectos y clientes quieren - en otras palabras, necesitan manejar de manera inteligente el ciclo de vida de sus clientes.

Hasta hace poco, la mayoría del software para CRM se enfocaba en la simplificación de la organización y en el manejo de la información de los clientes. Este software, llamado CRM operacional, se encarga de crear una base de datos de clientes que muestra una película consistente de la relación de cliente con la organización, y ofrece esta información en aplicaciones específicas. Esto incluye automatización de fuerza de ventas y aplicaciones del servicio al cliente en las cuales la compañía tiene contacto con el cliente.

De todas maneras, el aumento en el volumen de información almacenado acerca de los clientes y las cada vez más complejas relaciones con ellos han llevado al data mining a la vanguardia de convertir a las relaciones con los clientes en relaciones rentables. Data mining es un proceso que usa una variedad de técnicas de modelamiento y de análisis de datos para descubrir patrones y relaciones en los datos, que son usados para entender que es lo que sus clientes quieren y predecir que es lo que van a hacer. El data mining le ayuda a seleccionar los prospectos en los que su organización debe concentrar los esfuerzos, a ofrecer a sus clientes los productos adicionales adecuados y a identificar a los buenos clientes que pueden estar a punto de dejarlo. Esto trae como resultado una mejora en los ingresos, debido a que se mejora la capacidad de responder a cada contacto individual de la mejor manera posible, y a que se reducen los costos ya que se asignan los recursos de manera eficiente. Las aplicaciones de CRM que usan la minería de datos son llamadas CRM analítico

Este artículo describe diversos aspectos del CRM analítico y muestra cómo es usado para administrar el ciclo de vida de manera más efectiva. Hay que resaltar que los casos de estas compañías ficticias están basados en aplicaciones de data mining reales.

Data mining

El primer y más simple paso analítico en data mining es describir los datos. Por ejemplo, usted puede resumir los atributos estadísticos de los datos (como la media y las desviaciones estándar), revisar visualmente los datos usando tablas y gráficos y mirar la distribución de los valores de los campos de sus datos.

Pero la descripción por si sola no puede suministrar un plan de acción. Se debe construir un modelo predictivo basado en los patrones determinados a partir de los resultados ya conocidos y luego probar este modelo con los resultados por fuera de la

muestra original. Un buen modelo nunca debe confundirse con la realidad (usted sabe que un mapa no es una perfecta representación de la carretera real), pero puede ser una guía útil para entender su negocio.

Data mining puede ser usado para problemas de clasificación y regresión. En los problemas de clasificación, usted está pronosticando en que categoría cae cualquier cosa – por ejemplo, si una persona es buen o mal riesgo crediticio o cual de las diferentes ofertas es más probable que sea llamativa para alguien. En los problemas de regresión, usted está prediciendo un número, cómo cual es la probabilidad de que una persona responda a determinada oferta.

En CRM, el data mining es usado frecuentemente para asignar una calificación a un cliente o prospecto particular que indique la probabilidad de que un individuo se comporte de la manera que usted espera. Por ejemplo, la calificación podría medir la propensión a responder a una oferta particular, o la propensión a cambiarse a un producto/servicio ofrecido por su competencia. También es usado con frecuencia para identificar un grupo de características (llamado perfil) que segmente a los clientes en grupos con comportamientos similares, por ejemplo grupos que compren un producto en particular.

Un tipo especial de clasificación puede recomendar ítems basándose en la existencia de intereses similares de un grupo de clientes. Esto es llamado filtro colaborativo (collaborative filtering).

La tecnología de Data Mining, utilizada para resolver problemas de clasificación y filtro colaborativo, es descrita brevemente en el apéndice de este artículo.

Definiendo CRM

El CRM (Customer relationship management) a grosso modo simplemente significa administrar todas las interacciones con los clientes. En la práctica, esto requiere usar la información acerca de sus clientes y prospectos para interactuar más rentablemente con ellos en todas las etapas de su relación con ellos. Nos referimos acerca de estas etapas como el ciclo de vida del cliente.

El ciclo de vida de un cliente tiene tres etapas:

- Adquisición de clientes
- Aumento del valor de los clientes
- Retención de los buenos clientes

Data mining puede mejorar su rentabilidad en cada una de estas etapas cuando usted lo integra con los sistemas de CRM operacional o lo implementa como aplicaciones independientes.

Adquiriendo nuevos clientes a través del data mining

El primer paso en CRM es identificar los prospectos y convertirlos en clientes. Miremos como el data mining puede ayudar a administrar costos y a mejorar la efectividad de una campaña de adquisición de clientes.

Big Bank and Credit Card Company (BB&CC) anualmente lleva a cabo 25 campañas de correo directo, cada una de estas ofrece a un millón de personas la oportunidad de aplicar para obtener una tarjeta de crédito. La tasa de conversión mide la proporción de personas que se convierten en clientes de la tarjeta de crédito, la cual es de aproximadamente de uno por ciento por campaña para BB&CC.

Conseguir que las personas llenen la aplicación para la tarjeta de crédito es tan sólo el primer paso. Luego, el BB&CC debe decidir si el aplicante es un buen riesgo y aceptarlo como cliente o denegar esta solicitud. No es sorprendente que las personas que sean mal riesgo crediticio estén más interesados en aceptar la oferta que aquellos que sean un buen riesgo crediticio. Por lo tanto, el seis por ciento de la gente en la lista de envío respondió con la aplicación, tan sólo el 16 por ciento de estos son buen riesgo crediticio; aproximadamente uno por ciento de las personas en la lista de envío se convierten en clientes.

La tasa de respuesta de BB&CC de seis por ciento significa que tan solo 60,000 personas, de un millón de nombres, respondieron a la campaña. A menos que BB&CC cambie la naturaleza de la campaña - usando diferentes listas de correo, llegándole a los clientes de otra forma, cambiando los términos de la oferta - no va a recibir más de 60,000 respuestas. Y de estas 60,000 respuestas, tan sólo 10,000 tienen un perfil de riesgo lo suficientemente bueno para convertirse en clientes. El reto que el BB&CC afronta es llegarle a estas 10,000 personas de la manera más eficiente posible.

BB&CC gasta aproximadamente \$1.00 por pieza, para un costo total de \$1,000,000, para la campaña. En el transcurso de los próximos dos años, los clientes obtenidos a través de esta campaña generan aproximadamente \$1,250,000 en ganancias para el banco (o aproximadamente \$125 cada uno), para un retorno neto de \$250,000 del correo.

Data mining puede mejorar este retorno. Aunque el data mining no identificará exactamente los 10,000 eventuales clientes para la tarjeta de crédito, el data mining le ayudará a focalizar sus esfuerzos de mercadeo de manera costo - efectiva.

Como primera medida, BB&CC se envió un correo de prueba a 50,000 prospectos y cuidadosamente analizó los resultados, construyendo un modelo predictivo que muestre quien responderá (usando árboles de segmentación) y un modelo de credit scoring (usando redes neuronales). Luego, el BB&CC combinó estos dos modelos para encontrar las personas que eran buen riesgo crediticio y que además era muy probable que respondieran a la oferta.

BB&CC aplicó el modelo a las 950,000 personas restantes en la lista de correo, de las cuales 700,000 personas fueron seleccionadas para el envío. El resultado? De las 750,000 piezas enviadas (incluyendo los correos de prueba), BB&CC recibió 9,000 aplicaciones aceptables para tarjetas de crédito. En otras palabras, la tasa de respuesta se elevó de uno por ciento a 1.2 por ciento, un incremento de 20 por ciento. Mientras que los envíos focalizados le llegaron a 9,000 de los 10,000 prospectos - ningún modelo es perfecto - llegarle a los 1,000 prospectos

faltantes no es rentable. Si se hubieran enviado correos a las 250,000 personas restantes en la lista de envío, el costo de \$250,000 hubiera tenido como resultado otros \$125,000 de ganancia bruta, para una pérdida neta de \$125,000.

La siguiente tabla resume los resultados:

	Viejo	Nuevo	Diferencia
Número de piezas enviado	1,000,000	750,000	(250,000)
Costo del envío	1,000,000	750,000	(250,000)
Número de respuestas	10,000	9,000	(1,000)
Ganancia bruta por respuesta	\$125	\$125	\$0
Ganancia bruta	\$1,250,000	\$1,125,000	(\$125,000)
Ganancia neta	\$250,000	\$375,000	\$125,000
Costo del modelo	0	\$40,000	\$40,000
Ganancia final	\$250,000	\$335,000	\$85,000

Es importante resaltar que la ganancia neta del envío se incremento en \$125,000. Aun cuando usted incluya el costo de \$40,000 por concepto del costo del software para Data Mining y los computadores y los recursos humanos utilizados para este proyecto de modelamiento, la ganancia neta aumentó en \$85,000. Esto se translada a un retorno de la inversión (ROI) por el modelamiento de más de 200 por ciento, que excede de lejos el ROI requerido por el BB&CC's para un proyecto.

Incrementando el valor de sus clientes existentes:

Venta cruzada vía data mining

Cannons and Carnations (C&C) es una compañía que se especializa en vender morteros antiguos (antique mortars) y cañones que sirven como materas de flores para poner al aire libre. Adicionalmente ofrece una línea de materas de madera para flores hechas de pistolas antiguas de largo - calibre y una colección de mosquetes que se han convertido en floreros de flores de tallo largo. El catalogo de C&C es enviado aproximadamente a 12 millones de hogares.

Cuando un cliente llama a C&C para poner una orden, C&C identifica a la persona que llama usando un identificador de llamadas cuando es posible; de no ser así, el representante de ventas de C&C le pide el número telefónico o el número de cliente que aparece en el label del correo. A continuación, el representante de ventas mira el cliente en la base de datos y luego procede a tomar el pedido.

Cuando alguien llama a C&C, es una excelente oportunidad para la empresa para realizar una venta cruzada, o de vender algo adicional. Pero C&C descubrió que si la primera sugerencia falla y el representante de ventas sugiere un segundo ítem, el cliente se puede molestar y colgar el teléfono sin haber llegado a ordenar algo. Adicionalmente, hay algunos clientes que se disgustan cuando se intenta realizar una venta cruzada.

Antes de implementar el data mining, C&C era reacio a practicar la venta cruzada. Sin el modelo, las probabilidades hacer la

recomendación adecuada eran de uno a tres. Y ya que el hacer cualquier recomendación es inaceptable para algunos clientes, C&C quería estar completamente seguro de que nunca se hará una recomendación cuando no se debe. En una campaña de prueba, C&C tenía una tasa de venta de menos de uno por ciento y recibía un alto número de quejas. C&C se negaba a continuar con esta estrategia de venta cruzada para obtener una ganancia tan insignificante.

La situación cambio dramáticamente una vez C&C usó el data mining. Ahora el modelo de data mining opera sobre los datos. Al usar la información de los clientes que está almacenada en la base de datos y el nuevo orden, es más fácil para los representantes de ventas hacerle recomendaciones a los clientes. C&C vendió de manera exitosa un producto adicional a dos por ciento de los clientes y no tuvieron prácticamente queja alguna.

Desarrollar esta habilidad requirió de un proceso similar al que fue usado para resolver el problema de adquisición de clientes de tarjeta de crédito. Con esta situación, se necesitaban dos modelos.

El primer modelo pronostica si alguien se ofendería por recomendaciones adicionales de productos. C&C descubrió como reaccionan sus clientes llevando a cabo una corta encuesta telefónica. Para ser conservadores, C&C contó a cualquiera que hubiera declinado a participar en la encuesta, como alguien que hubiera considerado las recomendaciones como intrusivas. Más adelante, para verificar esta suposición, C&C hizo recomendaciones a un subgrupo pequeño, pero significativo de aquellos que se negaron a responder la encuesta. Para sorpresa de C&C, se detecto que la suposición no estaba justificada. Esto le permitió a C&C hacer más recomendaciones y mayores aumentos en las ganancias. El segundo modelo pronostico que oferta sería la que tendría mayor aceptación.

Resumiendo, data mining ayudo a C&C a entender mejor las necesidades de sus clientes. Cuando los modelos de data mining fueron incorporados en una campaña típica de CRM de venta cruzada, los modelos ayudaron a C&C a aumentar su rentabilidad en dos por ciento.

Incrementando el valor de sus clientes existentes:

Personalización vía data mining

Big Sam's Clothing (cuyo lema es: " Equipos resistentes al aire libre para los ciudadanos") desarrollaron un Web site para complementar su catalogo. Cada vez que usted ingresa a la página de Big Sam, el sitio lo saluda mostrando el letrero " Hola Sr. Pérez!" . Tan pronto usted se ha ordenado algo, o se ha registrado, usted será saludado por su nombre. Si usted tiene un registro de ordenes realizadas a través de la página, Big Sam también le dará información acerca de cualquier producto nuevo que pueda ser de su interés. Cuando usted mira un producto en particular, tal cómo una chaqueta impermeable Big Sam sugiere otro ítem que pueda complementar su compra.

Cuando Big Sam hizo el lanzamiento de su sitio, no había ningún tipo de personalización. La página era tan sólo una versión on-line de su catalogo impreso - realizado de manera agradable y

eficiente, pero no aprovechaba la oportunidad de venta que la página Web representa.

Data mining aumentó de manera considerable las ventas de la página Web de Big Sam. Los catálogos frecuentemente agrupan productos por tipo para simplificar la labor del cliente de seleccionar productos. En cambio, en un almacén on-line, los grupos de productos pueden ser bien diferentes, y estos grupos se basan en complementar el ítem que está bajo consideración. El sitio puede tener en cuenta no sólo el ítem que usted está buscando, sino también lo que aparece en su carro de compras, lo que conlleva a recomendaciones aún más personalizadas.

Como primera medida se usó clustering para descubrir que productos se agrupaban unos con otros de manera natural. Algunos de los clusters eran obvios, como camisetas y pantalones. Habían otros sorprendentes, como por ejemplo libros acerca de hiking en el desierto y kits contra mordeduras de serpientes. Ellos usaron estos agrupamientos para realizar recomendaciones cuando alguien se interesara en algún producto.

A continuación Big Sam construyó un perfil para ayudar a identificar los clientes que estarían interesados en los nuevos productos que son adicionados con frecuencia al catálogo. Big Sam aprendió que guiar a la gente hacia estos productos escogidos del catálogo, no sólo resultó en un incremento significativo en las ventas, sino también que fortaleció el relacionamiento con los clientes. Las encuestas establecieron que Big Sam era visto como un consejero de confianza de ropa y equipos.

Para extender aún más el alcance, Big Sam implementó un programa a través del cual los clientes podían elegir si recibían correos electrónicos acerca de nuevos productos que los modelos de data mining habían pronosticado que les interesarían. Mientras que los clientes percibían esto como otro ejemplo de mejora proactiva en el servicio, Big Sam descubrió que este era un programa de mejora en las ganancias.

El esfuerzo para lograr la personalización pagó el esfuerzo de Big Sam's, quienes experimentaron un aumento medible y considerable en las ventas repetidas, número de ventas promedio por cliente y valor promedio de las ventas.

Retención de los buenos clientes vía data mining

Para casi todas las compañías, el costo de adquirir nuevos clientes excede el costo de mantener a los buenos clientes. Este era un reto para KnowService, un Internet Service Provider (ISP) que tiene la tasa de abandono promedio de la industria, ocho por ciento por mes. KnowService tiene un millón de clientes, esto significa que 80,000 clientes se van cada mes. El costo de reemplazar estos clientes es de \$200 por cada uno, o de \$16,000,000 - incentivo suficiente para empezar un programa de administración de retención.

La primera cosa que KnowService hizo, fue preparar los datos que se iban a usar para pronosticar que clientes se iban a ir. KnowService tuvo que seleccionar las variables de su base de datos de clientes, y, probablemente transformarlos. El mayor volumen de los usuarios de KnowService son clientes de dial-in

(lo opuesto a los clientes que siempre están conectados a través de una línea T1 o DSL) por lo tanto KnowService sabe cuanto dura cada usuario conectado a la Web. KnowService también sabe el volumen de datos que se transfirieron desde y hacia el computadora de cada uno de los usuarios, el número de cuentas de correo electrónico que tiene un usuario, el número de mensajes de correo electrónico enviados y recibidos, así como el nivel de servicio y facturación registrado en el historial del cliente. Adicionalmente, KnowService tiene los datos demográficos que los usuarios suministraron al afiliarse.

A continuación, KnowService tuvo que identificar quienes eran los "buenos" clientes. Esto no es una pregunta de data mining, sino una definición de negocio (cómo rentabilidad o valor del ciclo de vida) seguido por un cálculo. KnowService construyó un modelo para perfilar a sus clientes rentables y a los no rentables. KnowService usó este modelo no sólo para retener a sus clientes, también para identificar a los clientes que no eran rentables en ese momento, pero que probablemente lo iban a ser en un futuro.

Luego KnowService construyó un modelo para pronosticar cuales de sus clientes más rentables lo abandonarían. Como en casi todos los problemas en los que se usa el data mining, determinar que datos usar y cómo combinar los datos existentes es un gran reto cuando se está desarrollando el modelo. Por ejemplo, KnowService necesitaba mirar datos de series de tiempo tales como uso mensual. En lugar de usar raw series de tiempo, se suavizaron los datos tomando rolling promedios de tres meses. KnowService también calculó el cambio en los promedios de cada tres meses y usó esto como predictor. Algunos de los factores que eran buenos predictores, como disminución en el uso, eran más síntomas que causas que podrían ser atacadas de manera directa. Otros predictores, como el número promedio de llamadas de servicio, y la variación en el promedio de llamadas de servicio, eran un indicador de problemas en la satisfacción de los clientes que era importante investigar.

Pronosticar quien abandonaría la empresa (churn) no era suficiente. Basándose en los resultados de modelamiento, KnowService identificó algunos programas y ofertas potenciales que se creía podrían incentivar a las personas a quedarse. Por ejemplo, algunos churners estaban excediendo la mayor cantidad de uso disponible para una tarifa fija mensual y estaban pagando excedentes a estas tarifas por el mayor uso. KnowService le ofreció a estos usuarios un servicio con una tarifa fija mayor que incluía más tiempo de conexión. A algunos usuarios les fue ofrecido más espacio en el disco duro para guardar páginas Web personales. Luego KnowService construyó modelos que pronosticarían cual sería la oferta más llamativa para un usuario en particular.

Para resumir, el proyecto de churn usó tres modelos. Un modelo identificaba las personas con más probabilidad de abandono, el siguiente modelo escogió a las personas rentables que probablemente iban a abandonar a la organización, pero que valía la pena mantener y el tercer modelo cruzó a los que probablemente abandonarían con la oferta más apropiada. El resultado neto fue una reducción en la tasa de abandono de KnowService, pasó de ocho por ciento a 7.5 por ciento, lo que le

permitió a KnowService ahorrar \$1,000,000 por mes en costos de adquisición de clientes.

KnowService descubrió que su inversión en data mining si sirvió - se mejoró el relacionamiento con los clientes y aumentó sustancialmente la rentabilidad.

Aplicando el data mining al CRM

Para poder construir buenos modelos para su sistema de CRM, hay ciertos pasos que usted debe seguir. El proceso de modelamiento de data mining de The Two Crows descrito a continuación es similar a otros modelos de proceso como el modelo de CRISP-DM, la diferencia radica principalmente en el énfasis que se le da a los diferentes pasos.

Tenga en mente que aunque los pasos aparecen en una lista, el proceso de data mining no es lineal - inevitablemente usted necesitará retornar a pasos previos. Por ejemplo, lo que usted aprende en el paso de la "exploración de datos" (paso 3) puede requerir que usted adicione nuevos datos a la base de datos de data mining. Los modelos iniciales que usted crea pueden ofrecerle pistas que lo lleven a crear nuevas variables.

Los pasos básicos de data mining para un CRM efectivo son:

1. Definir un problema de negocios
2. Construir un base de datos de mercadeo
3. Explorar los datos
4. Preparar los datos para el modelamiento
5. Construir el modelo
6. Evaluar el modelo
7. Desplegar el modelo y sus resultados

Repasemos estos pasos uno a uno para entender mejor el proceso.

1. Definición del problema de negocio. Cada aplicación de CRM tiene más de un objetivo de negocios para el que usted necesita construir un modelo apropiado. Dependiendo de su objetivo específico, como por ejemplo "aumentar la tasa de respuesta" o "incrementar el valor de una respuesta" usted construye un modelo muy diferente. Una definición efectiva del negocio incluye una forma de medir los resultados de su proyecto de CRM.

2. Construir una base de datos de marketing. Los pasos del dos al cuatro constituyen la clave de la preparación de los datos. Juntos, estos pasos toman más tiempo y esfuerzos que todos los demás pasos combinados. Pueden existir iteraciones repetidas repeated iterations de los pasos de preparación de los datos y la construcción del modelo ya que usted puede encontrar algo en el modelo que le sugiere que modifique los datos. Los pasos que involucran la preparación de los datos pueden tomar entre del 50 al 90 por ciento del tiempo y esfuerzo para el proceso total de data mining.

Usted necesitará construir una base de datos de marketing ya que sus bases de datos operacionales y data warehouse corporativos frecuentemente no contienen los datos que usted necesita en la forma que usted necesita. Adicionalmente, sus aplicaciones de CRM pueden interferir con la velocidad y ejecución efectiva de los sistemas.

Cuando se construye la base de datos de marketing usted necesita limpiarla - si usted quiere buenos modelos debe tener datos limpios. Los datos que necesita, pueden estar en diversas bases de datos como por ejemplo la de clientes, la de producto, o en la base de datos transaccional. Esto significa que usted necesita integrar y consolidar los datos en una base de datos de marketing unificada y conciliar las diferencias en los valores de los datos provenientes de diversas fuentes. Los datos que no se concilian adecuadamente son una fuente enorme de problemas en la calidad. En algunas ocasiones hay grandes diferencias en la manera en la que se definen y se usan los datos en las diferentes bases de datos. Algunas inconsistencias pueden ser fáciles de descubrir, cómo diferentes direcciones para un mismo cliente. Sin embargo, estos problemas son con frecuencia difíciles de detectar. Por ejemplo, el mismo cliente puede tener diferentes nombres, o lo que puede ser peor, múltiples números de identificación de cliente.

3. Exploración de datos. Antes de poder construir un buen modelo predictivo, usted debe entender sus datos. Empiece recolectando resúmenes numéricos (lo que incluye estadísticas descriptivas, como promedios, desviaciones estándares y demás) y mirando la distribución de los datos. Usted puede generar cross tabulations (tablas de pivoteo), para los datos multidimensionales.

Las herramientas de graficación y visualización son una ayuda vital en la preparación de los datos, y su importancia para lograr un análisis efectivo no debe ser subestimada. La visualización de los datos generalmente ofrece el "¡Ajá!" que lleva a nuevos descubrimientos y finalmente al éxito. Algunas formas de Some common and very useful graphical data displays are histograms or box plots which display distributions of values. You may also want to look at scatter plots in two or three dimensions of different pairs of variables. The ability to add a third, overlay variable greatly increases the usefulness of some types of graphs.

4. Prepare los datos para el modelamiento. Este es el último paso de la preparación de los datos, antes de empezar a construir modelos y es el paso en el que el "arte" se ve involucrada. Hay cuatro partes importantes involucradas en este paso:

Primero, seleccione las variables en las que se basará para construir el modelo. El escenario ideal es que usted tome todas las variables que tiene, las alimente en la herramienta de minería de datos y deje que esta herramienta encuentre a aquellos que son los mejores predictores. En la practica, esto es muy complicado. Una de las razones, es que el tiempo que se necesita para construir un modelo aumenta con el número de

variables involucradas. Otra de las razones es que incluir sin cuidado columnas extrañas puede llevar a obtener modelos con menos, en lugar de más, poder predictivo.

El siguiente paso es construir nuevos predictores que provengan de los datos sin modificar (raw). Por ejemplo, pronosticar el riesgo crediticio usando la proporción de deuda-ingreso y no únicamente la deuda y el ingreso como variables predictoras puede llevar a obtener resultados más acertados y además más fáciles de entender.

A continuación, usted puede seleccionar un subgrupo o una muestra de sus datos con la cual construirá modelos. Si usted tiene bastantes datos, usar todos sus datos puede tomar mucho tiempo o puede necesitar la compra de un computador más grande de los que usted puede querer. Al trabajar con una muestra seleccionada al azar de manera apropiada, generalmente resulta en la no-pérdida de información para la mayoría de problemas de CRM. Teniendo la opción de investigar unos cuantos modelos contruidos con todos los datos disponibles, o investigar más modelos contruidos a partir de una muestra, el último enfoque generalmente ayuda a desarrollar un modelo más robusto y preciso del problema.

Por último, se necesita transformar variables de acuerdo a los requerimientos del algoritmo que se eligió para construir su modelo.

5. Construcción del modelo de data mining. Lo más importante para recordar acerca de la construcción es que es un proceso iterativo. Usted necesita explorar modelos alternos para encontrar el más útil para resolver su problema de negocios. Lo que usted aprende cuando está buscando un buen modelo puede llevarlo a devolverse y hacer algunos cambios a los datos que se están usando, o incluso modificar la descripción del problema.

La mayoría de las aplicaciones de CRM están basadas en un protocolo llamado aprendizaje supervisado. Usted empieza con información de clientes para la cual el resultado esperado ya se conoce. Por ejemplo, usted puede tener datos históricos provenientes de una lista de correo anterior que es muy similar a la que se está usando en la actualidad. O, usted puede haber llevado a cabo una prueba de correo para determinar como responderá la gente a una oferta. A continuación usted divide estos datos en dos grupos. En el primer grupo, usted entrena o estima su modelo. Luego, usted prueba los datos restantes. Un modelo está construido cuando el ciclo de entrenamiento y prueba es completado.

6. Evalúe sus resultados. Probablemente el parámetro más utilizado para evaluar sus resultados es la exactitud. Suponga que tiene una oferta a la que únicamente uno por ciento de la gente respondió. Un modelo que pronostica que "nadie responderá" tiene una precisión del 99 por ciento y un 100 por ciento de inutilidad. Otra medida que es usada con frecuencia es el levantamiento (lift). El levantamiento (Lift) mide la mejora alcanzada por un modelo predictivo. Sin embargo, lift no tiene en cuenta costos ni ingresos, por lo que es más recomendable mirar las ganancias o el retorno de la inversión (ROI). Dependiendo de si se escoge maximizar el lift, las ganancias o

el retorno de la inversión (ROI), usted elige un porcentaje diferente de su lista de correo a los cuales les enviará información.

7. Incorporando el data mining a su solución de CRM. Al construir una aplicación de CRM

el data mining es con frecuencia una parte reducida, aunque crítica, del resultado final. Por ejemplo, los patrones predictivos, gracias al data mining, pueden ser combinados con el conocimiento de expertos e incorporados en aplicaciones grandes, usadas por diferentes tipos de personas.

La forma en la que el data mining es incorporado a la aplicación es determinado por la naturaleza de su interacción con clientes. Existen dos maneras de interactuar con clientes: ellos lo contactan a usted (inbound) o usted los contacta a ellos (outbound). Los requerimientos de deployment son bastante diferentes.

Las interacciones tipo outbound se caracterizan por que su compañía es la que origina el contacto, como por ejemplo campañas de correo directo. Usted selecciona las personas que está interesado en contactar, aplicando el modelo a su base de datos de clientes. Otro tipo de campaña outbound es una campaña publicitaria. En este caso, usted cruza los perfiles de los buenos prospectos mostrados por el modelo, con los perfiles de las personas a las que le llegara su aviso publicitario.

Para transacciones tipo inbound, como por ejemplo una compra telefónica, una orden puesta a través de Internet o una llamada a su línea de servicio al cliente, la aplicación debe responder en tiempo real. Por lo tanto, el modelo de data mining está incrustado en la aplicación y recomienda de manera activa acciones a seguir.

En cualquier caso, un punto clave con el que usted se debe enfrentar al aplicar el modelo a datos nuevos son las transformaciones que usted utilizó al construir el modelo. Es por esto que si los datos ingresados (de la transacción o de la base de datos) tienen los campos de edad, ingreso, sexo, pero el modelo requiere la proporción edad-ingreso y la variable género se ha modificado a dos variables binarias, usted debe transformar los datos ingresados de acuerdo a esto. La facilidad con la que usted pueda incluir estas transformaciones se convierte en uno de los factores más importantes en cuanto a productividad se refiere cuando usted quiere desplegar varios modelos de manera ágil.

Conclusiones

La implementación de CRM (Customer Relationship Management) es de vital importancia para competir efectivamente en el mercado actual. Entre más efectivamente se use la información que tiene de sus clientes, para satisfacerlos, más rentable será. (Este artículo no se concentra en el tema de las necesidades de los clientes. ¿Cuáles son sus necesidades? ¿Les hemos preguntado? ¿O simplemente les estamos sugiriendo productos basándonos en su comportamiento pasado?). El CRM Operacional necesita el CRM Analítico con los modelos predictivos de data mining como su base. El camino para un negocio exitoso necesita que usted

entienda a sus clientes y sus requerimientos, y el data mining es una guía esencial para esto.

Apéndice: tecnología para data mining

Arboles de segmentación

Los arboles de segmentación ofrecen una manera de representar una serie de reglas que llevan a una clase o valor. Por ejemplo, usted quiere ofrecerle a un posible cliente un producto en especial. El gráfico muestra un árbol de segmentación sencillo que soluciona el problema, mientras que ilustra todos los componentes básicos de los árboles de segmentación: el nodo de decisiones, las ramas y las hojas.

El primer componente de un árbol de clasificación es el nodo principal, o nodo-raíz, el cuál especifica la prueba que se llevará a cabo. Cada rama llevará a otro nodo, o al final del árbol, llamado nodo - hoja. Al navegar en el árbol de segmentación, usted puede asignar valores o clases a un caso, decidiendo que rama tomar, empezando en el nodo raíz y moviéndose a cada nodo subsecuente, hasta que se llega al nodo - hoja. Cada nodo usa los datos del caso para elegir la rama apropiada. Los arboles de segmentación son usados comúnmente en data mining para examinar los datos, generar un árbol y las reglas correspondientes, para de esta forma hacer predicciones. Una variedad de algoritmos pueden ser usados para construir arboles de segmentación, como por ejemplo CHAID (Chi-squared Automatic Interaction Detection), CART (Classification And Regression Trees), Quest y C5.0.

Redes Neuronales

Las redes neuronales son particularmente interesantes ya que ofrecen la manera de modelar de manera eficiente problemas largos y complicados, en los que puede haber cientos de variables predictoras que pueden tener múltiples interacciones. Las redes neuronales son usadas con frecuencia para regresiones, pero pueden ser usadas también en problemas de clasificación.

Las redes neuronales difieren en la filosofía de algunos de los métodos estadísticos en varias formas. Primero, una red neuronal usualmente tiene más parámetros que un modelo estadístico típico. Por ejemplo, una red neuronal con 100 inputs y 50 nodos escondidos tendrá más de 5,000 parámetros. Ya que hay tantos parámetros y tantas combinaciones de parámetros resultan en predicciones similares, los parámetros se vuelven difíciles de interpretar y la red sirve como una "caja negra" predictora. De cualquier modo, esto es aceptable en aplicaciones de CRM. Un banco puede asignar la probabilidad de banca rota a una cuenta, y no se preocupa por averiguar las causas.

Uno de los mitos de las redes neuronales, es que cualquier dato, sin importar la calidad, puede ser usado para obtener predicciones razonables, y examinar los datos para encontrar la verdad. Sin embargo, las redes neuronales requieren de tanta preparación, como con cualquier otro método, es decir que requieren bastante preparación de datos. Las implementaciones más exitosas de redes neuronales (o arboles de decisión, o regresión logística, o cualquier otro método) requieren de una limpieza, selección, preparación y preprocesamiento, cuidadoso

de los datos. Por ejemplo, las redes neuronales requieren que todas las variables sean numéricas. Es por esto que las variables categóricas como "estado" se codifican en múltiples variables dicotómicas (ej.: "California," "New York"), cada una con un valor de "1" (si) o "0" (no). La multiplicación de variables, como consecuencia de esto, se llama explosión categórica.

Clustering

Clustering divide las bases de datos en grupos diferentes. La meta del clustering es encontrar grupos que son muy diferentes unos de otros, y cuyos miembros son muy similares unos de otros. A diferencia de la clasificación, al empezar, usted no sabe que clusters van a aparecer, o por que tipo de atributos se aglomeraran los datos. Como consecuencia, los clusters deben ser interpretados por una persona que tenga conocimiento del negocio. Luego de haber hallado los clusters que segmentan de manera razonable su base de datos, estos clusters pueden ser usados para clasificar datos nuevos.

Es importante no confundir clustering con segmentación. La segmentación hace referencia al problema general de identificar grupos que tienen características comunes. Clustering es una forma de segmentar los datos en grupos que no se habían definido con anterioridad, mientras que la clasificación es una manera de segmentar datos, asignándolos en grupos que ya se habían definido previamente.

Data Mining hace la diferencia

SPSS Inc. ofrece soluciones que descubren lo que sus clientes quieren y predicen lo que van a hacer a continuación. La compañía ofrece soluciones que se encuentran en la intersección del CRM (Customer Relationship Management) y la inteligencia de negocios, las cuales les permiten a los clientes interactuar con los clientes de manera más rentable. Las soluciones de SPSS integran y analizan datos de mercadeo, de clientes y datos operacionales en mercados verticales claves alrededor del mundo, esto incluye entre otros: telecomunicaciones, salud, sector financiero, manufactura, retail, investigación de mercados y sector público. Con base en Chicago, SPSS tiene más de 40 oficinas alrededor del mundo.

Contáctenos

Para ponerse en contacto con SPSS:

- Visítenos en Internet: www.spss.com/la

SPSS es una marca registrada y otros productos de SPSS nombrados son registrados por SPSS Inc. Los nombres son marcas registradas de sus respectivos dueños