



Revista Cubana de Ciencias Informáticas

ISSN: 1994-1536

ISSN: 2227-1899

Editorial Ediciones Futuro

Salazar Gómez, Lisset; Vazquez Sánchez, Angel Alberto; Cañizares González, Roxana
Estimar conocimiento latente en grandes volúmenes de
datos utilizando el algoritmo Bayesian Knowledge Tracing
Revista Cubana de Ciencias Informáticas, vol. 15, Esp., 2021, Octubre-Diciembre, pp. 165-180
Editorial Ediciones Futuro

Disponible en: <https://www.redalyc.org/articulo.oa?id=378370462011>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org



Sistema de Información Científica Redalyc
Red de Revistas Científicas de América Latina y el Caribe, España y Portugal
Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso
abierto

Tipo de artículo: Artículo original
Temática: Programación paralela y distribuida
Recibido: 30/06/2021 | Aceptado: 01/10/2021

Estimar conocimiento latente en grandes volúmenes de datos utilizando el algoritmo Bayesian Knowledge Tracing

Estimating latent knowledge in large volumes of data using the Bayesian Knowledge
Tracing algorithm

Lisset Salazar Gómez ^{1*} <https://orcid.org/0000-0003-2159-0209>

Angel Alberto Vazquez Sánchez ¹ <https://orcid.org/0000-0002-3130-7983>

Roxana Cañizares González ¹ <https://orcid.org/0000-0002-3715-7103>

¹ Universidad de las Ciencias Informática. Km 2 ½ carretera de San Antonio, Torrens, La Habana.
{lsgomez, aavazquez, rcanizares}@uci.cu

*Autor para la correspondencia. (lsgomez@uci.cu)

RESUMEN

Ante la masividad de los datos que se generan en la educación, se han tenido que cambiar los métodos tradicionales para el descubrimiento de conocimientos. Uno de los algoritmos es el Bayesian Knowledge Tracing (BKT) que permite Estimar Conocimiento Latente (ECL). La ECL no es más que la forma de medir el conocimiento de un estudiante sobre habilidades y conceptos específicos, que es evaluada por sus patrones de corrección sobre esas habilidades. Dicho algoritmo está diseñado para ser utilizado en volúmenes de datos pequeños, afectándose su rendimiento ante la presencia de grandes volúmenes de datos. Para dar solución al

problema se presentará como resultado la transformación del algoritmo BKT teniendo en cuenta la programación paralela y distribuida. Se utilizaron herramientas de procesamiento en paralelo como el marco de trabajo Apache Spark en un entorno de minado. Se valida la propuesta de solución mediante pruebas para medir rendimiento y eficacia, usando métricas como speedup, eficiencia, error medio cuadrático del diferencial de probabilidades y error medio cuadrático del diferencial del área bajo la curva ROC; para las pruebas fueron empleadas bases de datos educacionales.

Palabras clave: Estimación del Conocimiento Latente (ECL); Minería de Datos Educacionales (MDE); Rastreo del Conocimiento Bayesiano (BKT).

ABSTRACT

Given the massive amount of data generated in education, the traditional methods for knowledge discovery have had to be changed. One of the algorithms is the Bayesian Knowledge Tracing (BKT) that allows to Estimate Latent Knowledge (LKE). LKE is nothing more than a way of measuring a student's knowledge about specific skills and concepts, which is evaluated by his or her correction patterns on those skills. This algorithm is designed to be used on small volumes of data, affecting its performance in the presence of large volumes of data. In order to solve the problem, the transformation of the BKT algorithm will be presented as a result, taking into account parallel and distributed programming. Parallel processing tools such as the Apache Spark framework were used in a mining environment. The proposed solution is validated through tests to measure performance and effectiveness, using metrics such as speedup, efficiency, mean squared error of the differential of probabilities and mean squared error of the differential of the area under the ROC curve; educational databases were used for the tests.

Keywords: Latent Knowledge Estimation (LKE); Educational Data Mining (EDM); Bayesian Knowledge Tracking (BKT).

Introducción

La aplicación de la minería de datos en la educación es un campo de investigación interdisciplinario emergente, también conocido como minería de datos (Romero y Ventura, 2013) . Al análisis y exploración de grandes volúmenes de datos en el contexto educativo es llamado Minería de Datos Educacional (MDE) (Martinez Torres et al., 2014) , que tiene como objetivo promover nuevos descubrimientos y avances en el terreno educativo mediante el uso de la información almacenadas en las plataformas educativas.

En (Romero y Ventura, 2013) , (Romero y Ventura, 2010) , (Baker y Yacef, 2009) y (Romero et al., 2010) la definen como una disciplina emergente, preocupada por el desarrollo de métodos para explorar los tipos únicos de datos que provienen de entornos educativos y utilizan esos métodos para comprender mejor a los estudiantes y los entornos en los que aprenden. Un análisis realizado en el área, arroja que el número de publicaciones en la MDE ha aumentado exponencialmente desde el 2005, siendo el mayor pico en el 2014 (Aristizábal Fúquene, 2017) , manteniéndose los avances de investigación hasta la fecha. Específicamente se ha dedicado un gran número de las investigaciones hacia la predicción, el cual permite predecir el rendimiento de un estudiante (Romero y Ventura, 2010) .

Dentro de la MDE, se utilizan diferentes técnicas de minería de datos, de ellas están las predictivas que tienen como objetivo predecir comportamiento de un aspecto de los datos (Larusson y White, 2014) . Estas pueden utilizarse para predecir si el alumno posee una determinada competencia sobre una habilidad. A esto último es lo que se conoce como Estimación de Conocimiento Latente (ECL), llamado así porque el conocimiento no es una variable directamente observable (Martinez Torres et al., 2014).

El área de la ECL es de particular importancia dentro de la MDE, debido a que aumentar el conocimiento de los estudiantes es la meta primaria de la educación. Por tanto, si el conocimiento puede ser medido, puedes saber dónde los estás haciendo mejor, puedes informar a los instructores (o cualquier otro interesado en el proceso) sobre el mismo y además puedes realizar decisiones pedagógicas automáticas (Baker y Corbett, 2014) (Larusson y White, 2014) .

Existen diferentes métodos que permiten ECL que surgen en el ambiente de la MDE (Feng y Heffernan, 2007) , entre los que se encuentran Análisis de los Factores de Rendimiento (PFA, *Performance Factors Analysis*), Teoría de Respuesta al Ítem (IRT, *Item Response Theory*) y el Rastreo del Conocimiento Bayesiano (BKT, *Bayesian Knowledge Tracing*). Cada uno de los algoritmos trata la capacidad de estimar el conocimiento latente de diferentes maneras (Larsson y White, 2014) . En el caso del IRT, predice la respuesta de una persona ante un ítem determinado. El PFA no es una expresión directa de la cantidad de habilidad latente, excepto por la probabilidad de responder correctamente. Por otro lado, el BKT expresa la probabilidad de que el estudiante domine la habilidad latente, y además muestra la probabilidad de que el estudiante responda correctamente la próxima vez que enfrente un problema donde deba aplicar dicha habilidad. Por lo que se diferencian en la manera de ECL, siendo el BKT el único algoritmo que puede determinar la probabilidad de dominio sobre una habilidad.

El BKT determina en qué medida un estudiante conoce una determinada aptitud o habilidad a partir de su rendimiento pasado con esa habilidad. Proporciona un conocimiento sobre habilidades de un sistema y predice comportamientos futuros sobre dichas habilidades, en pos de mejorar los sistemas de enseñanza-aprendizaje (Feng y Heffernan, 2007) . Esta información resulta de gran utilidad para determinar en qué medida una plataforma educativa cumple con su objetivo, para informar a los profesores o incluso para realizar acciones correctoras pedagógicas de manera automática (Martinez Torres et al., 2014).

Dicho algoritmo está diseñado para ser utilizado en volúmenes de datos pequeños, afectándose su rendimiento ante la presencia de grandes volúmenes de datos (Carlos Márquez Vera, 2012) , (Ballesteros Román, Sánchez-Guzmán y García Salcedo, 2013) , (Pardos et al., 2013) . Por lo que se plantea como objetivo de la investigación adaptar el algoritmo Bayesian Knowledge Tracing utilizando técnicas de programación paralela y distribuida, para disminuir los tiempos de ejecución manteniendo la eficacia en la estimación del conocimiento latente en datos educacionales masivos.

Métodos

Algoritmo BKT

Para predecir $P(L_j)$ para un estudiante individual se puede emplear el algoritmo de seguimiento del conocimiento. En este caso se busca $P(L_j \vee O_j)$, la probabilidad de que el estudiante ha aprendido la habilidad justamente luego de completar el paso j dado el rendimiento del estudiante O_j en los pasos previos, donde $O_j = \{o_1, o_2, \dots, o_j\}$ es el rendimiento del estudiante en las primeras j oportunidades y o_i puede ser correcto o incorrecto (van De Sande 2013)□. Estas probabilidades condicionales responden a la recurrencia:

$$P(L_{j-1}|O_j) = \frac{P(L_{j-1}|O_{j-1})(1-P(S))}{P(L_{j-1}|O_{j-1})(1-P(S)) + [1-P(L_{j-1}|O_{j-1})]P(G)}, O_j = \text{correcto} \quad (1)$$

$$P(L_{j-1}|O_j) = \frac{P(L_{j-1}|O_{j-1})P(S)}{P(L_{j-1}|O_{j-1})P(S) + [1-P(L_{j-1}|O_{j-1})](1-P(G))}, O_j = \text{incorrecto} \quad (2)$$

$$P(L_j|O_j) = P(L_{j-1}|O_j) + [1 - P(L_{j-1}|O_j)]P(T) \quad (3)$$

Existen una variedad de algoritmos que permiten realizar el ajuste de los parámetros necesarios para realizar la estimación del conocimiento latente del algoritmo BKT. Entre estos algoritmos se encuentran:

1. Maximizar la Expectación (EM, por sus siglas en inglés) (Dellaert, 2002) , (van De Sande, 2013) .
2. Fuerza Bruta (BF, por sus siglas en inglés) (d Baker, Corbett y Aleven, 2008) , (Yudelson et al., 2013).
3. Probabilidad Empírica (EP, por sus siglas en inglés) (Hawkins, Heffernan y Baker, 2014)
4. Ajuste Aleatorio (RF, por sus siglas en inglés).

5. Recocido Simulado (SAF, por sus siglas en inglés) (Miller, Baker y Rossi, 2014)
6. Baum-Welch (BW, por sus siglas en inglés) (Shen, 2008)

Características del algoritmo Bayesian Knowledge Tracing

El algoritmo BKT está regido por las expresiones 1 y 2. Estas expresiones en su forma funcional se pueden apreciar en la **Fig 1**. En la misma la línea puntuada representa la frontera entre soluciones crecientes y decrecientes. Desde que la curva “step j correct” está por encima de la línea puntuada, la ecuación 1 provoca que la secuencia converja al punto fijo en 1. De igual manera, desde que la curva “step j incorrect” está por debajo de la línea, la ecuación 1 causa que la secuencia $\{P(L_j \vee O_j)\}$ converge al punto fijo menor.

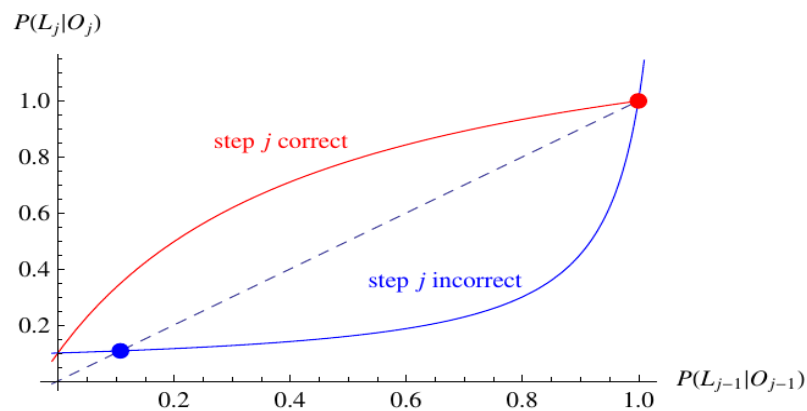


Fig 1 - Gráfico de la relación de recurrencia para el algoritmo BKT, de las ecuaciones (1) y (2). Los parámetros de este modelo son $P(S)=0.5$, $P(G)=0.3$, $P(T)=0.1$ y $P(L0) = 0.36$.

Fuente: (van De Sande 2013) .

Esta relación recurrencia (ecuaciones 1 y 2) no puede ser resuelta analíticamente, se puede aprender mucho de sus soluciones mediante la realización de un análisis de punto fijo (van de Sande, 2013) . El objetivo de un análisis de punto fijo es determinar la calidad del comportamiento de la secuencia $\{P(L_j \vee O_j)\}_{j=0}^n$ como una función de j. Si $P(L_j|O_j) > P(L_{j-1} \vee O_{j-1})$, entonces se puede decir que crece con j. Por otro lado, si $P(L_j|O_j) < P(L_{j-1} \vee O_{j-1})$, entonces se puede decir que decrece con j. Por lo que, la frontera entre

soluciones crecientes y decrecientes es $P(L_j | O_j) = P(L_{j-1} \vee O_{j-1})$, mostrada como la línea punteada en la fig. 1. Un punto fijo es un valor de $P(L_j \vee O_j)$ tal que la relación de recurrencia obedece $P(L_j | O_j) = P(L_{j-1} \vee O_{j-1})$.

En (van de Sande, 2013) también se definen dos tipos de puntos fijos:

1. Punto fijo estable: Si $P(L_j | O_j)$ está cerca del punto fijo, entonces $P(L_j | O_j)$ converge al punto fijo mientras j se incrementa.
2. Punto fijo inestable: Si $P(L_j | O_j)$ está cerca del punto fijo, entonces $P(L_j | O_j)$ se aleja del punto fijo al incrementarse j .

Aplicando estas ideas en las ecuaciones 1 y 2. En la ecuación 1 se encuentra un punto fijo estable en 1 y un punto fijo inestable en:

$$\frac{-P(G)P(T)}{1 - P(G) - P(S)} \quad (4)$$

Similarmente, la ecuación 2 tiene un punto fijo inestable en 1 y un punto fijo estable en:

$$\frac{(1 - P(G))P(T)}{1 - P(S) - P(G)} \quad (5)$$

Para que $P(L_j | O_j)$ permanezca en el intervalo $[0,1]$ para cualquier valor inicial de $P(L_0) \in [0,1]$ y cualquier secuencia de pasos O_j correctos/incorrectos, necesitamos que el punto fijo (5) esté en el intervalo $[0,1]$ y el punto inestable (4) se mantenga negativo. Esto nos da las siguientes restricciones en los valores permitidos para el modelo de los parámetros (Beck et al. 2008)□:

$$P(G) + P(S) < 1 \quad (6)$$

$$0 < P(T) < 1 - \frac{P(S)}{1 - P(G)} \quad (7)$$

Otras restricciones propuestas para este modelo son:

1. $P(S) < 0.5$ y $P(G) < 0.5$ (d Baker, Corbett y Alevan 2008)

2. $P(S) < 0.1$ y $P(G) < 0.3$ (Corbett y Anderson 1994)

La idea conceptual detrás de este algoritmo es:

1. Dominar una habilidad generalmente conlleva a un rendimiento correcto.
2. Un buen rendimiento implica que un estudiante domina la habilidad relevante.

Por tanto, mediante la búsqueda de dónde el estudiante muestra un buen rendimiento se puede inferir que domina la habilidad (S. Baker 2015).

Para el experimento se utilizarán diferentes conjuntos de datos públicos pertenecientes al repositorio de base de datos educacionales PSLC Datashop, disponible en <http://pslcdatashop.org>. Los datos estarán organizados de forma tal que en cada columna se tenga la siguiente información:

1. First attemp – Valor del resultado del primer intento del estudiante (1 – intento correcto, 0 – intento incorrecto).
2. Anon Student Id – Identificador del estudiante.
3. Concatenación de los campos Problem Hierarchy – Problem Name – Step Name. Identifica el problema donde fue aplicada la habilidad.
4. Knowledge Component – Habilidad a estimar.

Los conjuntos de datos cuentan con las siguientes características:

Tabla 1 - Características de los datasets.

Conjunto de datos	Número de estudiantes	Número de componentes de conocimiento	Número de problemas	Cantidad de filas
OSU, Honors Physics: Mechanics, Fall 2011 (dataset 1)	314	154	44	323,912
USNA Physics Fall 2006 (dataset 2)	66	250	251	345,536
ElemChinese (dataset 3)	221	22	868	812,329
Assistments Math 2006-2007 (5046 Students) (dataset 4)	5,046	318	1872	1,451,003

Para la realización de las pruebas se desplegó la propuesta de solución en un clúster de computadoras usando la herramienta Apache Spark en su versión 2.2.0. Se aplicó el modo independiente (Standalone) de despliegue, donde se utilizan las funciones nativas de la herramienta. El clúster estaba conformado por una estación de trabajo que actuó como máster y dos estaciones de trabajo utilizadas como nodos trabajadores. Se debe tener en cuenta que las pruebas fueron realizadas en un entorno no dedicado de nodos que pertenecen a la misma subred.

Para ejecutar el algoritmo secuencial se utilizó una computadora con las siguientes características de hardware: tipo de procesador Intel(R) Core(TM) i7-4790 CPU @ 3.60GHz, con 8 núcleos y memoria principal de 8 Gb DDR3 1600 y de software: sistema operativo Ubuntu 18.04 LTS 64 bits, versión del núcleo: 4.15.0-43-generic y software necesario Java OpenJDK versión "8" Update 192 y Spark 2.2.0.

Para las pruebas del algoritmo adaptado se utilizó un entorno de minado que tiene las siguientes características, en la estación de trabajo máster como hardware el tipo de procesador Intel(R) Core(TM) i7-

4790 CPU @ 3.60GHz, cantidad de núcleo 8 y memoria principal 8 Gb DDR3 1600. En las estaciones de trabajos usadas como nodos trabajadores se utilizaron dos clúster, que tienen como características de hardware, tipo de procesador Intel(R) Celeron(R) CPU G1830 @ 2.80GHz, cantidad de núcleo 2 cada una y memoria principal de cada una 4 Gb DDR3. La característica de la red de datos es 100 Mb/s.

Todos los clúster tienen como características de software, como sistema operativo Ubuntu 18.04 LTS 64 bits, versión del núcleo: 4.15.0-43-generic y de software instalado Java OpenJDK versión "8" Update 192 y Spark 2.2.0

En la prueba de los algoritmos secuencial y paralelo se utilizarán 4 conjunto de datos (dataset) de diferentes tamaños. Para el algoritmo secuencial se harán de 10 ejecuciones por cada conjunto de datos recogiendo el tiempo de ejecución y obteniendo la probabilidad de dominio de la habilidad por cada estudiante y el AUC (Área Bajo la Curva). Lo mismo se hará con el algoritmo paralelo, pero utilizando el entorno minado, con 10 ejecuciones por cada conjunto de datos, recogiendo el tiempo de ejecución, la probabilidad de dominio de la habilidad por cada habilidad y el AUC.

A partir de esos valores entonces se procederá al cálculo de aceleración (en lo adelante speedup) y eficiencia. Para la eficacia, se utilizará el Error Cuadrático Medio (ECM), para verificar que no exista diferencia significativa entre los resultados arrojados por el algoritmo secuencial y el paralelo. Las pruebas realizadas para el algoritmo adaptado con los conjuntos de datos seleccionados en el entorno de minado configurado, los valores obtenidos de aceleración y eficiencia permiten afirmar que el algoritmo adaptado presenta una mejora importante de tiempo de ejecución respecto al algoritmo secuencial.

Para comprobar la eficacia del algoritmo adaptado se tomará en cuenta dos métricas a medir que son el Error Cuadrático Medio (ECM) de los valores calculados de la probabilidad de dominio de las habilidades y la ECM del área bajo la curva ROC (AUC). Ambas métricas serán medidas inicialmente por el algoritmo secuencial y luego por el paralelo. Una vez realizada las pruebas se comprueba que no existen diferencias significativas entre los resultados para los diferentes conjuntos de datos. Lo que demuestra que la eficacia

de usar el algoritmo paralelo es similar a la obtenida al usar el algoritmo secuencial. Además, las diferencias entre la métrica de AUC para ambos algoritmos (secuencial y paralelo) nunca es mayor que 0,006598 lo que es un indicador de que la métrica se comporta de forma similar para ambos casos. Luego de realizado todas las pruebas se puede afirmar que los tiempos de ejecución son mejores para el algoritmo paralelo y la eficacia se mantiene con valores similares para las métricas utilizadas.

Resultados y discusión

El algoritmo propuesto utiliza los siguientes elementos para su funcionamiento:

1. Carga de datos: En este paso se realiza la carga de datos desde las diferentes fuentes que proporcionan datos.
2. Pre-procesar los datos: Se encarga de modificar el conjunto de datos para que estos estén listos para la ejecución del algoritmo propuesto.
3. Ajustar parámetros del algoritmo. El algoritmo propuesto utiliza un grupo de parámetros para su funcionamiento que necesitan ser ajustados antes de poder aplicarlos.
4. Ejecutar estimación en paralelo: Se realiza la estimación para todos los estudiantes en el conjunto de datos de forma paralela.
5. Visualizar resultados: Se visualizan y salvan los resultados obtenidos.

Carga de datos

En el primer paso, se tiene como opciones de orígenes de datos a:

1. Archivos CSV o TSV.
2. Hadoop Distributed File System (HDFS).

Por tanto, este procedimiento recibe como entrada el tipo de origen de los datos y los datos necesarios para cargarlo y consta de los siguientes pasos:

1. Configuración del SparkSession a través del objeto SparkConf.
2. A partir del SparkSession y del origen de datos, se obtiene el conjunto de datos (dataset) a utilizar.
3. El conjunto de datos obtenido es devuelto a la ejecución central del algoritmo.

Pre-procesar los datos

Para el correcto funcionamiento del algoritmo es necesario que el conjunto de datos posea la información referente al estudiante, el problema, la habilidad y el valor observado de respuesta a dicho problema. Por tanto, es preciso transformar el conjunto de datos obtenido en el paso anterior para obtener uno que tenga la estructura adecuada. Para ello se seguirán una serie de pasos para el pre-procesamiento de los datos (ver **Fig 2**).

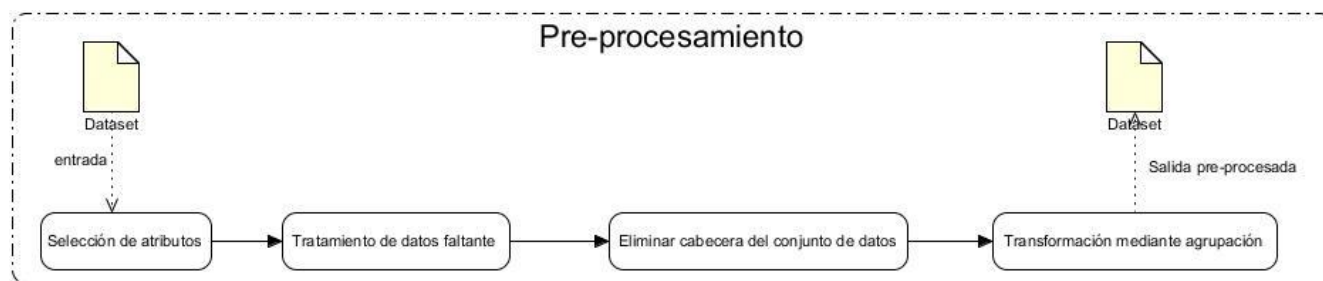


Fig.2 - Pre-procesamiento de los datos.

Fuente: Elaboración propia.

Este proceso de transformación el conjunto de datos debe tener el siguiente formato (Yudelsson et al., 2013):

1. id_observación – Observación de la respuesta (correcto, incorrecto).
2. id_estudiante – Identificación del estudiante.
3. id_problema – Concatenación de jerarquía del problema – nombre del problema – nombre del paso.
4. id_habilidad – Habilidad o componente de conocimiento a medir.

Ajustar parámetros del algoritmo

Para este proceso se hará el ajuste de parámetro con el par estudiante habilidad, donde se ejecutará en paralelo el proceso. Este algoritmo es un proceso de dos pasos que involucra anotar los datos del rendimiento con conocimiento y usar esta información para computar los parámetros de BKT (Hawkins, Heffernan y Baker 2014) .

Ejecutar estimación en paralelo

Una vez terminado el paso anterior del algoritmo, por cada par estudiante-habilidad se posee cuales parámetros se deben usar para calcular el conocimiento latente que posee el estudiante para cada habilidad.

Este paso del algoritmo recibe un conjunto de datos con las transacciones de todos los estudiantes-habilidades (estudiante, habilidad, secuencia de observaciones, parámetros con sus valores ajustados). De forma paralela y distribuida se calcula para cada par estudiante-habilidad cuál es el valor de la probabilidad de que se domine dicha habilidad por el estudiante. Una vez terminado este paso se tiene un conjunto de datos con llave estudiante-habilidad donde se posee el valor de la probabilidad de que el estudiante domine dicha habilidad.

Visualizar resultados

Este paso se encarga de visualizar los resultados derivados del paso anterior. La información se representará en diferentes tipos de gráficos representando conocimientos diversos. Para la visualización se utilizará la biblioteca gráfica JFreeChart^a.

Unos de los gráficos representados es el de habilidades a través de un gráfico de burbuja. En el mismo se visualizará por habilidad la cantidad de estudiantes, la cantidad de problemas y el promedio de dominio de la habilidad para todos los estudiantes (radio de la burbuja).

Una vez obtenido los datos a representar, entonces se visualiza el gráfico de burbuja. En el mismo se muestra en cada burbuja la etiqueta que representa la información sobre la cantidad de problema, la cantidad de estudiantes y el promedio de probabilidad.

Conclusiones

El algoritmo BKT adaptado cuenta con 5 pasos fundamentales: carga de los datos, pre-procesamiento, ajuste de parámetros, ejecutar el algoritmo y la visualización, el cual se ejecuta sobre un entorno minado distribuido sobre el marco de trabajo Apache Spark de forma paralela, obteniéndose las probabilidades por habilidades de cada estudiante a partir de la secuencia de observaciones por cada habilidad. Con la ejecución de los algoritmos sobre los 4 dataset de diferentes tamaños se pudo afirmar que el algoritmo adaptado presenta una mejora importante de tiempo de ejecución respecto a la aceleración y eficiencia con el algoritmo secuencial. Así mismo, la eficacia mantiene valores similares para las métricas del ECM utilizadas entre las probabilidades de dominio de las habilidades y el AUC.

Referencias

- Baker, R.S. Y Corbett, A.T., 2014. Assessment Of Robust Learning With Educational Data Mining. *Research & Practice In Assessment*, Vol. 9, Pp. 38-50.
- Baker, R.S.J.D. Y Yacef, K., 2009. The State Of Educational Data Mining In 2009: A Review And Future Visions. *Jedm| Journal Of Educational Data Mining*, Vol. 1, No. 1, Pp. 3-17.
- Ballesteros Román, A., Sánchez-Guzmán, D. Y García Salcedo, R., 2013. Minería De Datos Educativa: Una Herramienta Para La Investigación De Patroness De Aprendizaje Sobre Un Contexto Educativo. *Latin-American Journal Of Physics Education*, Vol. 7, No. 4.
- Beck, J.E., Chang, K., Mostow, J. Y Corbett, A., 2008. Introducing The Bayesian Evaluation And Assessment Methodology. *International Conference On Intelligent Tutoring Systems*. S.L.: S.N., Pp. 383-394.

- Corbett, A.T. Y Anderson, J.R., 1994. Knowledge Tracing: Modeling The Acquisition Of Procedural Knowledge. *User Modeling And User-Adapted Interaction*, Vol. 4, No. 4, Pp. 253-278.
- D Baker, R.S.J., Corbett, A.T. Y Alevan, V., 2008. More Accurate Student Modeling Through Contextual Estimation Of Slip And Guess Probabilities In Bayesian Knowledge Tracing. *International Conference On Intelligent Tutoring Systems*. S.L.: S.N., Pp. 406-415.
- Dellaert, F., 2002. The Expectation Maximization Algorithm. . S.L.:
- Carlos Márquez Vera, C.R., 2012. *Predicción Del Fracaso Escolar Mediante Técnicas De Minería De Datos*. 2012. S.L.: Obtenido De [Http://Rita. Det. Uvigo. Es/201208/Uploads/Ieee-Rita](http://Rita.Det.Uvigo.Es/201208/Uploads/Ieee-Rita).
- Feng, M. Y Heffernan, N.T., 2007. Towards Live Informing And Automatic Analyzing Of Student Learning: Reporting In Assistent System. *Journal Of Interactive Learning Research*, Vol. 18, No. 2, Pp. 207-230.
- Hawkins, W.J., Heffernan, N.T. Y Baker, R.S.J.D., 2014. Learning Bayesian Knowledge Tracing Parameters With A Knowledge Heuristic And Empirical Probabilities. *International Conference On Intelligent Tutoring Systems*. S.L.: S.N., Pp. 150-155.
- Larusson, J.A. Y White, B., 2014. *Learning Analytics: From Research To Practice*. S.L.: Springer.
- Martinez Torres, M. Del R., Gutiérrez Reina, D., Toral, S.L. Y Barrero, F., 2014. Metodologías De Análisis De Los Big Data En Las Plataformas Educativas. *Tae 2014: Xi Congreso De Tecnologías, Aprendizaje Y Enseñanza De La Electrónica: Libro De Actas, Bilbao 11-12 Y 13 De Junio De 2014 (2014)*, P 79-83. S.L.: S.N.,
- Miller, W.L., Baker, R.S. Y Rossi, L.M., 2014. Unifying Computer-Based Assessment Across Conceptual Instruction, Problem-Solving, And Digital Games. *Technology, Knowledge And Learning*, Vol. 19, No. 1-2, Pp. 165-181.
- Pardos, Z.A., Bergner, Y., Seaton, D.T. Y Pritchard, D.E., 2013. Adapting Bayesian Knowledge Tracing To A Massive Open Online Course In Edx. *Edm*, Vol. 13, Pp. 137-144.
- Romero, C. Y Ventura, S., 2010. Educational Data Mining: A Review Of The State Of The Art. *Ieee Transactions On Systems, Man, And Cybernetics, Part C (Applications And Reviews)*, Vol. 40, No. 6, Pp. 601-618.
- Romero, C. Y Ventura, S., 2013. Data Mining In Education. *Wiley Interdisciplinary Reviews: Data Mining And Knowledge Discovery*, Vol. 3, No. 1, Pp. 12-27.

Romero, C., Ventura, S., Pechenizkiy, M. Y Baker, R.Sj., 2010. *Handbook Of Educational Data Mining*. S.L.: Crc Press.

S. Baker, R., 2015. Big Data And Education. *Teachers College, Columbia University*.

Shen, D., 2008. Some Mathematics For Hmm. *Massachusetts Institute Of Technology*,

Van De Sande, B., 2013. Properties Of The Bayesian Knowledge Tracing Model. *Journal Of Educational Data Mining*, Vol. 5, No. 2, Pp. 1-10.

Yudelso, M. V, Koedinger, K.R. Y Gordon, G.J., 2013. Individualized Bayesian Knowledge Tracing models. *International conference on artificial intelligence in education*. S.l.: s.n., pp. 171-180.

Contribuciones de los autores

1. Conceptualización: Lisset Salazar Gómez, Angel Alberto Vazquez Sánchez.
2. Análisis formal: Lisset Salazar Gómez, Angel Alberto Vazquez Sánchez, Roxana Cañizares Gonzalez.
3. Investigación: Lisset Salazar Gómez, Angel Alberto Vazquez Sánchez.
4. Metodología: Lisset Salazar Gómez.
5. Software: Lisset Salazar Gómez, Angel Alberto Vazquez Sánchez.
6. Validación: Lisset Salazar Gómez.
7. Visualización: Lisset Salazar Gómez, Angel Alberto Vazquez Sánchez.
8. Redacción – borrador original: Lisset Salazar Gómez.

^{aa} <http://jfreechart.org/jfreechart>