

Detección de la polaridad en citas periodísticas: una solución no supervisada *

A non supervised method for sentiment polarity detection on reported speech from news

A. Montejo-Ráez, E. Martínez-Cámara, M. T. Martín-Valdivia, L. A. Ureña-López

Departamento de Informática

Universidad de Jaén

Las Lagunillas s/n, Jaén - 23071

{amontejo,emcamara,maite,laurena}@ujaen.es

Resumen: El presente trabajo expone los resultados alcanzados mediante un método no supervisado para la detección de la polaridad en textos relativos a citas aparecidas en noticias en inglés, correspondientes al corpus 2010_JRC_1590_quotes. Este método, basado en la obtención de un subgrafo de WordNet obtenido mediante el algoritmo Page Rank y su ponderación mediante los valores correspondientes en SentiWordNet, propone una solución no supervisada que ofrece unos resultados competitivos sobre algunas técnicas actuales.

Palabras clave: Análisis de emociones, clasificación de la polaridad, SentiWordNet, Page Rank

Abstract: This paper shows the results obtained by a non supervised method in the task of sentiment polarity detection on news quotations (reported speech) from the 2010_JRC_1590_quotes corpus. This method, which obtains a subgraph of WordNet by applying Page Rank over it, weights such a graph with corresponding SentiWordNet polarity scores, offering a non-supervised solution comparable to state-of-the-art techniques.

Keywords: Sentiment Analysis, polarity classification, SentiWordNet, Page Rank

1. Introducción

El análisis de opiniones (*Sentiment Analysis*) estudia el tratamiento de opiniones en textos. Este estudio, enmarcado dentro del Procesamiento del Lenguaje Natural (PLN), está suscitando un interés creciente en la comunidad investigadora por diversas razones. Principalmente, el análisis de la ingente cantidad de información que genera la Web 2.0, donde cada usuario en Internet es un generador potencial de información a través de comentarios, bitácoras, envíos por microblogging (como Twitter) y una larga lista de posibles escenarios en los que expresar una opinión con fuerte carga subjetiva. Poder analizar de forma automática estas publica-

ciones permite “tomar la temperatura” de un amplio grupo de usuarios acerca de temáticas diversas, o determinados productos. Es por esto que el análisis de opiniones (AO) es objeto de interés por parte de los investigadores como de las empresas, pues permite abrir una ventana a la opinión en la Red, mas allá de las encuestas explícitas.

La presencia de opiniones no solo se restringe a las plataformas 2.0 que pueblan Internet, sino que también se encuentran opiniones en ámbitos más profesionales, como puede ser el periodístico o crítica especializada (literaria, cinematográfica, gastronómica ...). El tratamiento de estas opiniones también interesa tanto al mundo académico como al empresarial, ya que debido a su carácter profesional pueden crear opinión, que luego se transmite a través de las redes sociales.

Dentro del AO, la clasificación de *polaridad* simplifica este problema como sigue: dado un texto t , el objetivo es obtener una puntuación que nos indique si el texto expresa

* Esta investigación ha sido subvencionada parcialmente por el proyecto del Instituto de Estudios Gienenses Geocaching Urbano (RFC/IEG2010), el Fondo Europeo de Desarrollo Regional (FEDER), a través del proyecto TEXT-COOL 2.0 (TIN2009-13391-C04-02) por el gobierno español, y por la Comisión Europea bajo el Séptimo programa Marco (FP7 - 2007-2013) a través del proyecto FIRST (FP7-287607)

una opinión positiva o negativa, dentro de un rango donde 0 indicaría una carga subjetiva neutra, 1 una carga subjetiva positiva y -1 una carga subjetiva negativa:

$$p(t) \in [-1, 0, 1] \quad (1)$$

El método propuesto es capaz de calcular este valor a partir de un texto sin necesidad de utilizar conocimiento del dominio, conjunto de datos etiquetados y modelos previamente entrenados, sino ponderando el uso conjunto de las bases de conocimiento de WordNet y SentiWordNet mediante los valores obtenidos por el algoritmo Page Rank (Page et al., 1998).

El artículo se estructura como sigue. Primero mencionamos algunos de los trabajos más recientes en clasificación de la polaridad. El siguiente punto describe el método propuesto y sus fases. Después pasamos a detallar la configuración de los experimentos y a comentar los resultados obtenidos. Finalmente, cerramos el trabajo con unas reflexiones finales e información sobre los aspectos que darán continuidad a esta investigación.

2. *Estado del arte en la clasificación de polaridad*

En la resolución de un problema mediante aprendizaje automático se pueden seguir principalmente dos estrategias: supervisada y no supervisada. Ésta dicotomía se encuentra muy presente en la literatura sobre AO, aunque las técnicas supervisadas están más representadas que las no supervisadas. El trabajo que se suele tomar como referencia para los métodos supervisados es (Pang, Lee, y Vaithyanathan, 2002), el cual utiliza como característica la presencia o no de los términos para el cálculo de la polaridad. En cambio, (Martínez-Cámara et al., 2011) utilizando el algoritmo SVM como (Pang, Lee, y Vaithyanathan, 2002) pero sobre un corpus de críticas de cine en español, obtiene mejores resultados representando los unigramas con el valor TF-IDF.

En la estrategia supervisada no solo se han utilizado características léxicas, sino que también se ha intentado explotar la información sintáctica. Un ejemplo de ello son los trabajos (Mullen y Collier, 2004) y (Whitelaw, Garg, y Argamon, 2005), en los que se utilizan los adjetivos como características a la hora de la clasificación de la polaridad. El tratamiento de la negación también se ha

tenido en cuenta en los experimentos supervisados, siendo (Wiegand et al., 2010) un buen resumen de lo publicado en el aprovechamiento de esta característica. No es el objetivo de este artículo ser un listado de todos los trabajos que siguen un enfoque supervisado, por lo que si el lector quiere una mayor información puede consultar (Pang y Lee, 2008), (Liu, 2010), (Tsytarau y Palpanas, 2012).

Los métodos supervisados adolecen de la necesidad de disponer de conjuntos de datos etiquetados, y de la dificultad de adaptación de las técnicas entre dominios distintos. Éstas son las principales razones que llevan a los investigadores a estudiar métodos no supervisados o semi-supervisados, que no requieren de datos etiquetados para la construcción de un modelo de clasificación, y que su aplicación a distintos dominios sea mucho más sencilla.

Los métodos no supervisados en AO suelen fundamentarse en la detección de identificadores de subjetividad u opinión en los textos, y posteriormente calculan la polaridad aplicando alguna función basada en los indicadores encontrados. Para dicha identificación se suelen utilizar conjuntos de vocablos etiquetados por su polaridad. Existen tres estrategias para la elaboración de estos lexicones: manual, basado en diccionario y basado en corpus. El enfoque manual es bastante costoso, por lo que se suele aplicar como último paso de los enfoques automáticos a modo de revisión. El método basado en diccionario construye un lexicón etiquetado por medio de la combinación de un conjunto de palabras de semilla, y el uso de recursos léxicos, como es el caso de WordNet, para la aplicación de la lista inicial (Kim y Hovy, 2004), (Hu y Liu, 2004). Los métodos basados en diccionario generan listas de indicadores de opinión genéricas, que no dan buenos resultados en dominios muy específicos. Los métodos basados en corpus se adaptan mejor al dominio en el que trabajan, ya que se basan en las características propias del corpus para la ampliación del lexicón de palabras semilla inicial (Hatzivassiloglou y McKeown, 1997), (Kanayama y Nasukawa, 2006).

Los métodos anteriores más que no supervisados deberían considerarse semi-supervisados, debido a que necesitan de un cierto conocimiento inicial, es decir, una lista de indicadores de opinión etiquetada

para calcular la polaridad. Además, aunque el método basado en diccionario es el más genérico, tiene una cierta dependencia de los términos semilla que se hayan seleccionado. Reduciendo al mínimo el conocimiento inicial, y por lo tanto ganando capacidad de adaptación del dominio, se encuentran los métodos que basan la clasificación en cálculos estadísticos, y en las relaciones semánticas de los términos. Como referencia de este enfoque se encuentra el trabajo de Turney en 2002 (Turney, 2002). Actualmente, algunos métodos (Wu et al., 2010) intentan resolver el problema de la adaptación al dominio mediante un mapeado de conceptos comunes, buscando un espacio dimensional más allá del contenido léxico (Ji et al., 2011), explorando conexiones entre grafos (Wu, Tan, y Cheng, 2009) o mediante métodos probabilísticos (Tan et al., 2009).

La metodología seguida en la generación del corpus Emitoblog propone la generación de modelos de alta granularidad para el AO sobre distintos tipos de textos, propios de publicaciones coloquiales en el marco de la Web 2.0 (Boldrini et al., 2012). Este trabajo demuestra que es posible la transferencia de dominio de modelos entrenados sobre ciertos conjuntos de datos a otros textos de dominio diferente, gracias a la independencia relativa que permiten características detalladas acerca de los sentimientos. En cualquier caso, es una solución que necesita de entrenamiento, por lo que en mayor o menor medida seguirá viéndose afectada por las limitaciones en la transferencia de sentimientos según dominio.

Por otra parte, SentiWordNet (Baccianella, Esuli, y Sebastiani, 2008) es un recurso léxico construido a partir de otro recurso: WordNet (Fellbaum, 1998). Proporciona información acerca de la orientación semántica en cuanto a emoción de los “synsets”. Un synset es el ítem de información básico en WordNet, y representa un “concepto”, sin ambigüedad. La mayoría de las relaciones en este grafo léxico trata los synsets como nodos: hiperonimia, sinonimia, homonimia, antonimia... SentiWordNet devuelve, para cada synset, una tripleta de tres valores que miden la carga de “positividad”, “negatividad” u “objetividad” del mismo. La última versión de SentiWordNet (3.0) se ha generado a partir de las anotaciones manuales de versiones previas, propagando sobre el grafo dichos val-

ores de emoción mediante un algoritmo de tipo *random walk*. Este recurso ha sido utilizado en AO y representa una fuente de información independiente de dominio (Denecke, 2008; Ogawa, Ma, y Yoshikawa, 2011).

3. Expansión de conceptos y pesado SWN

El método propuesto permite la expansión de conceptos mediante un algoritmo de tipo Page Rank (Page et al., 1998) sobre el grafo de WordNet, multiplicando los valores obtenidos de los synsets por los pesos de polaridad de SentiWordNet. Con más detalle, el procesamiento es el siguiente:

1. Los textos son procesados, haciendo uso de la biblioteca NLKT¹, para extraer los lemas y el *part of speech* (POS) de los mismos.
2. Se construyen los “contextos” para cada frase que son pasados a la herramienta UKB² para el cálculo, mediante un algoritmo de tipo Random Walk (Agirre y Soroa, 2009), que al mismo tiempo que desambigua los términos, genera los synsets con mayor peso resultado del proceso iterativo de propagación propio de estos algoritmos.
3. Expandimos el texto reemplazando los términos por sus vectores PPV (*Personalized Page Rank vectors*), que consisten en una secuencia de synsets de WordNet.
4. Los synsets, con los pesos asociados, son mapeados a SentiWordNet, obteniendo para ellos la carga subjetiva.
5. Se realiza el cálculo final de la polaridad mediante la media de la suma de los productos entre el peso del synset tras Page Rank y la polaridad asociada en SentiWordNet.

Debido a que buscamos una combinación de los valores de SentiWordNet con los pesos obtenidos por Page Rank, es importante asegurarse que la fórmula final produce valores comparables. Para ello, los valores de Page Rank para los synsets se ajustan a una norma L_1 , de forma que estos vectores de conceptos sumen la unidad como valor máximo

¹<http://www.nltk.org>

²<http://ixa2.si.ehu.es/ukb/>

posible. La polaridad final es obtenida mediante la fórmula siguiente:

$$p(t) = \frac{\mathbf{r} \cdot \mathbf{s}}{|\mathbf{t}|} \quad (2)$$

donde p es el valor de polaridad final, $\mathbf{r}$ es el vector de synsets con los pesos obtenidos por Page Rank sobre WordNet, $\mathbf{s}$ es el vector de polaridades correspondientes de SentiWordNet y t es el conjunto de conceptos expandidos que se derivan de la frase.

4. Experimentación

4.1. El corpus JRC

El corpus sobre citas en noticias periodísticas con el que vamos a trabajar ha sido preparado por el Joint Research Center (Balahur et al., 2010) y puede ser descargado desde su web de recursos lingüísticos³. Consiste en 1,590 citas en inglés que se han extraído automáticamente desde diversas fuentes periodísticas *online* y anotadas manualmente con expresividad del sentimiento hacia entidades como personas y organizaciones mencionadas en dichas citas. Estas anotaciones han sido llevadas a cabo por cuatro expertos, a partir de los cuales se construyen juicios de polaridad positiva o negativa. De las 1,590 entradas, sólo 427 tienen carga subjetiva y con valor consensuado entre los anotadores.

Un ejemplo de texto en dicho corpus sería el siguiente (etiquetado como *negativo* por dos anotadores):

Elisabeth's imprisonment was premeditated; Fritzl had been planning her imprisonment in the dungeon cellar for two to three months before he actually lured his daughter down there.

4.2. Evaluación

Hemos seleccionado de entre las entradas del corpus aquellas que son subjetivas y han sido consensuadas entre los anotadores, si bien también incluimos los extractos de noticias con una única anotación (ya sea positiva o negativa), lo que nos proporciona un total de 718 citas. Dado el carácter no supervisado de nuestro método, todas éstas son consideradas para la evaluación, calculando los valores de precisión, cobertura y medida F1 sobre este proceso de clasificación binaria.

El sistema, tal y como se ha planteado en secciones anteriores, puede configurarse en base a dos variables:

1. *Tamaño de la expansión.* Este valor hace referencia al número de synsets que vamos a asociar a cada término del texto. De esta forma, un valor 1 indicaría que nos quedamos con el synset asociado al término, por lo que nos limitaríamos a realizar un proceso de desambiguación mediante Page Rank sobre WordNet a partir de las semillas que se obtienen del resto de términos en el texto. Cualquier valor superior a 1 ya nos lleva a una expansión en el conjunto de synsets que vamos a asociar al texto. Es aquí donde reside la fortaleza del método, pues sentidos con carga emocional relevante pueden entrar en el cálculo final de la polaridad aunque no estén representando sentidos (conceptos) explícitamente reflejados en el texto. Trabajamos de esta forma con una solución que tienen en cuenta conceptos “implícitos” en el contexto de los términos de cada fragmento periodístico.
2. *Decisión sobre neutros.* En nuestra propuesta, actualmente los valores sin una orientación positiva o negativa clara son desechados. Pero también sería interesante analizar la tendencia del método a considerar neutros los textos que deberían ser positivos o negativos. Para ello se incorpora un segundo parámetro que fuerza un sesgo sobre el cálculo final, evitando una evaluación de la polaridad neutra (es decir, de valor cero).

4.3. Resultados

4.3.1. Tamaño de la expansión

Para estudiar el efecto que el tamaño de la expansión tiene sobre el método hemos explorado valores desde 1 (la mera desambiguación) a 20 (20 synsets asociados por término). Antes de entrar a analizar los valores obtenidos con el valor de expansión óptimo para esta colección de datos, podemos visualizar dicho efecto en la Figura 1.

De esta figura se extraen algunas observaciones interesantes. La primera es que el efecto que produce generalmente mejora la capacidad de clasificación en términos de F1 (siendo equivalente para precisión y cobertura). A medida que añadimos más términos el comportamiento es estable y podemos decir que en este caso añadir más de 9 conceptos por término no lleva a mejora alguna. En

³http://langtech.jrc.it/JRC_Resources.html

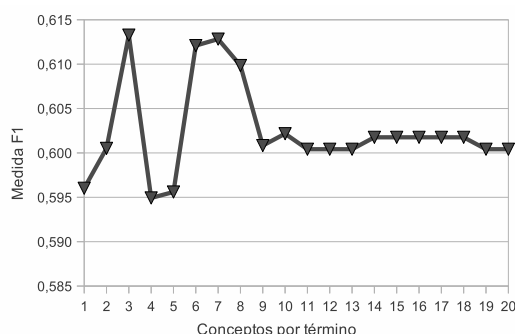


Figura 1: Efecto de la variación en el tamaño de la expansión de conceptos

cambio, sí es notable la variabilidad de la medida sobre los primeros valores. Creemos que dicho efecto tiene que ver con la inclusión de la antonimia y la polisemia como relaciones también consideradas en el grafo de WordNet. Dado que esta gráfica es sobre el total de las expansiones, es importante considerar un análisis futuro que estudie dicho efecto de forma más detallada, pues las expansiones pueden producir resultados muy dispares en función al término asociado.

Los resultados globales con el método propuesto sobre el corpus JRC Quotations quedan reflejados en el Cuadro 1, donde indicamos lo obtenido con expansiones de 1 a 10 conceptos por término (*vsize* en la tabla).

Precisión	Cobertura	F1	vsize
0,5991	0,5929	0,5960	1
0,6039	0,5971	0,6005	2
0,6199	0,6067	0,6132	3
0,6011	0,5888	0,5949	4
0,6022	0,5891	0,5955	5
0,6185	0,6057	0,6120	6
0,6192	0,6065	0,6128	7
0,6166	0,6031	0,6098	8
0,6075	0,5941	0,6008	9
0,6089	0,5955	0,6021	10

Cuadro 1: Resultados obtenidos

4.3.2. Tratamiento de valores neutros

En cuanto al tratamiento de los valores neutros, hemos considerado los tres escenarios posibles: considerar los neutros como positivos, considerar los neutros como negativos, o considerar los neutros como tales, no computando en las clases positiva o negativa. El Cuadro 2 resume los valores obtenidos con un valor de expansión de 3 (aquel que nos ha dado mejores resultados). Hemos observado que

es mejor considerar el neutro como tal, o lo que es lo mismo, el método propuesto no introduce ningún tipo de sesgo hacia una clase determinada, por lo que no es necesario realizar corrección alguna. Este comportamiento se mantiene independiente del tamaño de la expansión.

Precisión	Cobertura	F1	Neutro
0,6103	0,6008	0,6055	negativo
0,6198	0,6050	0,6123	positivo
0,6199	0,6067	0,6132	neutro

Cuadro 2: Consideración de valores neutros

4.3.3. Pesos de Page Rank

Llegados a este punto podemos plantearnos la conveniencia o no de la fórmula planteada, pues re-evalúa los valores de SentiWordNet con los valores para los synsets obtenidos mediante el algoritmo Page Rank (mediante una aproximación basada en Random Walk). El diagrama mostrado en la Figura 2 arroja luz sobre esta cuestión. Podemos reconocer claramente la contribución que los pesos de Page Rank aportan al valor final de polaridad, no solo en cuanto a la mejor sino en cuanto a la estabilización que introduce en base al tamaño de la expansión. También es importante señalar que estos efectos no son tan marcados cuando la expansión no va más allá de uno o dos términos. Lo cual puede estar relacionado con la discusión anterior acerca de los efectos de synsets provenientes de relaciones no deseables en el grafo de WordNet.

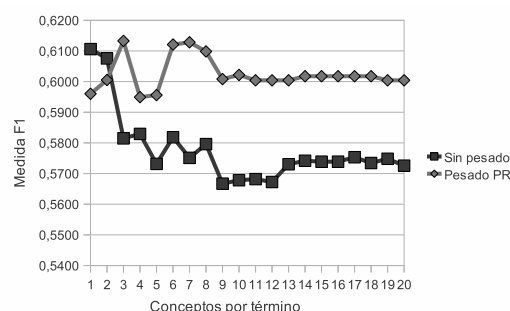


Figura 2: Efecto del pesado con los valores obtenidos por Page Rank

4.3.4. Comparación con otros trabajos

Existen dos trabajos destacados que hayan propuesto métodos de clasificación de

la polaridad evaluados con este corpus y que ya han sido mencionados anteriormente. En un primer trabajo (Balahur et al., 2010), precisamente donde se introduce este corpus a toda la comunidad científica, se estudian diversas variantes en base al uso de distintos recursos lingüísticos (SentiWordNet, Micro WordNet, WordNet Affect, etc.). Obtienen un valor de *accuracy* de 0,25 mediante el uso de SentiWordNet, mientras que nuestro valor máximo para esta misma medida es de 0,59 con una expansión de 6 conceptos con término. No hemos indicado los valores de *accuracy* pues consideramos más interesante para evaluar la clasificación de la polaridad los valores de precisión, cobertura y medida F. No obstante, hemos de resaltar que se logra un valor de 0,82 en esta misma medida usando determinadas bolsas de palabras específicas para este tipo de corpus.

En un trabajo posterior (Boldrini et al., 2012), encontramos otros resultados para la clasificación de este corpus, alcanzando un valor para F1 de 0,5340. Ya hemos indicado en el apartado anterior que el valor máximo obtenido en nuestro caso es de 0,6132. Con lo cual el sistema que presentamos obtiene una mejora de 7,92 %.

5. Conclusiones y trabajo futuro

El método propuesto plantea una solución muy prometedora para el AO. Los resultados obtenidos demuestran que es posible conseguir valores de precisión y cobertura similares o superiores a los de los métodos supervisados más extendidos, escapando a los problemas que plantea la transferencia de sentimiento entre distintos dominios o, dicho de otro modo, la validez de los modelos entrenados sobre un dominio para clasificar en otro. Además, aunque los experimentos han sido probados para el inglés, el método es completamente adaptable a cualquier idioma diferente.

Dada las bondades del método, nuestros objetivos se centran, por un lado, en seguir comprobando su validez sobre otros corpora y, por otro lado, en aplicar dicho método sobre idiomas distintos al inglés. Esto último es un reto importante, pero afortunadamente creemos que a día de hoy disponemos de recursos como Freeling (Padró et al., 2010) o el Multilingual Central Repository (Atserias et al., 2003) que pueden hacer viable la multilingüidad de nuestra propuesta sin recurrir a

traducciones (Denecke, 2008).

Entre otros retos, está el estudio y tratamiento de la negación, ya que las relaciones en el grafo de palabras de WordNet puede verse fuertemente afectado por un seguimiento incorrecto de las relaciones de antonimia, actualmente consideradas para calcular la expansión conceptual.

Bibliografía

- Agirre, Eneko y Aitor Soroa. 2009. Personalizing pagerank for word sense disambiguation. En *EACL '09: Proceedings of the 12th Conference of the European Chapter of the Association for Computational Linguistics*, páginas 33–41, Morristown, NJ, USA. Association for Computational Linguistics.
- Atserias, J., L. Villarejo, G. Rigau, E. Agirre, J. Carroll, B. Magnini, y P. Vossen. 2003. The MEANING Multilingual Central Repository. En Sojka, P and Pala, K and Smrz, P and Fellbaum, C and Vossen, P, editor, *GWC 2004: SECOND INTERNATIONAL WORDNET CONFERENCE, PROCEEDINGS*, páginas 23–30. 2nd International Global WordNet Conference (GWC 2004), Brno, CZECH REPUBLIC, JAN 20-23, 2004.
- Baccianella, Stefano, Andrea Esuli, y Fabrizio Sebastiani. 2008. Sentiwordnet 3.0 : An enhanced lexical resource for sentiment analysis and opinion mining. *Proceedings of the Seventh conference on International Language Resources and Evaluation LREC10*, 0:2200–2204.
- Balahur, Alexandra, Ralf Steinberger, Mijail Kabadjov, Vanni Zavarella, Erik Van Der Goot, Matina Halkia, Bruno Pouliquen, y Jenya Belyaeva. 2010. Sentiment analysis in the news. *English*, 10:2216–2220.
- Boldrini, Ester, Alexandra Balahur, Patricio Mart  nez-Barco, y Andr  s Montoyo. 2012. Using emotiblog to annotate and analyse subjectivity in the new textual genres. *Data Mining and Knowledge Discovery*, páginas 1–32. 10.1007/s10618-012-0259-9.
- Denecke, K. 2008. Using sentiwordnet for multilingual sentiment analysis. En *Data Engineering Workshop, 2008. ICDEW 2008. IEEE 24th International Conference on*, páginas 507–512, april.

- Fellbaum, Christiane, editor. 1998. *WordNet: An Electronic Lexical Database*. MIT Press, Cambridge, MA.
- Hatzivassiloglou, Vasileios y Kathleen R. McKeown. 1997. Predicting the semantic orientation of adjectives. En *Proceedings of the 35th Annual Meeting of the Association for Computational Linguistics and Eighth Conference of the European Chapter of the Association for Computational Linguistics*, ACL '98, páginas 174–181, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Hu, Minqing y Bing Liu. 2004. Mining and summarizing customer reviews. En *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*, KDD '04, páginas 168–177, New York, NY, USA. ACM.
- Ji, Yang-Sheng, Jia-Jun Chen, Gang Niu, Lin Shang, y Xin-Yu Dai. 2011. Transfer Learning via Multi-View Principal Component Analysis. *JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY*, 26(1):81–98, JAN.
- Kanayama, Hiroshi y Tetsuya Nasukawa. 2006. Fully automatic lexicon expansion for domain-oriented sentiment analysis. En *Proceedings of the 2006 Conference on Empirical Methods in Natural Language Processing*, EMNLP '06, páginas 355–363, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Kim, Soo-Min y Eduard Hovy. 2004. Determining the sentiment of opinions. En *Proceedings of the 20th international conference on Computational Linguistics*, COLING '04, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Liu, Bing. 2010. Sentiment analysis and subjectivity. En Nitin Indurkha y Fred J. Damerau, editores, *Handbook of Natural Language Processing, Second Edition*. CRC Press, Taylor and Francis Group, Boca Raton, FL. ISBN 978-1420085921.
- Martínez-Cámara, Eugenio, M. Teresa Martín-Valdivia, José M. Perea-Ortega, y L. Alfonso Ure na López. 2011. Técnicas de clasificación de opiniones aplicadas a un corpus en español. *Procesamiento de Lenguaje Natural*, 47.
- Mullen, Tony y Nigel Collier. 2004. Sentiment analysis using support vector machines with diverse information sources. En *EMNLP*, páginas 412–418. ACL. conf/emnlp/2004.
- Ogawa, Tatsuya, Qiang Ma, y Masatoshi Yoshikawa. 2011. News Bias Analysis Based on Stakeholder Mining. *IEICE TRANSACTIONS ON INFORMATION AND SYSTEMS*, E94D(3):578–586, MAR.
- Padró, Lluís, Miquel Collado, Samuel Reese, Marina Lloberes, y Irene Castellón. 2010. Freeling 2.1: Five years of open-source language processing tools. En Nicoletta Calzolari (Conference Chair) Khalid Choukri Bente Maegaard Joseph Mariani Jan Odijk Stelios Piperidis Mike Rosner, y Daniel Tapias, editores, *Proceedings of the Seventh International Conference on Language Resources and Evaluation (LREC'10)*, Valletta, Malta, may. European Language Resources Association (ELRA).
- Page, Lawrence, Sergey Brin, Rajeev Motwani, y Terry Winograd. 1998. The PageRank Citation Ranking: Bringing Order to the Web. Informe técnico, Computer Science Department, Stanford University.
- Pang, Bo y Lillian Lee. 2008. Opinion mining and sentiment analysis. *Found. Trends Inf. Retr.*, 2(1-2):1–135, jan.
- Pang, Bo, Lillian Lee, y Shivakumar Vaithyanathan. 2002. Thumbs up?: sentiment classification using machine learning techniques. En *Proceedings of the ACL-02 conference on Empirical methods in natural language processing - Volume 10*, EMNLP '02, páginas 79–86, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Tan, Songbo, Xueqi Cheng, Yuefen Wang, y Hongbo Xu. 2009. Adapting Naive Bayes to Domain Adaptation for Sentiment Analysis. En Boughanem, M and Berrut, C and Mothe, J and SouleDupuy, C, editor, *ADVANCES IN INFORMATION RETRIEVAL, PROCEEDINGS*, volumen 5478 de *Lecture Notes in Computer Science*, páginas 337–349. Google; Mateixware Informat Serv; Microsoft Res;

- Yahoo; Exalead; GDR 13; Univ Paul Sabatier; ARIA; Inforsid & Reg; Midi Pyrennees. 31st European Conference on Information Research, Toulouse, FRANCE, APR 06-09, 2009.
- Tsytsarau, Mikalai y Themis Palpanas. 2012. Survey on mining subjective data on the web. *Data Mining and Knowledge Discovery*, 24:478–514.
- Turney, Peter D. 2002. Thumbs up or thumbs down?: semantic orientation applied to unsupervised classification of reviews. En *Proceedings of the 40th Annual Meeting on Association for Computational Linguistics*, ACL '02, páginas 417–424, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Whitelaw, Casey, Navendu Garg, y Shlomo Argamon. 2005. Using appraisal groups for sentiment analysis. En *Proceedings of the 14th ACM international conference on Information and knowledge management*, CIKM '05, páginas 625–631, New York, NY, USA. ACM.
- Wiegand, Michael, Alexandra Balahur, Benjamin Roth, Dietrich Klakow, y Andrés Montoyo. 2010. A survey on the role of negation in sentiment analysis. En *Proceedings of the Workshop on Negation and Speculation in Natural Language Processing*, NeSp-NLP '10, páginas 60–68, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Wu, Qiong, Songbo Tan, y Xueqi Cheng. 2009. Graph ranking for sentiment transfer. En *Proceedings of the ACL-IJCNLP 2009 Conference Short Papers*, ACLShort '09, páginas 317–320, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Wu, Qiong, Songbo Tan, Miya Duan, y Xueqi Cheng. 2010. A Two-Stage Algorithm for Domain Adaptation with Application to Sentiment Transfer Problems. En Cheng, PJ and Kan, MY and Lam, W and Nakov, P, editor, *INFORMATION RETRIEVAL TECHNOLOGY*, volumen 6458 de *Lecture Notes in Computer Science*, páginas 443–453. Natl Taiwan Univ; Natl Sci Council, Republ China; Minist Educ, Republ China. 6th Asia Information Retrieval Societies Conference, Taipei, TAIWAN, DEC 01-03, 2010.