

独立成分分析を用いた文書分類
—SVMのための素性空間再構成—

高村 大也

松本 裕治

奈良先端科学技術大学院大学
情報科学研究科自然言語処理学講座
〒630-0101 奈良県生駒市高山町 8916-5
0743-72-5246, 5248

{hiroya-t,matsu}@is.aist-nara.ac.jp

独立成分分析を用いた素性空間再構成による新しい文書分類の方法を提案する。提案手法では、独立成分分析により変換された文書ベクトルを元のベクトルと結合することにより素性空間を再構成し、得られた新しいベクトルを入力として SVM にを用いて分類を行う。これにより、元の空間の情報を失うことなく Latent Semantic Space に重みを与えることができ、SVM の分類能力を向上させることが可能になる。実験では、ラベル付データの量に関わらず高精度での文書分類に成功した。

キーワード：文書分類, 独立成分分析, サポートベクターマシン, ラベル無しデータ, 素性空間

Feature Space Restructuring for SVMs
with Application to Text Categorization

Hiroya Takamura

Yuji Matsumoto

Nara Institute of Science and Technology
Graduate School of Information Science
8916-5 Takamaya, Ikoma, Nara 630-0101, JAPAN
+81-743-72-5246, 5248

{hiroya-t,matsu}@is.aist-nara.ac.jp

In this paper, we propose a new method of text categorization based on feature space restructuring for SVMs. In our method, independent components of document vectors are extracted using ICA and concatenated with the original vectors. This restructuring makes it possible for SVMs to focus on the latent semantic space without losing information given by the original feature space. Using this method, we achieved high performance in text categorization both with small number and large numbers of labeled data.

Keywords : Text Categorization, Independent Component Analysis, Support Vector Machine, unlabeled data, feature space

1 はじめに

従来の文書分類の方法は、その多くが十分な量のラベル付データに基づいたものである。しかし、ラベル付データの収集には多大な労力がかかり、実際の応用を考慮すると、文書分類はラベル付データが少ない場合においても良い精度を実現しなくてはならない。もちろん、少量のラベル付データでの分類方法も今までいくつか提案されているが (Nigam et al, 2000), さらなる改良が必要である。そのような方法の実現のためには、ラベル無しデータが与える貴重な情報を十分に活用することが重要である。本稿では、サポートベクターマシン (SVMs) (Vapnik, 1995) と独立成分分析 (ICA) (Herauld and Jutten, 1986; Bell and Sejnowski, 1995) に基づいた新たな分類手法を提案する。

SVM は、画像処理や自然言語処理など多くの分野で応用されてきた。SVM を文書分類に適用するという考えは (Joachims, 1998) で初めて導入され、良い結果を残している。しかし、ラベル付データが少ない場合は、しばしば高精度の分類に失敗することがある。このような問題に対しては、大きく分けて二つのアプローチがある。一つは学習アルゴリズムそのものを改良するもので (Joachims, 1999a; Glenn and Mangasarian, 2001), もう一つは素性選択など、データに働きかけるものである (Weston et al, 2000)。本稿で提案する手法は後者に属し、素性空間の再構成 (Feature Space Restructuring) を行う。k-NN 法 (Mitchell, 1997) などに対しては主成分分析 (PCA) により素性空間の次元を圧縮する方法がよく用いられる。しかし後に実験で示すように、従来の次元圧縮に基づく方法は、SVM に対して必ずしも良い結果を生まない。従来の方法と違い、我々の提案手法は ICA によって得られた独立成分を素性空間の次元を拡張するために用いている。

ICA は信号処理の分野で発展してきている手法で、混合信号から独立成分を抽出するために使用される。ICA は二つの仮定の上に成り立っている: (1) 信号源は統計的に互いに独立である (2) 観測信号は信号源の線形混合として複数の点において計測される。ICA は理論面、応用面ともに発展してきている (Bell and Sejnowski, 1997)。ICA の文書分類への応用は (Isbell and Viola, 1998) などでも試みられた。それらの研究において、文書は単語の

出現頻度を要素とするベクトルで表されている。しかし、抽出された独立成分は必ずしも分類すべきクラスに相当しないことが報告されている (Kolenda et al, 2000)。(Kaban and Girolami, 2000) では、文書ベクトルが潜在的に持つ独立成分数は人間によって付けられたクラス数より多いことが指摘されている。これらの事実は、ICA を文書分類に直接利用することは困難であるということを示す。

こういった観察を考慮に入れて、我々には以下のようなアプローチを採る。まず、ICA により文書ベクトルを圧縮し、次にそれを元のベクトルと結合 (concatenate) することによりベクトルを再構成する。そのベクトルを新たな素性ベクトルとして文書分類を行う。

ICA の類似手法としては PCA がある。よって我々は、通常の SVM, PCA によって再構成された素性ベクトルでの SVM, ICA によって再構成された素性ベクトルでの SVM について実験を行った。その結果、ラベル付データの量の大小に関わらず、提案手法 (ICA によって再構成された素性ベクトルでの SVM) が他の方法よりも良い結果を出すことが示された。

2 サポートベクターマシン

2.1 サポートベクターマシンの概観

サポートベクターマシン (SVM) は、Large-Margin Classifier (Smola et al, 2000) の一種である。与えられた素性ベクトルとラベルのペアの集合、

$$\begin{aligned} & (x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) \\ & \forall i, x_i \in \mathbf{R}^d, y_i \in \{-1, 1\} \end{aligned} \quad (1)$$

に対して、SVM はマージン (分離平面とベクトルとの距離) が最大になるような分離平面を構成する (図 1):

$$f(x) = w \cdot x + b. \quad (2)$$

最大マージンを達成することはノルム $\|w\|$ を最小にすると同値になる。この問題は次のように書ける。

$$\begin{aligned} \min. \quad & \frac{1}{2} \|w\|^2, \\ \text{s.t. } \forall i, y_i (x_i \cdot w + b) - 1 \geq 0. \end{aligned} \quad (3)$$

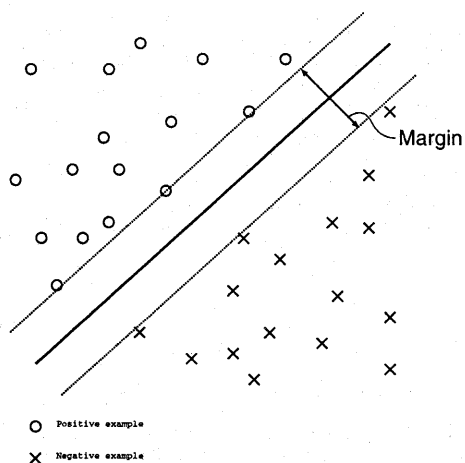


図 1: Support Vector Machine (実線が最適分離平面を表す).

この問題の解は次の双対問題を解くことによって得られる:

$$\begin{aligned} \max. \quad & \sum_i \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j y_i y_j \mathbf{x}_i \cdot \mathbf{x}_j \quad (4) \\ \text{s.t.} \quad & \sum_i \alpha_i y_i = 0, \\ & \forall i, \alpha_i \geq 0. \end{aligned}$$

ここで α_i 's はラグランジュ乗数である。(4)を最大にする α_i 's を用いて、最適な $\mathbf{w}$ は、

$$\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i. \quad (5)$$

と表される。(5)を(2)に代入することにより、

$$f(\mathbf{x}) = \sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b. \quad (6)$$

を得る。

テストデータは(6)の符合に従って分類される。

2.2 カーネル法

SVMは線形分類器であるので、その分離能力には限界がある。この限界を越えるために、通常はカーネル法がSVMと組み合わせて用いられる(Vapnik, 1995)。

カーネル法では、(4)や(6)における内積が、カーネル関数と呼ばれるより一般的な内積 $K(\mathbf{x}_i, \mathbf{x}_j)$ に置き換えられる。もちろんカー

ネル関数は一定の条件(Mercerの条件)を満たしていなくてはならない。よく使用されるカーネル関数としては、多項式カーネル $(\mathbf{x}_i \cdot \mathbf{x}_j + 1)^d$ ($d \in \mathbb{N}_+$)、RBFカーネル $\exp\{-\|\mathbf{x}_i - \mathbf{x}_j\|^2 / 2\sigma^2\}$ などがある。

カーネル関数により、素性ベクトルは(高次元の)ヒルベルト空間に写像され、その空間において線形分離される。この写像によって、SVMは線形分類器であるにも関わらず、非線形分離が可能になる。

カーネル法を用いることのもう一つの利点は、高次元の空間を扱うにも関わらず、高次元ベクトルを明示的に計算する必要がないことである。カーネル関数の値だけを計算すればよい。これは計算量の大幅な節約につながる。

2.3 トランスダクティブ SVM

トランスダクティブ サポートベクターマシン(TSVM)は、(Joachims, 1999a)において提案された手法で、(Vapnik, 1995)における *transductive learning* (帰納過程と演繹過程を同時に行う学習)の一つの具現化であるといえる。TSVMは少量のラベル付データのみに基づいた分類のために設計されている。TSVMのアルゴリズムの概要は次の通りである:

1. ラベル付データを用いて分離平面を構築する。
2. 現在の分離平面によりラベル無しデータを分類する。
3. 2で正例と判定された事例と負例のペアで、分離平面に十分近いものを選ぶ。
4. 3で選ばれたペアのラベルを交換する。ただし交換によりマージンが大きくなる場合のみこの操作を行う。
5. 終了条件が成立したら終了。そうでなければ2へ戻る。

これは、直前の繰り返しで構築された分類器によって付与されたラベルの付け替えを許して、マージン最大化を実現する方法である。

3 独立成分分析

独立成分分析(ICA)は、混合信号から独立信号源を抽出する方法である。ICAにおいては二つのことが仮定されている:(1)信号源 $\mathbf{s} \in \mathbb{R}^m$ は統計的に互いに独立である(2)観測信号 $\mathbf{x} \in \mathbb{R}^n$ は信号源の線形混合である。

つまり,

$$\mathbf{x} = A\mathbf{s}. \quad (7)$$

ここで A は混合行列と呼ばれる. $\mathbf{x}$ は時系列として観測され, それをもとに A と $\mathbf{s} = (s_1, \dots, s_m)$ が推定される. 言い替えると, ここでの我々の目的は, 分離行列 W を, $s_1, \dots, s_m$ が互いにできるだけ独立となるように選ぶことである:

$$\mathbf{s} = W\mathbf{x}. \quad (8)$$

計算は独立性を測る目的関数を用い, 最急勾配法によって行われる. 独立性の尺度としていくつかの関数が提案されているが, ここでは情報量最大化基準を用いる (Bell and Sejnowski, 1995). 学習規則は自然勾配 (Amari, 1998) を考慮して算出されたものを使用する:

$$\delta W = (I + (I - 2g(W\mathbf{x}))(W\mathbf{x})^t)W, \quad (9)$$

where, $g(u) = 1/(1 + \exp(-u))$.

これは, パラメータ空間がユークリッド空間でなく実際はリーマン空間であることを考慮に入れたモデルである. 自然勾配を考慮することにより, 収束速度の向上, 収束の安定性などの利点がある上に, 更新式は自然勾配を使わないものよりはるかに簡潔になる.

4 素性空間再構成を基にした文書分類

従来研究と同様, 我々も文書を表現するためにベクトル空間モデルを採用する (Salton and McGill, 1983). ベクトル空間モデルでは, 文書 $\mathbf{d}$ はその文書内の各単語の頻度を要素とするベクトル $(f_1, \dots, f_d)$ で表される.

まず, PCA もしくは ICA を用いて文書ベクトルの次元を圧縮する. PCA の場合は多くの従来研究 (Deerwester et al, 1990) と全く同様である. (Isbell and Viola, 1998) においては, ICA が次元圧縮に使用されて情報検索でよい結果を残している. 圧縮段階では, 彼らの方法を用いる. 彼らの枠組では, 各ドキュメント $\mathbf{d}$ が, トピックに対応すると考えられる信号源 $\mathbf{s}$ の線形混合で表されると仮定されている. つまり, 各単語が”マイク”の役割を担い, 各時刻 (つまり各文書) においてその文書内の各単語頻度を混合信号として受信する. この定式化は以下のように表される.

$$\mathbf{d} = A\mathbf{s}. \quad (10)$$

ここで A は混合行列である. A と $\mathbf{s}$ は共に未知であるが, 独立性の仮定に従って求めることができる. 信号源 $\mathbf{s}$ は該当する文書の圧縮表現ということになる. PCA の場合もほとんど同様で, 独立成分が主成分となるだけである.

信号源ベクトル $\mathbf{s}$ を計算した後, 元のベクトル $\mathbf{d}$ と圧縮されたベクトル $\mathbf{s}$ を連結 (concatenate) する:

$$\hat{\mathbf{d}} = \begin{bmatrix} \mathbf{d} \\ \mathbf{s} \end{bmatrix}. \quad (11)$$

$\hat{\mathbf{d}}$ を入力として SVM により分類を行う. 本来 SVM は二値分類器であるので, ここでは one-versus-rest 法により多値分類器に拡張する.

5 理論的考察

5.1 カーネル関数としての妥当性

提案された素性空間再構成法は, 元の空間である特殊なカーネル関数を使うことと同等である. このことを線形の場合について説明する. 二つのベクトル $\mathbf{d}_1, \mathbf{d}_2$ が与えられた時, 再構成された空間のカーネル関数 K は次のように表される:

$$\begin{aligned} K(\hat{\mathbf{d}}_1, \hat{\mathbf{d}}_2) &= \hat{\mathbf{d}}_1^t \hat{\mathbf{d}}_2 \\ &= \mathbf{d}_1^t \mathbf{d}_2 + \mathbf{s}_1^t \mathbf{s}_2 \\ &= \mathbf{d}_1^t \mathbf{d}_2 + \mathbf{d}_1^t A^t A \mathbf{d}_2. \end{aligned} \quad (12)$$

上式の二つの項がどちらもカーネルであること, また任意の二つのカーネルの和がカーネルになること (Vapnik, 1995) を考慮すると, 再構成は元の空間における特殊なカーネルを使用することに相当することが示される.

線形カーネルの場合について説明したが, 一般的によく使用される多項式カーネル $(\mathbf{x}_i \cdot \mathbf{x}_j + 1)^d (d \in \mathbb{N}_+)$ については全く同様の議論が成立する.

5.2 素性空間再構成の解釈

(12) は, それぞれ ICA あるいは PCA によって決定される Latent Semantic Index に重みを与えていることになる. ここで重みを与えることと圧縮することは異なることに注意されたい. 一般の次元圧縮法では, Latent Semantic Space だけが考慮されるが, 我々の方法では元の空間が依然として計算結果に影響を与えている.

表 1: 実験で使用された文書

カテゴリー	文書数
earn	2673
acq	1435
trade	225
crude	223
money-fx	176
interest	140

提案手法におけるこの性質は、元の空間によって与えられる情報を失わずに Latent Semantic Space に注目することを可能にする。

6 実験

提案手法の有効性を示すために、いくつかの実験を行なった。

使用したデータは Reuters-21578 データセットである。最も頻度の高い 6 つのカテゴリーを訓練セットの中から選んだ (表 1)。これを部分的にラベル付データとして扱うことにより訓練データとテストデータに分割した。データ中に 2 回以上出現した単語のみを使い、また stemming とストップワード削除も行った。計算には SVM-light (Joachims, 1999b) を使用した。

実験は大きく二種類に分かれる。一つは固定したラベル付データ数に対し、各カテゴリーにおける分類性能の振舞いを調べる目的で行なわれた (6.1 節)。もう一つはラベル付データの数が増加していった時の分類性能の振舞いを調べるためである (6.2 節)。

結果は F 値を用いて評価する。カテゴリーによる分類性能の違いを吸収するために、F 値のマイクロ平均とマクロ平均を算出した (Yang, 1999)。

マイクロ平均は、まずカテゴリーを無視して全体の適合率と再現率を計算し、それらを用いて F 値を計算することにより算出される。マクロ平均は、まず各カテゴリーの F 値を計算し、それら F 値の単純平均を計算することにより算出される。マイクロ平均は、事例数の大きいカテゴリーに大きな影響を受け、マクロ平均は事例数の小さいカテゴリーに大きな影響を受ける。

用いたカーネル関数は線形カーネルである。抽出された独立成分あるいは主成分の数は 50 に固定した。

6.1 ラベル付データ数固定での各カテゴリーにおける分類性能

ここでは、それぞれ 100, 500, 1000, 2000 個の事例をラベル付きとして扱い、残りをラベル無しとして扱った。ランダムに選ばれた異なるラベル付データに対しそれぞれ 10 回実験を行い、平均値を算出した。結果を表 2, 3, 4, 5 に示す。列 "Method" には、使用された再構成方法が記されている。"Original" は元の文書ベクトル、"PCA" と "ICA" は圧縮されたベクトル、"Original+PCA" と "Original+ICA" は 4 節で説明した再構成されたベクトルを示す。

提案手法は、ラベル付データ数が 100 及び 500 のときはほとんどのカテゴリーで、ラベル付データ数が 1000 及び 2000 のときはすべてのカテゴリーで最も高い F 値を示している。また、マイクロ平均とマクロ平均に関してはすべてのラベル付データ数の場合について、提案手法が最高値を記録している。このことは、提案手法がカテゴリーやラベル付データのサイズに関わらず高い性能を持つことを示している。

6.2 ラベル付データ数を増加させた時の振舞い

ラベル付データ数を増加させていったときに各方法がどのように振舞うかを調べるために、この実験を行なった。ラベル付データ数は 100 から 2000 まで、100 ずつ増加させた。結果を図 2, 3 に示す。"PCA" などはラベル付データ数が小さいときのみ、また "Original" はラベル付データ数が大きいときのみ良い値を示すのに対し、提案手法はラベル付データ数が小さいときも大きいときも良い値を示していることがわかる。

7 おわりに

SVM のための素性空間再構成方法を提案した。提案手法では ICA により独立成分を抽出し、元のベクトルと結合させる。この手法が、文書分類においてラベル付データ数に関わらず高い精度を与えることを実験的に示した。

本稿で提案した素性空間再構成方法は他の機械学習アルゴリズムにも適用可能である。ただし、素性空間の次元が大きくなるので、計算量の面で高次元を扱うことに支障のないものでなくてはならない。このような観点で見ると、提案手法はカーネル関数を用いたアルゴリズムと相性がよいものと期待される。

表 2: F-値 (100 ラベル付データ)

Method	Original	Original(TSVM)	PCA	ICA	Original+PCA	Original+ICA
earn	92.96	84.00	91.13	86.60	92.97	92.88
acq	85.88	81.42	85.67	80.86	85.91	87.48
trade	36.52	65.59	72.41	72.28	36.68	70.73
crude	65.69	70.90	79.75	80.67	65.93	82.87
money-fx	32.46	45.01	52.69	54.37	32.47	48.62
interest	51.30	52.69	64.44	63.48	51.30	64.84
ミクロ平均	83.63	79.48	85.98	82.14	83.66	87.40
マクロ平均	60.80	66.60	74.34	73.04	60.87	74.56

表 3: F-値 (500 ラベル付データ)

Method	Original	Original(TSVM)	PCA	ICA	Original+PCA	Original+ICA
earn	96.49	93.97	94.38	93.45	96.49	96.70
acq	93.23	91.57	89.18	87.45	93.22	93.41
trade	86.31	80.81	87.42	86.58	86.37	91.70
crude	83.33	79.78	81.36	78.28	83.43	87.12
money-fx	62.94	64.88	72.83	73.45	63.17	73.99
interest	59.31	52.02	73.37	72.18	59.31	70.41
マイクロ平均	92.17	89.75	90.54	89.33	92.19	93.48
マクロ平均	80.26	77.17	83.09	81.89	80.34	85.55

表 4: F-値 (1000 ラベル付データ)

Method	Original	Original(TSVM)	PCA	ICA	Original+PCA	Original+ICA
earn	97.15	95.52	96.07	95.53	97.15	97.26
acq	94.60	93.77	92.18	91.44	94.60	94.84
trade	91.19	86.11	87.13	86.87	91.23	93.25
crude	87.99	80.03	80.93	78.75	87.99	89.41
money-fx	73.68	68.85	72.96	72.68	69.96	80.99
interest	75.34	57.26	72.83	68.25	75.34	79.27
マイクロ平均	94.23	91.79	92.31	91.54	94.09	94.90
マクロ平均	86.65	80.25	83.68	82.25	86.04	89.17

表 5: F-値 (2000 ラベル付データ)

Method	Original	Original(TSVM)	PCA	ICA	Original+PCA	Original+ICA
earn	97.48	95.92	97.18	97.12	97.48	97.55
acq	95.39	94.39	94.78	94.80	95.39	95.65
trade	93.81	86.33	88.61	85.28	93.81	95.90
crude	89.88	80.35	82.63	78.56	89.88	90.25
money-fx	77.44	70.60	74.84	70.69	77.49	81.56
interest	82.71	62.15	73.99	68.46	82.76	83.02
マイクロ平均	95.19	92.43	93.93	93.26	95.20	95.58
マクロ平均	89.45	81.62	85.33	82.48	89.47	90.65

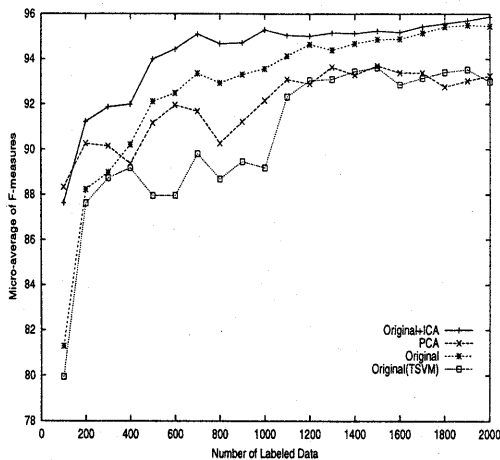


図 2: Micro-average

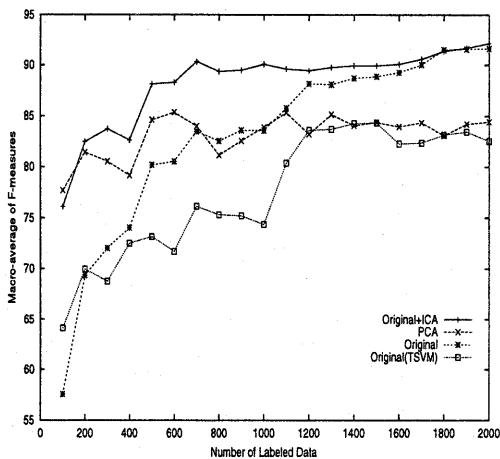


図 3: Macro-average

また、文書分類だけでなく、多義語の曖昧性解消など、他のタスクへの応用も期待される。

今後の発展としては、まず抽出すべき独立成分数の自動決定が重要である。本稿では独立成分数を 50 とおいたが、これには理論的背景がない。この独立成分数に依存して分類性能に影響が生じる可能性もあるので、Minimum Description Length (Rissanen, 1987) や Akaike Information Criterion

(Akaike, 1974) などといったモデル選択基準などを用いて最適な独立成分数を決定するような枠組が必要になるだろう。

また、Latent Semantic Space を重みづけする際に、本稿では圧縮されたベクトルをそのまま使用しているが、圧縮されたベクトルの定数倍を使用する方法も考えられる。ここにも最適化すべき問題が潜んでいる。

References

- Akaike, H. 1974. A New Look at the Statistical Model Identification. *IEEE Trans. Autom. Control*, vol. AC-19, pp. 716-723.
- Amari, S. 1998. Natural Gradient Works Efficiently in Learning. *Neural Computation*, vol. 10-2, pp. 251-276.
- Bell, A. J. and Sejnowski, T. J. 1995. An Information Maximization Approach to Blind Separation and Blind Deconvolution. *Neural Computation*, 7, 1129-1159.
- Bell, A. J. and Sejnowski, T. J. 1997. The 'Independent Components' of Natural Scenes are Edge Filters. *Vision Research*, 37(23), pp. 3327-3338.
- Deerwester, S., Dumais, T., Landauer, T., Furnas, W. and Harshman, A. 1990. Indexing by Latent Semantic Analysis. *Journal of the Society for Information Science*, 41(6), pp. 391-497.
- Glenn, F. and Mangasarian, O. 2001. Semi-Supervised Support Vector Machines for Unlabeled Data Classification. *Optimization Methods and Software*, pp. 1-14.
- Herault, J. and Jutten, J. 1986. Space or Time Adaptive Signal Processing by Neural Network Models. *Neural networks for computing: AIP conference proceedings* 151.
- Isbell, C. and Viola, P. 1998. Restructuring Sparse High Dimensional Data for Effective Retrieval. *Advances in Neural Information Processing Systems*, volume 11.
- Joachims, T. 1998. Text Categorization with Support Vector Machines: Learning with Many Relevant Features. *Proceedings of the European Conference on Machine Learning*.
- Joachims, T. 1999. Transductive Inference for Text Classification using Support Vec-

- tor Machines. *Machine Learning - Proc. 16th Int'l Conf. (ICML '99)*, pp. 200-209.
- Joachims, T. 1999. Making large-Scale SVM Learning Practical. *Advances in Kernel Methods - Support Vector Learning*.
- Kaban, A. and Girolami, M. 2000. Unsupervised Topic Separation and Keyword Identification in Document Collections: A Projection Approach *Technical Report* available in <http://cis.paisley.ac.uk/research/reports/index.html>
- Kolenda, T, Hansen, L., K. and Sigurdsson, S. 2000. Independent Components in Text. *Advances in Independent Component Analysis*, Springer-Verlag.
- Mitchell, T. 1997. *Machine Learning*, McGraw Hill.
- Nigam, K., McCallum, A., Thrun, S. and Mitchell, T. 2000. Text Classification from Labeled and Unlabeled Documents using EM. *Machine Learning*, 39(2/3). pp. 103-134.
- Rissanen, J. 1987. Stochastic Complexity. *Journal of Royal Statistical Society, Series B*, 49(3), pp. 223-239.
- Salton, G. and McGill, M. J. 1983. *Introduction to Modern Information Retrieval*. McGraw-Hill Book Company, New York.
- Smola, A., Bartlett, P., Schölkopf, B. and Schuurmans, D. 2000. *Advances in Large Margin Classifiers*. MIT Press
- Vapnik, V. 1995. *The Nature of Statistical Learning Theory*. Springer.
- Weston, J., Mukherjee, S., Chapelle, O., Pontil, M., Poggio, T. and Vapnik, V. 2000. Feature Selection for SVMs. *In Advances in Neural Information Processing Systems*, volume 13.
- Yang, Y. An Evaluation of Statistical Approaches to Text Categorization. *Information Retrieval*, volume 1, 1-2, pp. 69-90.