

梵文『法華經』の系統分類

逢坂雄美, 山崎守一

仙台電波工業高等専門学校
〒989-3124 仙台市青葉区上愛子北原

大乘仏教研究において最も重要な文献の一つである法華經の梵文写本は、発見された地域により、ネパール・チベット、カシミール、西域の3種に分類できる。これら3種の写本間には類似性と共に著しい相違も見られると同時に、同一地域で発見された写本間ですら相違が見られるのが実情である。特にネパール系写本間の相違が著しい。このように錯綜した写本の系統解明は系統ごとの校訂本作成にとって必須である。系統解明のために、31の写本の各詩偈の類似性、つまり語順対応・文字対応情報を抽出して、主成分分析及びクラスター分析を適用した。これらの系統は、①ネパール系の10の紙写本群、②9本のネパール系の貝葉写本と2つの紙写本を含む11写本群、③数本の写本からなるいくつかの小群に整理できた。この小群にはカシミール本と西域本が属している。この結果は、言語学的研究から特徴を抽出したデータに対して、同様な手法を適用して得られた結論と殆ど同じであった。

Classification of Saddharmapūṇḍarīka Based on Statistical Analysis

Yumi OUSAKA and Moriichi YAMAZAKI

Sendai National College of Technology
Kami-Ayashi, Aobaku, Sendai 989-3124

The texts of the Saddharmapūṇḍarīka, which is one of the most important manuscripts in the study of the Buddhism, are classified into three kinds of texts, Nepal, Kashmir and Central Asia. In order to emendate the critical texts to each kind, it is indispensable to classify these many manuscripts, which have many differences as well as many analogous features, into proper groups. For clarification of these texts, we apply the principal component analysis and also the cluster analysis to the data, which describe the similarity between the verses of different manuscripts, that is, the informations on the correspondences of words and also character sequences. As a result, we can successfully classify 32 manuscripts into two large groups and several small groups: a group consisting of ten paper manuscripts for Nepal, one group comprising of nine palm and two paper manuscripts. The Kashmir and Central Asia manuscripts belong to these small groups. These results grossly correspond to the conclusion obtained by applying the same method to the data, which are selected on the basis of the linguistic discussions on the manuscripts.

1. 序論

本論文では、大乘仏教の研究において最も重要な文献の一つである法華經の梵文写本の系統分類について、統計解析の手法に基づき議論する。

梵文法華經の写本は、発見された地域により、1. ネパール・チベット、2. カシミール（ギルギット）、3. 西域（中央アジア）の3種に分類できる（表1参照）。ネパール・チベット系写本は、2、3と比較して、完本の数が多いのが特徴である。材質は貝葉（パーム）と紙とに分かれ、貝葉写本は11世紀、紙写本は18～19世紀に集中している。これら3種の写本間には相違が見られ、特に1と

3との間には著しい相違がある。これらの系統分類をすることは、系統ごとの校訂本作成に不可欠である。

これまで、ケルン・南条本[1]は法華經原典として研究者に多大の便宜を供してきた。しかし、バルフ[2]の指摘するごとく、ネパール系写本とカシュガル写本を混淆し、両系統に相違がある場合は、カシュガル写本が中期インド・アリアン語(Middle Indo-Aryan = MIA)の要素を残していると見做して、その読みを採用したことは多くの問題を含んでいる。また、荻原・土田本[3]にしても、校訂上の個々の問題には多くの示唆を与えつつも、エジャートン[4]等によって批判されたよ

表1. 写本分類

写本	材質	編纂年	写本	材質	編纂年
Kn	紙	1908-12.	T3	紙	近代
(1) ネパール・チベット写本	K	貝葉 1070年	T4	紙	近代 1799 / 1800, 1806年
	Pk	貝葉 1082年	T5	紙	
	C1	紙 近代	T6	貝葉	11世紀
	C2	紙 近代	T7	貝葉	11世紀
	C3	貝葉 11世紀	T8	紙	
	C4	貝葉 1039, 1036 / 1037年	T9	紙	近代
	C5	貝葉 1064 / 1065, 1063 / 1064年	A1	紙	1680, 1679 / 1680年
	C6	貝葉 1093, 1091 / 1092年	A2	紙	1713, 11711 / 1712年
	B	貝葉 11 / 12 世紀	A3	紙	
	R	紙 18世紀	N1	貝葉	
	P1	紙 19世紀	N2	貝葉	
	P2	紙 1826年	N3	貝葉	
	P3	紙 19世紀	(2) カシミール写本		
	T2	貝葉 11世紀	D2	白樺樹皮	
			D3	白樺樹皮	
			(3) 中央アジア写本		
			O		

うに、韻律の解釈に多くの問題を残している。

しかしながら、30を超える梵文写本を言語学的研究のみにて、グループごとに分類することは容易ではない。なぜならば、同一の詩節であっても、Aという詩脚で分類できた系統とBという詩脚で分類できた系統とが合致しないということが起こり得る。梵文法華經の写本において、詩節や詩脚の数が増えれば増えるほど系統が複雑に錯綜するようになるのが実情である。このように錯綜した系統解明は、言語学的研究のみにては不可能であり、多変量解析が適している。

山崎は「方便品」(vv. 42-70) のネパール系諸写本を底本として、写本の様態を概観し、次いで同一詩脚に於ける異読、つまり異なって使用された語彙や語順の違う語句を拾い出した[5]（表2参照）。我々は、そのデータに対して数量化第3類（主成分分析に相当）及びクラスター解析を適用して、ネパール系諸写本の系統分類（貝葉写本の系統解明と各紙写本がそれら貝葉写本のどの系統に属するか）の解明を試みた。

数量化第3類を考慮して得られたクラスター解析結果を図1に示す。この結果、27本

表2. それぞれの語句について各写本の損欠を表示している。空白は写本の欠損を、「無」はその詩脚は存在するが相当語が欠けていることを表す。下線の引かれた写本は貝葉である。

		K	Pk	C1	C2	C3	C4	C5	C6	B	R	P1	P2	P3	T2	T3	T4	T5	T6	T7	T8	T9	A1	A2	A3	N1	N2	N3
42a	śārisuta- śāriputra-	○	○																									
42c	nāyikā tāyin-	○	○	○	○																							
46b	-buddha- -kalpa-	○	○	○	○	○	○	○	○																			
48b	dr̥st̥vā dr̥stā	○	○	○	○	○	○	○	○																			
49b	satva sarva tatva	○	○	○		○	○		○																			
50d	vadāmi teṣām prakāśayāmi	○	○	○	○	○	○	○	○																			
51a	-sāmpad- śuddha virā	○	○			○	○	○	○	○																		
51b	-rūpāya -rūpakarunānvi-	○	○			○	○	○	○	○																		
53b	agrām agryam mahyam	○																										
54b	pī hi 無	○	○	○	○	○	○		○																			
54d	upadarsayanti kathenti loke	○	○	○		別語																						
55c	kārya- kāya- yāna-	○				○	○																					
55d	nayanti(/-e/-a) (hi) yānti vadanti	○	○	○		○																						
61d	ca pra- na pra- pra-	○		○	○																							
62a	śārisutā śāriputra sutā śāriputrā	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○	○
64b	saṭṣu gatiṣū gatiṣu saṭṣu	○	○	○	○																							
64c	kaṭasim vi-vrdh- kaṭasim vi-vrt- kaṭasim vi-vrdh- gatim ca viṣ- gatim ca vidh- gatismi vi-vrdh-	○	○																									
65c	niśrayitvā niścayitvā	○		○																								
67c	dr̥st̥vā dr̥stā	○		○	○	○	○	○	○																			
68b	praśāntā pravṛttā praśāstā	○		○	○	○	○	○	○																			
68c	pūriya pūrayi	○		○	○	○	○	○	○																			

(4 写本) からなる B 群の 2 つに分割できた。B 群は 2 つの系統 (P1, P2, T2) と (T3, T6, K, N2, T7, N2, B, N3, C6, C4, A1, C5, Pk, C3) に分類できる。つまり、2 つの紙写本 P1, P2 は 1 つのパーム写本 T2 と、また 2 つの紙写本 T3, A1 は 12 のパーム写本 T6, K, N2, T7, N2, B, N3, C6, C4, C5, Pk, C3 と密接な関連を持っていることを示した。表 2 をより簡単化 (2 値化) し

たデータに対しても、クラスター分析法によりほぼ同等な結論が得られている [6].

しかしながら、基本となる表2のデータでは以下の問題点がある。この表では、例えば42aの詩偈では、言語学の観点から語彙 *sārisuta* と *sāriputra* が各写本に存在するか否かのみのも異読の特徴を抽出している。表3に示すように、この詩偈にはこれら以外の語彙にも、各写本ごとに差異が存在する。例えば、写本KとC1

を比較してみるとその内容は極めて類似した文字順になっているにもかかわらず、表2のデータでは完全に異なる範疇に分類されてしまう。つまり、語彙の異読が強く強調されたデータになり、これに基づくデータ解析では実際の類似性からかけ離れた結果を導き出す危険性がある。それ故、写本のテキスト自体を直接解析して、このような問題点が、系統分類にどの程度影響を与えるかを検討し、図

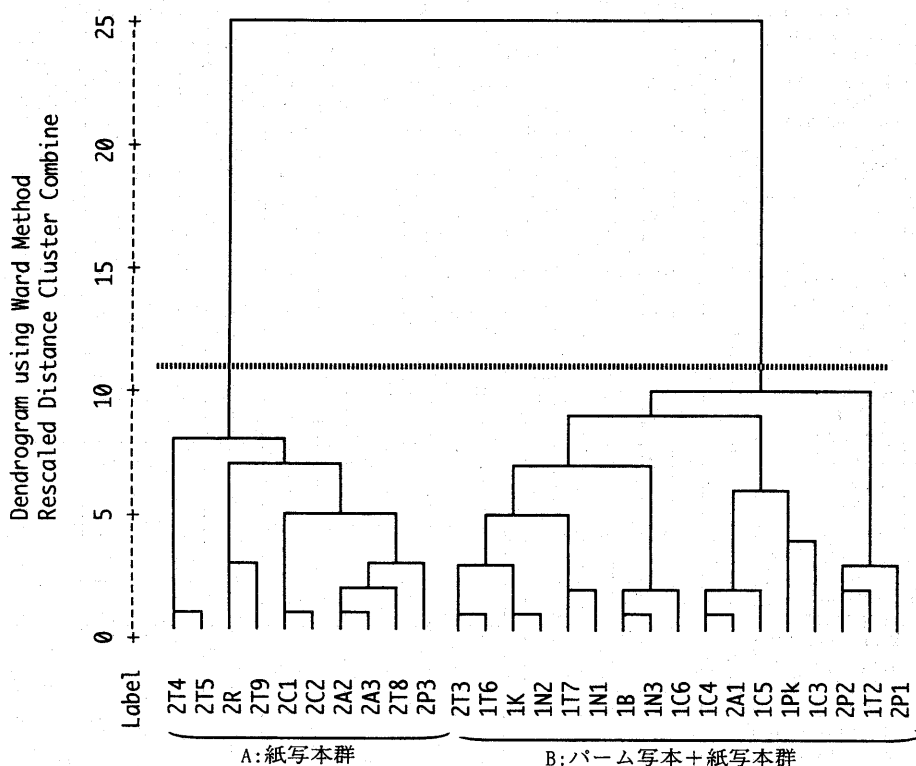


図1. 27種のネパール系写本の言語学的特徴を抽出した結果 (表2) に対するクラスター分析.

1の結果の妥当性について吟味する。もし、図1の結果が妥当であるならば、大量のデータを解析する必要までもなく、言語学的に裏付けられたデータに基づく多変量解析による系統分類の有効性が確立されることになる。

本論文では、ネパール・チベット系のみならず、カシミール (ギルギット)、西域 (中央アジア) を含めた31種の写本のテキスト自体

を直接解析して、系統的な分類について議論する。当該写本の古原型を保持していると考えられている1番から70番までの詩偈 (vv. 1~70), つまり140の半詩偈のテキストを解析対象とした。得られた結果と、写本の特徴を抽出した表2に基づき得られた分類結果と、を比較し、分類の妥当性について議論する。

第2節で、解析用データ作成について議論する。第3、4節では、それぞれ主成分分析、クラスター分析に基づく系統分類について議論する。第5節では、問題点等について検討する。

2. 解析用データ作成

各詩偈の類似性に着目してテキスト自体を直接解析するために、最初に表3のように、各写本の同一半詩偈における語順対応が分かるように整理した。写本N2及びOでは、2番目の単語‘me’が欠けているので、この欠落を示すために、*を挿入した。また、D1とD3では当該詩偈全体が欠落しているので、必要個数だけ（この場合は7個の）*を挿入し、全ての半詩偈において語彙の対応がつくように整理した。

その後、単語の対応関係を保ちながら、各語彙中の共通の文字を比較して類似性を計測した。この数値をKnとKを例に取りながら説明する。表3より分かるように第1番目の語彙では、‘r’、‘n’、‘h’、‘i’の4文字が同じであり、第2番目から7番目までの語彙で共通な文字数は、2、7、7、7、6、13と求められ、総共通文字数は46となる。その後、写本Knを基準としたKの類似性を示す数値を[総共通文字数/Knの文字数(46/50=0.92)]で定義する。このようにして求められた各写本の類似性のデータを表4の編みの列に示す。この結果と表3

表3. Knと31写本の第42ab半詩偈

Kn	śṛṇohi	me	śārisutā	yathaiṣa	saṃbuddha	dharmah	puruṣottamehi
K	śṛṇauhi	me	śārisutā	yathaiṣa	saṃbuddha	dharmmah	puruṣottamehi
Pk	śṛṇohi	me	śārisutā	yathaiṣa	saṃbuddha	dharmma	puruṣottamehi
C1	śṛṇohi	me	śāriputra	yathaiṣa	saṃbuddha	dharmmah	puruṣottamehi
C2	śṛṇohi	me	śāriputra	yathaiṣa	saṃbuddha	dharmmah	puruṣottamehi
C3	śṛṇohi	me	śāriputrē	yathaiṣa	saṃbuddha	dharmah	puruṣottamehi
C4	śṛṇohi	me	śārisutā	yathaiṣa	saṃbuddha	dharmah	puruṣottamehi
C5	śṛṇohi	me	śārisutā	yathaiṣa	saṃbuddha	dharmah	puruṣottamehi
C6	śṛṇohi	me	śārisuta	yathaiva	saṃbuddha	dharmah	puruṣottamehi
B	śṛṇohi	me	śārisutā	yathaivaḥ	saṃbuddha	dharmah	puruṣottamehi
R	śṛṇohi	me	śāriputra	yaṣa	saṃbuddha	dharmma	puruṣottamehi
P1	śṛṇohi	me	śārisutā	yathaivaṃ	saṃbuddha	dharmah	puruṣottamehi
P2	śṛṇohi	me	śārisutā	yathaivaṃ	saṃbuddha	dharmmah	puruṣottamehi
P3	śṛṇohi	me	śārisuta	yathaiṣa	saṃbuddha	dharmmah	puruṣottamehi
T2	śṛṇohi	me	śārisutā	yathaivaṃ	saṃbuddha	dharmah	puruṣottamehi
T3	śṛṇohi	me	śāriputrā	yathaiva	saṃbuddha	dharmmah	puruṣottamehi
T4	śṛṇohi	me	śārisuta	yathaiṣa	saṃbuddha	dharmma	puruṣottamehi
T5	śṛṇohi	me	śārisuta	yathaiṣa	saṃbuddha	dharmma	puruṣottamehi
T6	ṇohi	me	śāriputrā	yathaivaṃ	saṃbuddha	dharmah	puruṣottamehi
T7	śṛṇohi	me	śārisutā	yathaiva	saṃbuddha	dharmah	puruṣottamehi
T8	śṛṇohi	me	śārisuta	yathaiva	saṃbuddha	dharmah	puruṣottamehi
T9	śṛṇohi	me	śāriputra	yathaiṣa	saṃbuddha	dharmma	puruṣottamehi
A1	śṛṇohi	me	śārisutā	yathaiṣa	saṃbuddha	dharmmah	puruṣottamehi
A2	śṛṇohi	me	śāriputra	yathaiṣa	saṃbuddha	dharmmah	puruṣottamehi
A3	śṛṇohi	me	śārisuta	yathaiva	saṃbuddha	dharmmah	puruṣottamehi
N1	śṛṇohi	me	śārisutā	yathai	saṃbuddha	dharmma	puruṣottamehi
N2	śṛṇohi	*	śārisutā	yathaiva	saṃbuddha	dharmah	puruṣottamehi
N3	śṛṇohi	me	śārisuta	yathaiva	saṃbuddha	dharmah	puruṣottamehi
D1	*	*	*	*	*	*	*
D2	śṛṇohi	me	śārisutā	yathaiṣa	saṃbuddha	dharmah	puruṣottamehi
D3	*	*	*	*	*	*	*
O	śṛṇohi	*	śārisutā	yathaiṣa	saṃbuddha	dharmam	puruṣottamena

表4. Knと31の写本の類似性の評価値

h.v. texts	1ab	1cd	..	42ab	..	70ab	70cd
Kn	1.00	1.00	..	1.00	..	1.00	1.00
K	0.97	0.91	..	0.92	..	0.80	0.92
PK	1.00	0.91	..	0.94	..	0.90	0.96
C1	1.00	1.00	..	0.96	..	0.92	1.00
C2	1.00	1.00	..	0.96	..	0.93	1.00
C3	0.76	0.27	..	0.96	..	0.88	0.88
C4	1.00	0.82	..	0.96	..	0.93	0.98
C5	1.00	0.77	..	0.96	..	0.87	1.00
C6	0.97	0.86	..	0.94	..	0.83	0.84
B	1.00	0.82	..	0.98	..	0.87	0.94
R	1.00	0.95	..	0.86	..	0.93	0.98
P1	0.97	0.95	..	0.94	..	0.90	1.00
P2	0.97	0.95	..	0.98	..	0.88	0.98
P3	1.00	0.95	..	0.98	..	0.88	1.00
T2	0.03	0.05	..	0.92	..	0.87	0.56
T3	1.00	0.95	..	0.96	..	0.90	0.96
T4	0.97	1.00	..	0.94	..	0.92	0.98
T5	0.97	0.91	..	0.94	..	0.88	0.96
T6	1.00	0.95	..	0.90	..	0.90	0.92
T7	0.03	0.05	..	0.98	..	0.87	0.94
T8	1.00	0.95	..	0.94	..	0.90	1.00
T9	0.97	0.91	..	0.92	..	0.88	1.00
A1	1.00	0.82	..	0.92	..	0.87	0.98
A2	1.00	1.00	..	0.96	..	0.90	0.92
A3	1.00	0.95	..	0.94	..	0.92	1.00
N1	1.00	0.82	..	0.92	..	0.85	0.96
N2	0.97	0.91	..	0.92	..	0.87	0.90
N3	1.00	0.95	..	0.96	..	0.90	0.58
D1	0.97	0.91	..	0.00	..	0.02	0.04
D2	0.03	0.05	..	1.00	..	0.95	0.94
D3	0.30	0.41	..	0.00	..	0.60	0.70
O	1.00	0.68	..	0.90	..	0.80	0.80

のテキスト例と比較すると、この数値が各テキストの類似性をよりよく表しているように思われる。

表4形式のデータは、当該写本の古原型を保持していると考えられている1番から70番までの詩偈(vv.1~70)、つまり140の半詩偈に対して求めた。この範囲内の1バイトの全文字数はおよそ22万であり、1写本当たりで平均7000文字、1半詩偈当たりで50文字相当である。統計解析の対象としては十分なだけの資料数であろう。また、解析対象の数値は、各半詩偈ごとに規格化されているために、1行の文字数の多寡に依存しないようになっている。以上のことから判断して、表4に統計解析の手法を適用して言語学的に妥当な結果が得られると予想される。

3. 主成分分析

ここでは前節で評価した表4に主成分分析法を適用し、梵文『法華経』の系統分類について議論する。以下の結果は、統計解析ツールSPSSを使用して得られた。

図2に第14主成分までの累積寄与率を示している。第1, 2, 3主成分の固有値はそれぞれ、65.5, 41.0, 11.6である。これらの主成分の累積寄与率への寄与は46.5%, 29.1%, 8.3%であり、第3主成分までで83.5%と相当高い数値になる。それ故、第3主成分までで分類に関して信頼できる議論ができるであろう。

第3主成分まで考慮した3次元の得点分布図(第5節の図6参照)を作成し検討したが、第2主成分までの特性を変更するようなデータは存在しなかった。それ故、第2主成分までの得点分布図で当該文献の分類について議論する。図3及び4が各写本の第1と第2主成分得点分布を示している。図4は図3の右側の点線で囲まれた領域の拡大図である。図3及び4は次の特徴を有する。

- (i) 図2の横軸(第1主成分得点)は、表4における各写本の全詩偈の数値の平均値

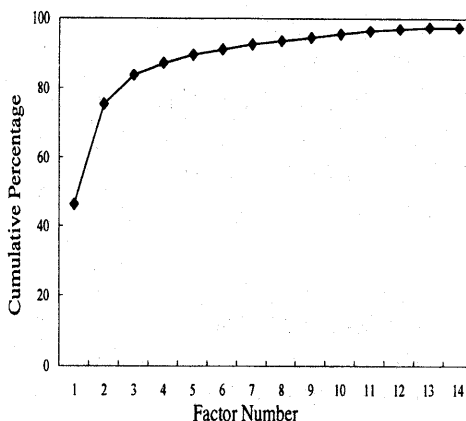


図2. 累積寄与率

に対応している。最も左側に位置する写本D3の平均値は0.265であり、写本D1, N3, D2, T2の順に右側に行くにつれその数値は、0.485, 0.506, 0.582, 0.671と大きくなり、Knでは1になる。又図3の左から右に向かって順に位置する写本O, C6, T7, K, P1ではその平均値が、0.796, 0.840, 0.841, 0.894, 0.881であり、右側に行くにつれ平均値の大きな写本が分布する傾向がある。なお、図4の点線の枠内では、その数値は0.906から0.94までの狭い範囲に多数の写本が分布している。

- (ii) これらの平均値の差異はどのような理由に起因するのであろうか? 写本D3, D1, N3, D2, T2では、考慮した141詩偈中それぞれ、71, 68, 61, 45, 35の詩偈が欠落している。これらのテキストに対する平均値が小さくなっているのは、明らかにこれらのテキストにおいて欠落している詩偈が多いことに起因している。一方これら以外の写本(図3の右側の点線中に入っている写本)では、欠落している詩偈が殆ど見られないのに、その平均値が変動しているのは、各写本の特性を反映していると考えられる。

- (iii) 図3の縦軸(第2主成分得点)の上の方には、詩偈番号の前半で表4の数値が小さい(つまりがKnとの類似性が低い)写

本が分布している。下の方では、逆に詩偈番号の後半部分でKnとの類似性が低い写本が分布している。また、第2主成分得点が0付近では、全詩偈に亘ってKnとの類似性が高い写本が分布している。

よって(i)～(iii)より、図3の右側の点線で囲まれた部分、つまり図4の写本間の位置の違いは、明らかに写本の異同に起因する、と結論することができる。この図から次の結論が得られる。

- ①中央アジア写本である写本Oはチベット版パーム写本C6と近い関係にある。
- ②チベット版パーム写本C6, T7, K, C5は、お互い同士にあまり親近性がなく、あったとしても希薄であろう。
- ③紙写本のうち幾つかは、残りのチベット版パーム写本, C3, N1, B, Pk, N2, C4, T6に似ていると思われる。(図4の点線で囲まれた部分参照。この特性については次節でもっと詳しく吟味する)

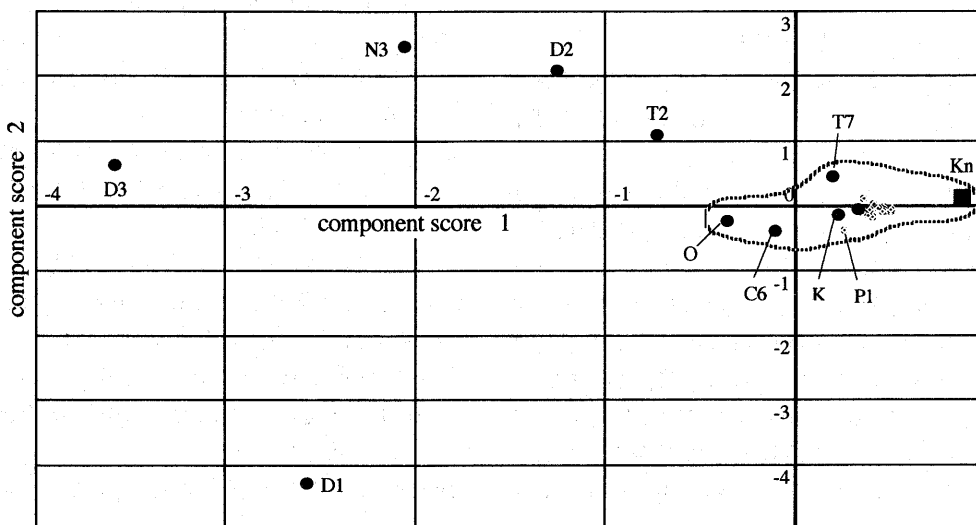


図3. Knと31の写本の第1, 2主成分得点分布。大きな(小さな)丸が頁葉(紙)写本を示している。

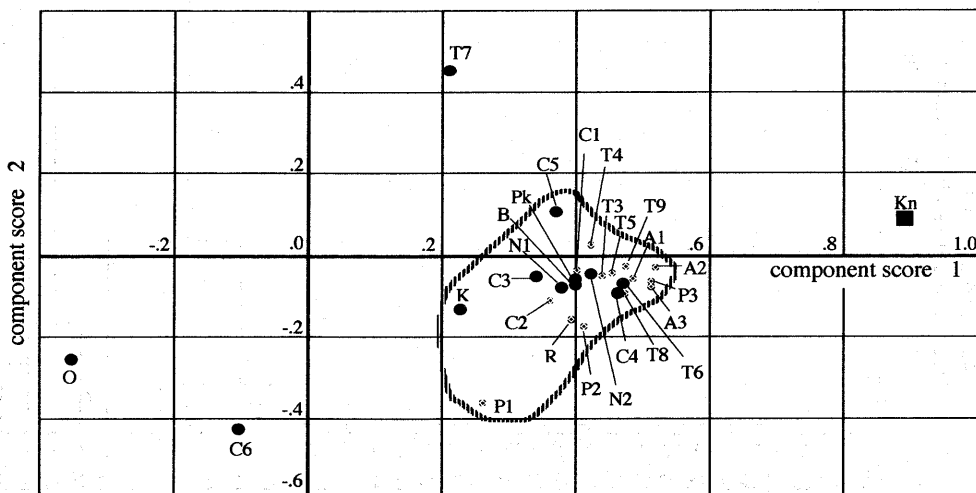


図4. 図3の点線の部分の拡大図

④表4の数値の基準にとった写本Knは、結果的に他の写本群から少し離れた位置に分布した。

①②の特性は、古い写本形式であるパーム群は、歴史的にも地理的にも極めて広い範囲に分布していることに起因していると思われる。③の特性は、紙写本群がパーム写本群の影響の下に編纂された可能性が高いことを示しているのであろう。以上のことを次節では、より詳しく議論する。

4. クラスタ分析

本節では、前節で得られた分類特性を基準にして、クラスタ分析により詳細に議論する。

クラスタ分析では、使用する連結法・距離(Measure)により種々の異なる結果が得られるという困難な点がある。適切な手法の決定は、主成分分析法による分類結果、図3と4に相当する結果が得られるかどうかを判定基準にする。特に、図4の結果(前節の④)より判断して、写本Knが一つだけに括られ、又、O、C6、T7が他の残りの写本群から充分離れているようなクラスタ分析法が採用される

べきであろう。多くの手法を適用し、検討した結果、連結法としてウード法が最適であり、距離としては都市ブロック距離が適していることを見いだした。

図5aは表4に対して、クラスタ分析を適用したときの結果を示している。写本Knが一つだけに括られるような切断は、点線の部分での分割により実現できる。明らかに、図の右側の5つの写本D1、D2、D3、N3、T2は残りの27の写本より、十分に離れている。この5つの写本は次の様に解釈される。

①カシミール系写本D2とネパール系写本T2は類似性が高い。また、D3、N3もこのグループに近い。

②D1は大枠では、貝葉写本群に分類されるのであろうが、どの写本群からも離れているように見える。

図5aの左側の27の写本群に対して、再度同じ手法を適用した結果を図5bに示す。この結果は以下の特性を有する。

③最も左側の部分(A群)は10写本群からなり、序論の図1における紙写本群と完全に一致している。

④B群は11写本群からなり、図1のB群

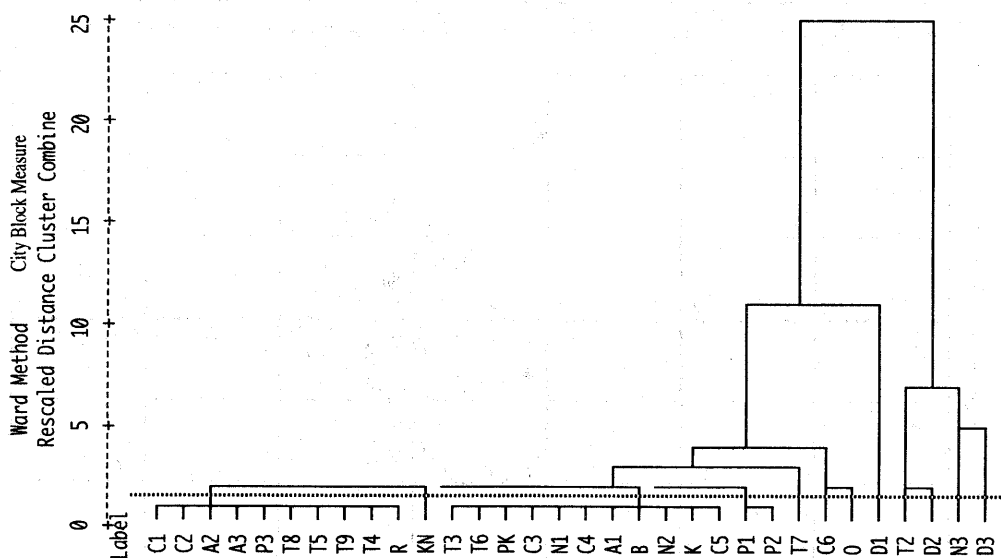


図5a. Knと31の写本に対するクラスタ分析

(17写本)から6写本(T7, N3, C6, P2, P1, T2)が欠けている。B群は基本的に貝葉写本群であり、その中に2つの紙写本T3とA1が混じっている。

- ⑤ 6写本の中、2写本P2, P1はそれだけでまとまっている。この中には、T2が所属していない(この点が図1と異なっている)。
- ⑥ 6写本の中、2写本N3, T2は欠落している詩偈が多いために、本分類ではB群より離れている(①参照)。
- ⑦ T7は他の写本から離れているが、このことは図6からも理解できる。。
- ⑧ 中央アジア系写本Oとネパール系写本C6は多少類似性がある。

これらのことは、図1の結果が基本的には妥当であることを示している。但し、テキスト全体で見た場合、欠落部分が多い写本については、慎重な吟味を必要とする。結局、分類特性の大枠を議論する場合には、大量のテキストデータを解析するまでもなく、言語学的に裏付けられたデータに統計解析手法を適用するだけで十分である、ことが結論として

得られる。

5. 討論

最初に、主成分分析の第3主成分まで考慮した場合、第3節の結果に対してどのような要素が付加されるか議論する。又、数量化理論第4類に基づく結果についても議論する。

図6は、図4に第3主成分を付加した時の各写本の得点分布図を示している。パーム写本C5, C6, K及び中央アジア写本Oの第3主成分得点は殆ど同じ(高さがほぼ一定)である。但し、パーム写本T7は第3主成分得点が小さくなっているが、この写本自体は元々他から離れている。従って、第3節で指摘したように、第3主成分には第2主成分までの特性(第3節の結論①~④)を変える因子がないことが分かる。

次に数量化理論第4類の解析について簡単に述べる。この手法を適用するために、各テキスト間の類似度を表す数値を表4と同じようにして求めた。この手法による類似度の解析では、残念ながら第1, 2, 3固有値成分による分類によって、明瞭な結論を得るこ

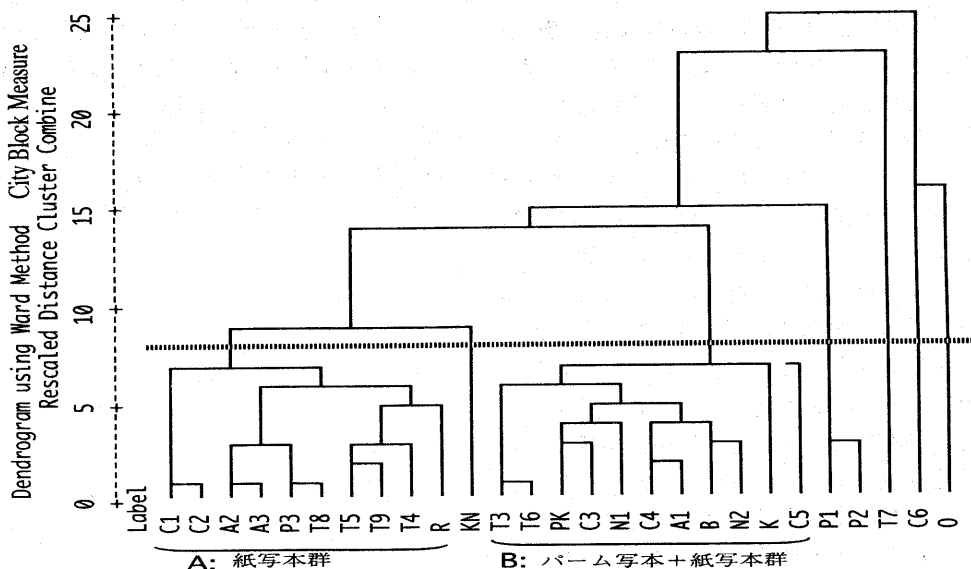


図5b. 27の写本に対するクラスター分析。図5aの右側5個の写本を除いた写本のクラスター分析。点線は図5aの点線と対応している。

とができなかった。この手法では、表4の数値でいうと（正確にはこの数値ではないが基本的には類似の性格を持つ）、その平均値の小さい写本、つまりD3、D1、N3、D2、T2の特性を強く抽出しすぎている。その結果、他の写本を分類できなくなっている。この理由で、第4類はこの解析には不適となっているのであろう。

本論文の結論は次のようにまとめられる。我々は、クラスター分析に主成分分析法を援用して、梵文法華經の系統分類について明確

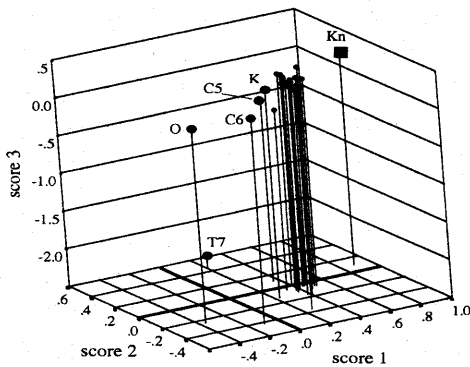


図6. 図4に第3主成分得点を付加した結果。

な結論（図5 a, b及び第4節の結論①～⑧）を得た。最初に大量のテキストデータを整理して、主成分分析に必要なデータを作成し、主成分分析を適用して32種の写本を大きく分類した。その後、同種のデータに対してクラスター分析を適用し、主成分分析結果を再現するような連結法（ウード法）と距離（都市間ブロック距離）を見つけた。得られた結果は、言語学的な議論から特性を抽出した表2（及び簡単化（2値化）した表）に対してクラスター分析を適用して得られた結果（図1）と基本的に一致した。結局、分類特性の大枠を議論する場合には、大量のテキストデータを解析するまでもなく、言語学的に裏付けられたデータを統計解析するだけで十分である、ことが分かった。逆に言うところのこ

とは、テキスト自体を直接解析する統計分類の手法の妥当性を示していると言えよう。

参考文献

- [1] Saddharmapuṇḍarika, ed. by H. Kern & B. Nanjio, *Bibliotheca Buddhica* X, St-petersbourg 1908-12, Reprint, Osnabrück 1970.
- [2] W. Baruch, *Beiträge zum Saddharmapuṇḍarikasūtra*, Leiden 1938.
- [3] Saddharmapuṇḍarika, *Romanized and Revised Text of the Bibliotheca Buddhica publication by Consulting a Sanskrit MS. and Tibetan and Chinese Translations*, by U. Wogihara and C. Tsuchida, Tokyo 1934-35.
- [4] F. Edgerton, "The Meter of the Saddharmapuṇḍarika", *Kuppuswami Sastri Commemoration Volume*, Madras 1936, pp. 39, 41.
- [5] 山崎守一, 「法華經伝承の一樣相 — Upāyakaśālyā-Parivartō Nāma Dvitiyāḥ v. 64c —」, 『法華文化研究』（平成12年3月）第26号, pp. 1-18.
山崎守一, 「梵文法華經校訂の試み—第2章「方便品」(vv. 42-70)を中心に—」, 『法華經の思想と展開』（法華經研究 XIII）（平成13年3月）, pp. 191-230.
山崎守一, 逢坂雄美, 「法華經ネパール系写本の系統分類—統計解析による試み—」, 『石上善應教授古稀記念論文集・仏教文化の基調と展開』（平成13年5月）, pp. 323-336.
- [6] 逢坂雄美, 山崎守一, 「梵文『法華經』の統計解析」, 『人文科学とコンピュータシンポジウム』（平成11年9月）, pp. 105-106.