

Universal Multiple-Octet Coded Character Set
International Organization for Standardization
Organisation Internationale de Normalisation
Международная организация по стандартизации

Doc Type: Working Group Document**Title: Response to Ngô Trung Việt on feedback from Cham experts****Source: Michael Everson, EGT (IE)****Status: Expert Contribution****Action: For consideration by JTC1/SC2/WG2 and UTC****Date: 1999-02-01**

Ngô Trung Việt was kind enough to forward a fax from Cham expert Bùi Khánh Thế to me, and to translate the text of that fax into English for me. I believe that the coding model which I have proposed for Cham will handle each of the cases which Bùi Khánh Thế has asked us to consider.

The “virama model” proposed for encoding Cham is the model normally used in the UCS (ISO/IEC 10646 and Unicode) for scripts derived from Ashoka’s Brahmi script. (Exceptions to this are Thai and Lao, which are based on the “graphic model” of the Thai industrial standard, and Tibetan, which uses the “subjoin model”, since Tibetan has special stacking features.)

1. Representation of consonant clusters. The virama model uses the virama (Cham character U+xx3F) to “kill” the inherent -a vowel of a consonant, and often causes a following consonant to change in some way to join with the first consonant. The font, not the underlying encoding, is responsible for presenting the resulting syllable clusters to the user in the correct shape. The relative order of the signs as *drawn* is therefore irrelevant to the encoding (inputting via the keyboard is a separate issue). Examples in Bùi Khánh Thế’s fax are easy to represent with the virama model:

Consider	$\text{𑄑} + \text{𑄓} + \text{𑄚} = \text{𑄑𑄚}$	$ka + \text{VIRAMA} + ya = kya$
	$\text{𑄑} + \text{𑄓} + \text{𑄛} = \text{𑄑𑄛}$	$ka + \text{VIRAMA} + ra = kra$
	$\text{𑄑} + \text{𑄓} + \text{𑄜} = \text{𑄑𑄜}$	$ka + \text{VIRAMA} + la = kla$
	$\text{𑄑} + \text{𑄓} + \text{𑄝} = \text{𑄑𑄝}$	$ka + \text{VIRAMA} + va = kwa$

In the graphic model used for Thai, the conjunct forms of *ya*, *ra*, *la*, *va* would be coded separately from the base consonants.

$\text{𑄑} + \text{𑄓𑄚} = \text{𑄑𑄚}$	$ka + Y = kYa (= kya)$
$\text{𑄑𑄓} + \text{𑄚} = \text{𑄑𑄚}$	$R + ka = Rka (= kra)$
$\text{𑄑} + \text{𑄓𑄜} = \text{𑄑𑄜}$	$ka + L = kLa (= kla)$
$\text{𑄑} + \text{𑄓𑄝} = \text{𑄑𑄝}$	$ka + W = kWa (= kwa)$

But note that the underlying representation for these is not the same as the phonetic order. This could make certain text operations, such as sorting, somewhat more difficult because *Rka* must sort under *ka*, not under *ra*. Graphic methods for encoding Sinhala and Myanmar were proposed to WG2 but were rejected, because Sinhala and Myanmar experts agreed that the virama model would work for them. It will be easier for implementors who have software that handles Brahmic scripts like Devanagari, Sinhala, or Myanmar, to adapt this software to Cham if Cham uses the virama model.

Compare	Devanagari	क + ण + र = क्र	ka + VIRAMA + ra = kra
	Sinhala	ක + ට් + ර = ක්‍ර	ka + VIRAMA + ra = kra
	Myanmar	က + ိ + ရ = ကြ	ka + VIRAMA + ra = kra
	Cham	ꨀ + ꨑ + ꨐ = ꨀꨑꨐ	ka + VIRAMA + ra = kra

2. Representation of final consonants. The final syllables described in Bùì Khánh Thế’s fax can also be correctly rendered by the font, by simply entering the basic characters in sequence:

ꨀ + ꨑ + ꨐ = ꨀꨑꨐ	ka + i + -m̃ = kim̃
ꨀ + ꨑꨐ + ꨐ = ꨀꨑꨐꨐ	ka + au + ng = kaung
ꨀ + ꨑꨐ + ꨐ = ꨀꨑꨐꨐ	ka + au + -m̃ = kaur̃m̃

3. Representation of other syllable-final consonants. The virama model handles other syllable-final consonants quite easily. One encodes them by using the ZERO-WIDTH NON-JOINER (ZWNJ).

ꨀ + ꨑ + ZWNJ = ꨀꨑ	ka + VIRAMA + ZWNJ = k.
ꨀ + ꨑ + ZWNJ = ꨀꨑ	ta + VIRAMA + ZWNJ = t.

If you write ꨀ + ꨑ + ZWNJ + ꨐ, you will get ꨀꨑꨐ (k.ra). Again, inputting software can allow Cham users to type according to their preference regardless of the underlying encoding.

4. Representation of medial and final vowels. The point raised by Bùì Khánh Thế here is a little bit more controversial, but again, the Brahmic model offers a fairly easy solution. The vowels appear to be decomposable in the writing system, but as linguistic entities they are indivisible units. In some other Brahmic scripts (such as Bengali, Malayalam, Oriya, Sinhala, Tamil), similar “multipart” vowels have been encoded in a similar fashion. The font can handle the sequences.

ꨀ + ꨑ̃ = ꨀꨑ̃	ka + ū = kū
ꨀ + ꨑ̃ = ꨀꨑ̃	ka + e = ke
ꨀ + ꨑ̃ = ꨀꨑ̃	ka + ai = kai
ꨀ + ꨑ̃ = ꨀꨑ̃	ka + au = kau
ꨀ + ꨑ̃ = ꨀꨑ̃	ka + ǔ = kǔ

One *could* encode these graphically.

ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	ka + ā + u = kāu = kū
ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	e + ka + o' = ek'o' = ke
ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	e + ka + ā = ekā = kai
ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	e + ka + ǔ = ekǔ = kau
ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	ka + ā + ǔ = kāǔ = kǔ

Or three of these could, in principle, be entered otherwise and still rendered correctly by the font:

ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	ka + e + o' = ke'o' = ke
ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	ka + e + ā = keā = kai
ꨀ + ꨑ̃ + ꨑ̃ = ꨀꨑ̃ꨑ̃	ka + e + ǔ = keǔ = kau

It is better for consistency to have the same coding conventions for each vowel in Cham. This facilitates lexical work, and, as noted above, also makes the software developed for other Brahmic scripts easier to adapt for Cham.

Assuming that the discussion here satisfies Bùi Khánh Thế, there is an important question which I hope he can answer: is the order of the characters in my proposal correct? I have based it on *Từ Điển Chăm Viết*, which I analyzed in order to put the characters into the correct alphabetical order. I am not sure that, algorithmically, the ordering given in *Từ Điển Chăm Viết* is perfect and error-free, but it is the order I endeavoured to follow. In particular I would like confirmation of the position of CHAM SIGN NG (U+xx0B).