

Providing some UTF-8 support via `inputenc`

David Carlisle Frank Mittelbach Chris Rowley*

v1.3c 2026-06-01 printed August 6, 2026

This file is maintained by the L^AT_EX Project team.
Bug reports can be opened (category `latex`) at
<https://latex-project.org/bugs.html>.

Contents

1	Introduction	2
1.1	Background and general stuff	2
1.2	More specific stuff	2
1.3	Notes	3
1.4	Basic operation of the code	3
2	Coding	4
2.1	Housekeeping	4
2.2	Parsing UTF-8 input	4
2.3	Mapping Unicode codes to L ^A T _E X internal forms	9
2.4	Loading Unicode mappings at begin document	14
3	Mapping characters — based on font (glyph) encodings	15
3.1	About the table itself	15
3.2	The mapping table	16
3.3	Notes	28
3.4	Mappings for OT1 glyphs	28
3.5	Mappings for OMS glyphs	29
3.6	Mappings for TS1 glyphs	29
3.7	Mappings for <code>latex.ltx</code> glyphs	29
3.8	Old <code>utf8.def</code> file as a temp fix for pT _E X and friends	30
4	A test document	35

*Borrowing heavily from tables by Sebastian Rahtz; some table and code cleanup by Javier Bezos

1 Introduction

1.1 Background and general stuff

For many reasons what this package provides is a long way from any type of ‘Unicode compliance’.

In stark contrast to 8-bit character sets, with 16 or more bits it can easily be very inefficient to support the full range.¹ Moreover, useful support of character input by a typesetting system overwhelmingly means finding an acceptable visual representation of a sequence of characters and this, for L^AT_EX, means having available a suitably encoded 8-bit font.

Unfortunately it is not possible to predict exactly what valid UTF-8 octet sequences will appear in a particular file so it is best to make all the unsupported but valid sequences produce a reasonably clear and noticeable error message.

There are two directions from which to approach the question of what to load. One is to specify the ranges of Unicode characters that will result in some sensible typesetting; this requires the provider to ensure that suitable fonts are loaded and that these input characters generate the correct typesetting via the encodings of those fonts. The other is to inspect the font encodings to be used and use these to define which input Unicode characters should be supported.

For Western European languages, at least, going in either direction leads to many straightforward decisions and a few that are more subjective. In both cases some of the specifications are T_EX specific whilst most are independent of the particular typesetting software in use.

As we have argued elsewhere, L^AT_EX needs to refer to characters via ‘seven-bit-text’ names and, so far, these have been chosen by reference to historical sources such as Plain T_EX or Adobe encoding descriptions. It is unclear whether this ad hoc naming structure should simply be extended or whether it would be useful to supplement it with standardized internal Unicode character names such as one or more of the following:²

```
\ltxutwochar <4 hex digits>

\ltxuchar {<hex digits>}
  B H U R R R

\ltxueightchartwo <2 utf8 octets as 8-bit char tokens>
\ltxueightcharthree <3 utf8 octets ...>
\ltxueightcharfour <4 utf8 octets ...>
```

1.2 More specific stuff

In addition to setting up the mechanism for reading UTF-8 characters and specifying the L^AT_EX-level support available, this package contains support for some default historically expected T_EX-related characters and some example ‘Unicode definition files’ for standard font encodings.

¹In fact, L^AT_EX’s current 8-bit support does not go so far as to make all 8-bit characters into valid input.

²Burkhard und Holger Mittelbach spielen mit mir! Sie haben etwas hier geschrieben.

1.3 Notes

This package does not support Unicode combining characters as $\TeX$ is not really equipped to make this possible.

No attempt is made to be useful beyond Latin, and maybe Cyrillic, for European languages (as of now).

1.4 Basic operation of the code

The `inputenc` package makes the upper 8-bit characters active and assigns to all of them an error message. It then waits for the input encoding files to change this set-up. Similarly, whenever `\inputencoding` is encountered in a document, first the upper 8-bit characters are set back to produce an error and then the definitions for the new input encoding are loaded, changing some of the previous settings.

The 8-bit input encodings currently supported by `inputenc` all use declarations such as `\DeclareInputText` and the like to map an 8-bit number to some $\LaTeX$ internal form, e.g. to `"a`.

The situation when supporting UTF-8 as the input encoding is different, however. Here we only have to set up the actions of those 8-bit numbers that can be the first octet in a UTF-8 representation of a Unicode character. But we cannot simply set this to some internal $\LaTeX$ form since the Unicode character consists of more than one octet; instead we have to define this starting octet to parse the right number of further octets that together form the UTF-8 representation of some Unicode character.

Therefore when switching to `utf8` within the `inputenc` framework the characters with numbers (hex) from `"C2` to `"DF` are defined to parse for a second octet following, the characters from `"E0` to `"EF` are defined to parse for two more octets and finally the characters from `"F0` to `"F4` are defined to parse for three additional octets. These additional octets are always in the range `"80` to `"BF`.

Thus, when such a character is encountered in the document (so long as expansion is not prohibited) a defined number of additional octets (8-bit characters) are read and from them a unique control sequence name is immediately constructed.

This control sequence is either defined (good) or undefined (likely); in the latter case the user gets an error message saying that this UTF-8 sequence (or, better, Unicode character) is not supported.

If the control sequence is set up to do something useful then it will expand to a $\LaTeX$ internal form: e.g. for the `utf8` sequence of two octets `"C3 "A4` we get `"a` as the internal form which then, depending on the font encoding, eventually resolves to the single glyph ‘latin-a-umlaut’ or to the composite glyph ‘latin-a with an umlaut accent’.

These mappings from (UTF-8 encoded) Unicode characters to $\LaTeX$ internal forms are made indirectly. The code below provides a declaration `\DeclareUnicodeCharacter` which maps Unicode numbers (as hexadecimal) to $\LaTeX$ internal forms.

This mapping needs to be set up only once so it is done at `\begin{document}` by looking at the list of font encodings that are loaded by the document and providing mappings related to those font encodings whenever these are available. Thus at most only those Unicode characters that can be represented by the glyphs available in these encodings will be defined.

Technically this is done by loading one file per encoding, if available, that is supposed to provide the necessary mapping information.

2 Coding

2.1 Housekeeping

The usual introductory bits and pieces:

```

1 <utf8>\ProvidesFile{utf8.def}
2 <test>\ProvidesFile{utf8-test.tex}
3 <+lcy> \ProvidesFile{lcyenc.dfu}
4 <+ly1> \ProvidesFile{ly1enc.dfu}
5 <+oms> \ProvidesFile{omsenc.dfu}
6 <+ot1> \ProvidesFile{ot1enc.dfu}
7 <+ot2> \ProvidesFile{ot2enc.dfu}
8 <+t1> \ProvidesFile{t1enc.dfu}
9 <+t2a> \ProvidesFile{t2aenc.dfu}
10 <+t2b> \ProvidesFile{t2benc.dfu}
11 <+t2c> \ProvidesFile{t2cenc.dfu}
12 <+ts1> \ProvidesFile{ts1enc.dfu}
13 <+x2> \ProvidesFile{x2enc.dfu}
14 <+all> \ProvidesFile{utf8enc.dfu}
15 <-utf8-2018> [2026-06-01 v1.3c UTF-8 support]
16 <*utf8>

```

This is a temporary fix for the e-pTeX / e-upTeX engines that do not yet have a `\ifincsname` primitive. Once this is available the extra file will be dropped.

```

17 \ifx\ifincsname\@undefined % old e-pTeX or e-upTeX engines
18 \input utf8-2018.def
19 \expandafter\@firstofone
20 \else
21 \expandafter\@gobble
22 \fi
23 \endinput
24 \makeatletter

```

We restore the `\catcode` of space (which is set to ignore in `inputenc`) while reading `.def` files. Otherwise we would need to explicitly use `\space` all over the place in error and log messages.

```

25 \catcode'\ \saved@space@catcode

```

2.2 Parsing UTF-8 input

A UTF-8 char (that is not actually a 7-bit char, i.e. a single octet) is parsed as follows: each starting octet is an active TeX character token; each of these is defined below to be a macro with one to three arguments nominally (depending on the starting octet). It calls one of `\UTFviii@two@octets`, `\UTFviii@three@octets`, or `\UTFviii@four@octets` which then actually picks up the remaining octets as the argument(s).

- When typesetting we pick up the necessary number of additional octets, check if they form a command that L^AT_EX knows about (via `\csname u8:\string #1\string #2...\endcsname`) and if so use that for typesetting. `\string` is needed as the octets may (all?) be active and we want the literal values in the name.
- If the UTF-8 character is going to be part of a label, then it is essentially becoming part of some csname and with the test `\ifincsname` we can find this out. If so, we render the whole sequence of octets harmless by using `\string` too when the starting octet executes (`\UTF@...\@octets@string`).
- Another possible case is that `\protect` has *not* the meaning of `\typeset@protect`. In that case we may do a `\write` or we may do a `\protected@edef` or ... In all such cases we want to keep the sequence of octets unchanged, but we can't use `\string` this time, since at least in the case of `\protect@edef` the result may later be typeset after all (in fact that is quite likely) and so at that point the starting octet needs to be an active character again (the others could be stringified). So for this case we use `\noexpand` (`(\UTF@...s@octets@noexpand)`).

`\UTFviii@two@octets` Putting that all together the code for a start octet of a two byte sequence would then look like this:

```

26 \long\def\UTFviii@two@octets{%
27   \ifincsname
28     \expandafter \UTF@two@octets@string
29   \else
30     \ifx \protect\@typeset@protect \else
31       \expandafter\expandafter\expandafter \UTF@two@octets@noexpand
32     \fi
33   \fi
34   \UTFviii@two@octets@combine
35 }
```

`\ifcsname` is tested first because that can be true even if we are otherwise doing typesetting. If this is the case we use `\string` on the whole octet sequence. `\UTF@two@octets@string` not only does this but also gets rid of the command `\UTFviii@two@octets@combine` in the input stream by picking it up as a first argument and dropping it.

If this is not the case and we are doing typesetting (i.e., `\protect` is `\typeset@protect`), then we execute `\UTFviii@two@octets@combine` which picks up all octets and typesets the character (or generates an error if it doesn't know how to typeset it).

However, if we are not doing typesetting, then we execute the command `\UTFviii@two@octets@noexpand` which works like `\UTF@two@octets@string` but uses `\noexpand` instead of `\string`. This way the sequence is temporarily rendered harmless, e.g., would display as is or stays put inside a `\protected@edef`. But if the result is later reused the starting octet is still active and so will be able to construct the UTF-8 character again.

`\UTFviii@three@octets` The definitions for the other starting octets are the same except that they pick up
`\UTFviii@four@octets` more octets after them.

```

36 \long\def\UTFviii@three@octets{%
```

```

37 \ifincsname
38   \expandafter \UTF@three@octets@string
39 \else
40   \ifx \protect\@typeset@protect \else
41     \expandafter\expandafter\expandafter \UTF@three@octets@noexpand
42   \fi
43 \fi
44 \UTFviii@three@octets@combine
45 }

46 \long\def\UTFviii@four@octets{%
47   \ifincsname
48     \expandafter \UTF@four@octets@string
49   \else
50     \ifx \protect\@typeset@protect \else
51       \expandafter\expandafter\expandafter \UTF@four@octets@noexpand
52     \fi
53   \fi
54   \UTFviii@four@octets@combine
55 }

```

\UTFviii@two@octets@noexpand These temporarily prevent the active chars from expanding.

```

56 \long\def\UTF@two@octets@noexpand#1#2#3{\unexpanded{#2#3}}
\UTFviii@three@octets@noexpand 57 \long\def\UTF@three@octets@noexpand#1#2#3#4{\unexpanded{#2#3#4}}
\UTFviii@four@octets@noexpand 58 \long\def\UTF@four@octets@noexpand#1#2#3#4#5{\unexpanded{#2#3#4#5}}

```

\UTFviii@two@octets@string And the same with \string for use in \csname constructions.

```

\UTFviii@three@octets@string 59 \long\def\UTF@two@octets@string#1#2#3{\detokenize{#2#3}}
\UTFviii@four@octets@string 60 \long\def\UTF@three@octets@string#1#2#3#4{\detokenize{#2#3#4}}
61 \long\def\UTF@four@octets@string#1#2#3#4#5{\detokenize{#2#3#4#5}}

```

\UTFviii@two@octets@combine From the arguments a control sequence with a name of the form u8:#1#2... is constructed where the #i (i > 1) are the arguments and #1 is the starting octet (as a T_EX active character token). Since some or even all of these characters are active we need to use \string when building the \csname.

The \csname thus constructed can of course be undefined but to avoid producing an unhelpful low-level undefined command error we pass it to \UTFviii@defined which is responsible for producing a more sensible error message (not yet done!!). If, however, it is defined we simply execute the thing (which should then expand to an encoding specific internal L^AT_EX form).

```

62 \long\def\UTFviii@two@octets@combine#1#2{\expandafter
63   \UTFviii@defined\csname u8:\string#1\string#2\endcsname}

64 \long\def\UTFviii@three@octets@combine#1#2#3{\expandafter
65   \UTFviii@defined\csname u8:\string#1\string#2\string#3\endcsname}

66 \long\def\UTFviii@four@octets@combine#1#2#3#4{\expandafter
67   \UTFviii@defined\csname u8:\string#1\string#2\string#3\string#4\endcsname}

```

\UTFviii@defined This tests whether its argument is different from \relax: it either calls for a sensible error message (not done), or it gets the \fi out of the way (in case the command has arguments) and executes it.

```

68 \def\UTFviii@defined#1{%
69   \ifx#1\relax

```

Test if the sequence is invalid UTF-8 or valid UTF-8 but without a L^AT_EX definition.

```
70 \if\relax\expandafter\UTFviii@checkseq\string#1\relax\relax
```

The endline character has a special definition within the inputenc package (it is gobbling spaces). For this reason we can't produce multiline strings without some precaution.

```
71 \UTFviii@undefined@err{#1}%
```

```
72 \else
```

```
73 \latex@error{Invalid UTF-8 byte sequence (\expandafter
```

```
74 \gobblefour\string#1)}}%
```

```
75 \UTFviii@invalid@help
```

```
76 \fi
```

```
77 \else\expandafter
```

```
78 #1%
```

```
79 \fi
```

```
80 }
```

```
\UTFviii@invalid@err
```

```
\UTFviii@invalid@help
```

```
81 \def\UTFviii@invalid@err#1{%
```

```
82 \latex@error{Invalid UTF-8 byte "\UTFviii@hexnumber{'#1}}%
```

```
83 \UTFviii@invalid@help}
```

```
84 \def\UTFviii@invalid@help{%
```

```
85 The document does not appear to be in UTF-8 encoding.\MessageBreak
```

```
86 Try adding \noexpand\UseRawInputEncoding as the first line of the file.\MessageBreak
```

```
87 or specify an encoding such as \noexpand\usepackage[latin1]{inputenc}.\MessageBreak
```

```
88 in the document preamble.\MessageBreak
```

```
89 Alternatively, save the file in UTF-8 using your editor or another tool}
```

```
\UTFviii@undefined@err
```

```
90 \def\UTFviii@undefined@err#1{%
```

```
91 \latex@error{Unicode character \expandafter
```

```
92 \UTFviii@splitcsname\string#1\relax
```

```
93 \MessageBreak
```

```
94 not set up for use with LaTeX}%
```

```
95 {You may provide a definition with\MessageBreak
```

```
96 \noexpand\DeclareUnicodeCharacter}%
```

```
97 }
```

\UTFviii@checkseq Check that the csname consists of a valid UTF-8 sequence.

```
\UTFviii@check@continue
```

```
98 \def\UTFviii@checkseq#1:#2#3{%
```

```
99 \ifnum'#2<"80 %
```

```
100 \ifx\relax#3\else1\fi
```

```
101 \else
```

```
102 \ifnum'#2<"C0 %
```

```
103 1 %
```

```
104 \else
```

```
105 \expandafter\expandafter\expandafter\UTFviii@check@continue
```

```
106 \expandafter\expandafter\expandafter#3%
```

```
107 \fi
```

```
108 \fi}
```

```

109 \def\UTFviii@check@continue#1{%
110   \ifx\relax#1%
111   \else
112   \ifnum'#1<"80 1\else\ifnum'#1>"BF 1\fi\fi
113   \expandafter\UTFviii@check@continue
114   \fi
115 }

```

`\UTFviii@loop` This bit of code derived from `xmlltex` defines the active character corresponding to starting octets to call `\UTFviii@two@octets` etc as appropriate. The starting octet itself is passed directly as the first argument, the others are picked up later en route.

The `\UTFviii@loop` loops through the numbers starting at `\count@` and ending at `\@tempcnta - 1`, each time executing the code in `\UTFviii@tmp`.

Store current settings so can restore after the loops without using a group (gh/762).

```

116 \edef\reserved@a{%
117 \catcode'\noexpand\~=\the\catcode'\~\relax
118 \catcode'\noexpand\"=\the\catcode'\\"~\relax
119 \uccode'\noexpand\~=\the\uccode'\~\relax
120 \count@=\the\count@\relax
121 \@tempcnta=\the\@tempcnta\relax
122 \let\noexpand\reserved@a\relax}

123 \catcode'\~13
124 \catcode'\\"12

125 \def\UTFviii@loop{%
126   \uccode'\~\count@
127   \uppercase\expandafter{\UTFviii@tmp}%
128   \advance\count@\@ne
129   \ifnum\count@<\@tempcnta
130   \expandafter\UTFviii@loop
131   \fi}

```

Handle the single byte control characters. C0 controls are valid UTF-8 but defined to give the “Character not defined error” They may be defined with `\DeclareUnicodeCharacter`.

```

132   \def\UTFviii@tmp{\protected\edef~{\noexpand\UTFviii@undefined@err{:\string~}}}
133 % 0 ~@ null
134   \count@1
135   \@tempcnta9
136 % 9 ~I tab
137 % 10 ~J nl
138 \UTFviii@loop
139   \count@11
140   \@tempcnta12
141 \UTFviii@loop
142 % 12 ~L
143 % 13 ~M
144   \count@14
145   \@tempcnta32
146 \UTFviii@loop

```

Bytes with leading bits 10 are not valid UTF-8 starting bytes


```

147     \count@"80
148     \@tempcnta"C2
149     \def\UTFviii@tmp{\protected\edef~{\noexpand\UTFviii@invalid@err\string~}}
150 \UTFviii@loop

```

Setting up 2-byte UTF-8: The starting bytes is passed as an active character so that it can be reprocessed later!

```

151     \count@"C2
152     \@tempcnta"E0
153     \def\UTFviii@tmp{\protected\edef~{\noexpand\UTFviii@two@octets\noexpand~}}
154 \UTFviii@loop

```

Setting up 3-byte UTF-8:

```

155     \count@"E0
156     \@tempcnta"F0
157     \def\UTFviii@tmp{\protected\edef~{\noexpand\UTFviii@three@octets\noexpand~}}
158 \UTFviii@loop

```

Setting up 4-byte UTF-8:

```

159     \count@"F0
160     \@tempcnta"F5
161     \def\UTFviii@tmp{\protected\edef~{\noexpand\UTFviii@four@octets\noexpand~}}
162 \UTFviii@loop

```

Bytes above F4 are not valid UTF-8 starting bytes as they would encode numbers beyond the Unicode range

```

163     \count@"F5
164     \@tempcnta"100
165     \def\UTFviii@tmp{\protected\edef~{\noexpand\UTFviii@invalid@err\string~}}
166 \UTFviii@loop

```

Restore values after the loops.

```

167 \reserved@a

```

For this case we must disable the warning generated by `inputenc` if it doesn't see any new `\DeclareInputText` commands.

```

168 \@inpenc@test

```

If this file (`utf8.def`) is not being read while setting up `inputenc`, i.e. in the preamble, but when `\inputencoding` is called somewhere within the document, we do not need to input the specific Unicode mappings again. We therefore stop reading the file at this point.

```

169 \ifx\@begindocumenthook\undefined
170   \makeatother

```

The `\fi` must be on the same line as `\endinput` or else it will never be seen!

```

171   \endinput \fi

```

2.3 Mapping Unicode codes to L^AT_EX internal forms

`\DeclareUnicodeCharacter` The `\DeclareUnicodeCharacter` declaration defines a mapping from a Unicode character code point to a L^AT_EX internal form. The first argument is the Unicode number as hexadecimal digits and the second is the actual L^AT_EX internal form.

We start by making sure that some characters have the right `\catcode` when they are used in the definitions below.

```

172 \begingroup
173 \catcode'\="=12
174 \catcode'\<=12
175 \catcode'\.=12
176 \catcode'\,=12
177 \catcode'\;=12
178 \catcode'\!=12
179 \catcode'\~=13

180 \gdef\DeclareUnicodeCharacter#1#2{%
181   \count@"#1\relax
182   \wlog{ \space\space defining Unicode char U+#1 (decimal \the\count@)}%
183   \begingroup

```

Next we do the parsing of the number stored in `\count@` and assign the result to `\UTFviii@tmp`. Actually all this could be done in-line, the macro `\parse@XML@charref` is only there to extend this code to parsing Unicode numbers in other contexts one day (perhaps).

```

184   \parse@XML@charref

```

Here is an example of what is happening, for the pair "C2 "A3 (which is the utf8 representation for the character £). After `\parse@XML@charref` we have, stored in `\UTFviii@tmp`, a single command with two character tokens as arguments:

```

[tC2 and tA3 are the characters corresponding to these two octets]
\UTFviii@two@octets tC2tA3

```

what we actually need to produce is a definition of the form

```

\def\u8:tC2tA3 {\LaTeX internal form}.

```

So here we temporarily redefine the prefix commands `\UTFviii@two@octets`, etc. to generate the csname that we wish to define; the `\strings` are added in case these tokens are still active.

```

185   \def\UTFviii@two@octets##1##2{\csname u8:##1\string##2\endcsname}%
186   \def\UTFviii@three@octets##1##2##3{\csname u8:##1%
187                                     \string##2\string##3\endcsname}%
188   \def\UTFviii@four@octets##1##2##3##4{\csname u8:##1%
189                                     \string##2\string##3\string##4\endcsname}%

```

Now we simply:-) need to use the right number of `\expandafters` to finally construct the definition: expanding `\UTFviii@tmp` once to get its contents, a second time to replace the prefix command by its `\csname` expansion, and a third time to turn the expansion into a csname after which the `\gdef` finally gets applied. We add an irrelevant `\IeC` and braces around the definition, in order to avoid any space after the command being gobbled up when the text is written out to an auxiliary file (see `inputenc` for further details

```

190   \expandafter\expandafter\expandafter
191   \expandafter\expandafter\expandafter
192   \expandafter
193   \gdef\UTFviii@tmp{\IeC{#2}}%
194   \endgroup
195 }

```

`\parse@XML@charref` This macro parses a Unicode number (decimal) and returns its UTF-8 representation as a sequence of non-active T_EX character tokens. In the original code it

had two arguments delimited by ; here, however, we supply the Unicode number implicitly.

```
196 \gdef\parse@XML@charref{%
```

We need to keep a few things local, mainly the `\uccode`'s that are set up below. However, the group originally used here is actually unnecessary since we call this macro only within another group; but it will be important to restore the group if this macro gets used for other purposes.

```
197 % \begingroup
```

The original code from `xmltex` supported the convention that a Unicode slot number could be given either as a decimal or as a hexadecimal (by starting with `x`). We do not do this so this code is also removed. This could be reactivated if one wants to support document commands that accept Unicode numbers (but then the first case needs to be changed from an error message back to something more useful again).

```
198 % \uppercase{\count@ \if x\noexpand#1\else#1\fi#2}\relax
```

As `\count@` already contains the right value we make `\parse@XML@charref` work without arguments. In the case single byte UTF-8 sequences, only allow definition if the character is already active. The definition of `\UTFviii@tmp` looks slightly strange but is designed for the sequence of `\expandafter` in `\DeclareUnicodeCharacter`.

```
199 \ifnum\count@<"80\relax
200   \ifnum\catcode\count@=13
201     \uccode'\~=\count@\uppercase{\def\UTFviii@tmp{\@empty\@empty~}}%
202   \else
203     \@latex@error{Cannot define non-active Unicode char value < 0080}%
204     \eha
205     \def\UTFviii@tmp{\UTFviii@tmp}%
206   \fi
```

The code below is derived from `xmltex` and generates the UTF-8 byte sequence for the number in `\count@`.

The reverse operation (just used in error messages) has now been added as `\decode@UTFviii`.

```
207 \else\ifnum\count@<"800\relax
208   \parse@UTFviii@a,%
209   \parse@UTFviii@b C\UTFviii@two@octets.,%
210 \else\ifnum\count@<"10000\relax
211   \parse@UTFviii@a;%
212   \parse@UTFviii@a,%
213   \parse@UTFviii@b E\UTFviii@three@octets.{,;%}
214 \else
```

Test added here for out of range values, the 4-octet definitions are still set up so that `\DeclareUnicodeCharacter` does something sensible if the user scrolls past this error.

```
215   \ifnum\count@>"10FFFF\relax
216     \@latex@error
217       {\UTFviii@hexnumber\count@\space too large for Unicode}%
218       {Values between 0 and 10FFFF are permitted}%
219   \fi
```

```

220     \parse@UTFviii@a;%
221     \parse@UTFviii@a,%
222     \parse@UTFviii@a!%
223     \parse@UTFviii@b F\UTFviii@four@octets.{!,,}%
224     \fi
225     \fi
226     \fi
227 %   \endgroup
228 }

```

\parse@UTFviii@a ...so somebody else can document this part :-)

```

229 \gdef\parse@UTFviii@a#1{%
230     \@tempcnta\count@
231     \divide\count@ 64
232     \@tempcntb\count@
233     \multiply\count@ 64
234     \advance\@tempcnta-\count@
235     \advance\@tempcnta 128
236     \uccode'#1\@tempcnta
237     \count@\@tempcntb}

```

\parse@UTFviii@b ...same here

```

238 \gdef\parse@UTFviii@b#1#2#3#4{%
239     \advance\count@ "#10\relax
240     \uccode'#3\count@
241     \uppercase{\gdef\UTFviii@tmp{#2#3#4}}

```

\decode@UTFviii In the reverse direction, take a sequence of octets(bytes) representing a character in UTF-8 and construct the Unicode number. The sequence is terminated by \relax.

In this version, if the sequence is not valid UTF-8 you probably get a low level arithmetic error from \numexpr or stray characters at the end. Getting a better error message would be somewhat expensive. As the main use is for reporting characters in messages, this is done just using expansion, so \numexpr is used, A stub returning 0 is defined if \numexpr is not available.

```

242 \ifx\numexpr\@undefined
243 \gdef\decode@UTFviii#1{0}
244 \else

```

If the input is malformed UTF-8 there may not be enough closing) so add 5 so there are always some remaining then cleanup and remove any remaining ones at the end. This avoids \numexpr parse errors while outputting a package error.

```

245 \gdef\decode@UTFviii#1\relax{%
246     \expandafter\UTFviii@cleanup
247     \the\numexpr\dec@de@UTFviii#1\relax)))))\@empty}
248 \gdef\UTFviii@cleanup#1)#2\@empty{#1}
249 \gdef\dec@de@UTFviii#1{%
250     \ifx\relax#1%
251     \else
252         \ifnum'#1>"EF
253             (((('1-"F0)%

```

```

254 \else
255   \ifnum'#1>"DF
256   ((('#1-"E0)%
257 \else
258   \ifnum'#1>"BF
259   ((('#1-"C0)%
260 \else
261   \ifnum'#1>"7F
262   )*64+('#1-"80)%
263 \else
264   +'#1 %
265 \fi
266 \fi
267 \fi
268 \fi
269 \expandafter\dec@de@UTFviii
270 \fi}
271 \fi

```

\UTFviii@hexnumber Convert a number to a sequence of uppercase hex digits. If **\numexpr** is not available, it returns its argument unchanged.

```

272 \ifx\numexpr\@undefined
273 \global\let\UTFviii@hexnumber\@firstofone
274 \global\UTFviii@hexdigit\hexnumber@
275 \else
276 \gdef\UTFviii@hexnumber#1{%
277 \ifnum#1>15 %
278 \expandafter\UTFviii@hexnumber\expandafter{\the\numexpr(#1-8)/16\relax}%
279 \fi
280 \UTFviii@hexdigit{\numexpr#1\ifnum#1>0-((#1-8)/16)*16\fi\relax}%
281 }

```

Almost but not quite **\hexnumber@**.

```

282 \gdef\UTFviii@hexdigit#1{\ifcase\numexpr#1\relax
283 0\or1\or2\or3\or4\or5\or6\or7\or8\or9\or
284 A\or B\or C\or D\or E\or F\fi}
285 \fi

```

\UTFviii@splitcsname Split a csname representing a unicode character and return the character and the **\UTFviii@hexcodepoint** unicode number in hex.

```

286 \gdef\UTFviii@hexcodepoint#1{U+%
287 \ifnum#1<16 0\fi
288 \ifnum#1<256 0\fi
289 \ifnum#1<4096 0\fi
290 \UTFviii@hexnumber{#1}%
291 }%
292 \gdef\UTFviii@splitcsname#1:#2\relax{%

```

Need to pre-expand the argument to ensure cleanup in case of mal-formed UTF-8.

```

293 #2 (\expandafter\UTFviii@hexcodepoint\expandafter{%
294 \the\numexpr\decode@UTFviii#2\relax})%
295 }

```

```

296 \endgroup
297 \@onlypreamble\DeclareUnicodeCharacter

```

These are preamble only as long as we don't support Unicode charrefs in documents.

```

298 \@onlypreamble\parse@XML@charref
299 \@onlypreamble\parse@UTFviii@a
300 \@onlypreamble\parse@UTFviii@b

```

2.4 Loading Unicode mappings at begin document

The original plan was to set up the UTF-8 support at `\begin{document}`; but then any text characters used in the preamble (as people do even though advised against it) would fail in one way or the other. So the implementation was changed and the Unicode definition files for already defined encodings are loaded here.

We loop through all defined font encodings (stored in `\cdp@list`) and for each load a file *nameenc.dfu* if it exist. That file is then supposed to contain `\DeclareUnicodeCharacter` declarations.

```

301 \begingroup
302   \def\cdp@elt#1#2#3#4{%
303     \wlog{Now handling font encoding #1 ...}%
304     \lowercase{%
305       \InputIfFileExists{#1enc.dfu}}%
306       {\wlog{... processing UTF-8 mapping file for font %
307         encoding #1}%

```

The previous line is written to the log with the newline char being ignored (thus not producing a space). Therefore either everything has to be on a single input line or some special care must be taken. From this point on we ignore spaces again, i.e., while we are reading the *.dfu* file. The `\endgroup` below will restore it again.

```

308       \catcode'\ 9\relax}%
309       {\wlog{... no UTF-8 mapping file for font encoding #1}}%
310   }
311   \cdp@list
312 \endgroup

```

However, we don't know if there are font encodings still to be loaded (either with `fontenc` or directly with `\input` by some package). Font encoding files are loaded only if the corresponding encoding has not been loaded yet, and they always begin with `\DeclareFontEncoding`. We now redefine the internal kernel version of the latter to load the Unicode file if available.

```

313 \def\DeclareFontEncoding#1#2#3{%
314   \expandafter
315   \ifx\csname T@#1\endcsname\relax
316     \def\cdp@elt{\noexpand\cdp@elt}%
317     \xdef\cdp@list{\cdp@list\cdp@elt{#1}%
318       {\default@family}{\default@series}%
319       {\default@shape}}%
320     \expandafter\let\csname#1-cmd\endcsname\@changed@cmd
321     \begingroup
322       \wlog{Now handling font encoding #1 ...}%
323       \lowercase{%

```

```

324         \InputIfFileExists{#1enc.dfu}}%
325         {\wlog{... processing UTF-8 mapping file for font %
326             encoding #1}}%
327         {\wlog{... no UTF-8 mapping file for font encoding #1}}%
328     \endgroup
329 \else
330     \@font@info{Redefining font encoding #1}%
331 \fi
332 \global\@namedef{T@#1}{#2}%
333 \global\@namedef{M@#1}{\defaultM#3}%
334 \xdef\LastDeclaredEncoding{#1}%
335 }
336 </utf8>

```

3 Mapping characters — based on font (glyph) encodings

This section is a first attempt to provide Unicode definitions for characters whose standard glyphs are currently provided by the standard L^AT_EX font-encodings T1, OT1, etc. They are by no means completed and need checking.

For example, one should check the already existing input encodings for glyphs that may in fact be available and required, e.g. `latin4` has a number of glyphs with the `\=` accent. Since the T1 encoding does not provide such glyphs, these characters are not listed below (yet).

The list below was generated by looking at the current L^AT_EX font encoding files, e.g., `t1enc.def` and using the work by Sebastian Rahtz (in `ucharacters.sty`) with a few modifications. In combinations such as `\^i` the preferred form is that and not `\i`.

This list has been built from several sources, obviously including the Unicode Standard itself. These sources include Passive T_EX by Sebastian Rahtz, the `unicode` package by Dominique P. G. Unruh (mainly for Latin encodings) and `text4ht` by Eitan Gurari (for Cyrillic ones).

Note that it strictly follows the Mittelbach principles for input character encodings: thus it offers no support for using utf8 representations of math symbols such as $\times$ or $\div$ (in math mode).

3.1 About the table itself

In addition to generating individual files, the table below is, at present, a one-one (we think) partial relationship between the (ill-defined) set of LICRs and the Unicode slots "0080 to "FFFF. At present these entries are used only to define a collection of partial mappings from Unicode slots to LICRs; each of these mappings becomes full if we add an exception value ('not defined') to the set of LICRs.

It is probably not essential for the relationship in the full table to be one-one; this raises questions such as: the exact role of LICRs; the formal relationships on the set of LICRs; the (non-mathematical) relationship between LICRs and Unicode (which has its own somewhat fuzzy equivalences); and ultimately what a character is and what a character representation and/or name is.

It is unclear the extent to which entries in this table should resemble the closely related ones in the 8-bit `inputenc` files. The Unicode standard claims that the

first 256 slots ‘are’ ASCII and Latin-1.

Of course, $\text{\TeX}$ itself typically does not treat even many perfectly ‘normal text’ 7-bit slots as text characters, so it is unclear whether $\text{\LaTeX}$ should even attempt to deal in any consistent way with those Unicode slots that are not definitive text characters.

3.2 The mapping table

Note that the first argument must be a hex-digit number greater than 00BF and at most 10FFFF.

There are few notes about inconsistencies etc at the end of the table.

```

337 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00A0}{\nobreakspace}
338 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00A1}{\textexclamdown}
339 <all, ts1, ly1>\DeclareUnicodeCharacter{00A2}{\textcent}
340 <all, ts1, t1, ot1, ly1>\DeclareUnicodeCharacter{00A3}{\textsterling}
341 <all, x2, ts1, t2c, t2b, t2a, ly1, lcy>\DeclareUnicodeCharacter{00A4}{\textcurrency}
342 <all, ts1, ly1>\DeclareUnicodeCharacter{00A5}{\textyen}
343 <all, ts1, ly1>\DeclareUnicodeCharacter{00A6}{\textbrokenbar}
344 <all, x2, ts1, t2c, t2b, t2a, oms, ly1>\DeclareUnicodeCharacter{00A7}{\textsection}
345 <all, ts1>\DeclareUnicodeCharacter{00A8}{\textasciidieresis}
346 <all, ts1, utf8>\DeclareUnicodeCharacter{00A9}{\textcopyright}
347 <all, ts1, ly1, utf8>\DeclareUnicodeCharacter{00AA}{\textordfeminine}
348 <*all, x2, t2c, t2b, t2a, t1, ot2, ly1, lcy>
349 %\DeclareUnicodeCharacter{00AB}{\guillemotleft} % wrong Adobe name
350 \DeclareUnicodeCharacter{00AB}{\guillemetleft}
351 </all, x2, t2c, t2b, t2a, t1, ot2, ly1, lcy>
352 <all, ts1>\DeclareUnicodeCharacter{00AC}{\textlnot}
353 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00AD}{\textasciitilde}
354 <all, ts1, ly1, utf8>\DeclareUnicodeCharacter{00AE}{\textregistered}
355 <all, ts1>\DeclareUnicodeCharacter{00AF}{\textasciimacron}
356 <all, ts1, ly1>\DeclareUnicodeCharacter{00B0}{\textdegree}
357 <all, ts1>\DeclareUnicodeCharacter{00B1}{\textpm}
358 <all, ts1>\DeclareUnicodeCharacter{00B2}{\texttwosuperior}
359 <all, ts1>\DeclareUnicodeCharacter{00B3}{\textthreesuperior}
360 <all, ts1>\DeclareUnicodeCharacter{00B4}{\textasciicircum}
361 <all, ts1, ly1>\DeclareUnicodeCharacter{00B5}{\textmu} % micro sign
362 <all, ts1, oms, ly1>\DeclareUnicodeCharacter{00B6}{\textparagraph}
363 <all, oms, ts1, ly1>\DeclareUnicodeCharacter{00B7}{\textperiodcentered}
364 <all, ot1>\DeclareUnicodeCharacter{00B8}{\c}
365 <all, ts1>\DeclareUnicodeCharacter{00B9}{\textonesuperior}
366 <all, ts1, ly1, utf8>\DeclareUnicodeCharacter{00BA}{\textordmasculine}
367 <*all, x2, t2c, t2b, t2a, t1, ot2, ly1, lcy>
368 %\DeclareUnicodeCharacter{00BB}{\guillemotright} % wrong Adobe name
369 \DeclareUnicodeCharacter{00BB}{\guillemetright}
370 </all, x2, t2c, t2b, t2a, t1, ot2, ly1, lcy>
371 <all, ts1, ly1>\DeclareUnicodeCharacter{00BC}{\textonequarter}
372 <all, ts1, ly1>\DeclareUnicodeCharacter{00BD}{\textonehalf}
373 <all, ts1, ly1>\DeclareUnicodeCharacter{00BE}{\textthreequarters}
374 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00BF}{\textquestiondown}
375 <all, t1, ly1>\DeclareUnicodeCharacter{00C0}{\@tabacckludge'A}
376 <all, t1, ly1>\DeclareUnicodeCharacter{00C1}{\@tabacckludge'A}
377 <all, t1, ly1>\DeclareUnicodeCharacter{00C2}{\^A}
378 <all, t1, ly1>\DeclareUnicodeCharacter{00C3}{\~A}

```



```

379 <all, t1, ly1>\DeclareUnicodeCharacter{00C4}{\~A}
380 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00C5}{\r A}
381 <all, t1, ot1, ly1, lcy>\DeclareUnicodeCharacter{00C6}{\AE}
382 <all, t1, ly1>\DeclareUnicodeCharacter{00C7}{\c C}
383 <all, t1, ly1>\DeclareUnicodeCharacter{00C8}{\@tabacckludge'E}
384 <all, t1, ly1>\DeclareUnicodeCharacter{00C9}{\@tabacckludge'E}
385 <all, t1, ly1>\DeclareUnicodeCharacter{00CA}{\^E}
386 <all, t1, ly1>\DeclareUnicodeCharacter{00CB}{\~E}
387 <all, t1, ly1>\DeclareUnicodeCharacter{00CC}{\@tabacckludge'I}
388 <all, t1, ly1>\DeclareUnicodeCharacter{00CD}{\@tabacckludge'I}
389 <all, t1, ly1>\DeclareUnicodeCharacter{00CE}{\^I}
390 <all, t1, ly1>\DeclareUnicodeCharacter{00CF}{\~I}
391 <all, t1, ly1>\DeclareUnicodeCharacter{00D0}{\DH}
392 <all, t1, ly1>\DeclareUnicodeCharacter{00D1}{\~N}
393 <all, t1, ly1>\DeclareUnicodeCharacter{00D2}{\@tabacckludge'O}
394 <all, t1, ly1>\DeclareUnicodeCharacter{00D3}{\@tabacckludge'O}
395 <all, t1, ly1>\DeclareUnicodeCharacter{00D4}{\~O}
396 <all, t1, ly1>\DeclareUnicodeCharacter{00D5}{\~O}
397 <all, t1, ly1>\DeclareUnicodeCharacter{00D6}{\~O}
398 <all, ts1>\DeclareUnicodeCharacter{00D7}{\texttimes}
399 <all, t1, ot1, ly1, lcy>\DeclareUnicodeCharacter{00D8}{\O}
400 <all, t1, ly1>\DeclareUnicodeCharacter{00D9}{\@tabacckludge'U}
401 <all, t1, ly1>\DeclareUnicodeCharacter{00DA}{\@tabacckludge'U}
402 <all, t1, ly1>\DeclareUnicodeCharacter{00DB}{\~U}
403 <all, t1, ly1>\DeclareUnicodeCharacter{00DC}{\~U}
404 <all, t1, ly1>\DeclareUnicodeCharacter{00DD}{\@tabacckludge'Y}
405 <all, t1, ly1>\DeclareUnicodeCharacter{00DE}{\TH}
406 <all, t1, ot1, ly1, lcy>\DeclareUnicodeCharacter{00DF}{\ss}
407 <all, t1, ly1>\DeclareUnicodeCharacter{00E0}{\@tabacckludge'a}
408 <all, t1, ly1>\DeclareUnicodeCharacter{00E1}{\@tabacckludge'a}
409 <all, t1, ly1>\DeclareUnicodeCharacter{00E2}{\~a}
410 <all, t1, ly1>\DeclareUnicodeCharacter{00E3}{\~a}
411 <all, t1, ly1>\DeclareUnicodeCharacter{00E4}{\~a}
412 <all, t1, ly1>\DeclareUnicodeCharacter{00E5}{\r a}
413 <all, t1, ot1, ly1, lcy>\DeclareUnicodeCharacter{00E6}{\ae}
414 <all, t1, ly1>\DeclareUnicodeCharacter{00E7}{\c c}
415 <all, t1, ly1>\DeclareUnicodeCharacter{00E8}{\@tabacckludge'e}
416 <all, t1, ly1>\DeclareUnicodeCharacter{00E9}{\@tabacckludge'e}
417 <all, t1, ly1>\DeclareUnicodeCharacter{00EA}{\~e}
418 <all, t1, ly1>\DeclareUnicodeCharacter{00EB}{\~e}
419 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00EC}{\@tabacckludge'i}
420 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00ED}{\@tabacckludge'i}
421 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00EE}{\~i}
422 <all, t1, ot1, ly1>\DeclareUnicodeCharacter{00EF}{\~i}
423 <all, t1, ly1>\DeclareUnicodeCharacter{00F0}{\dh}
424 <all, t1, ly1>\DeclareUnicodeCharacter{00F1}{\~n}
425 <all, t1, ly1>\DeclareUnicodeCharacter{00F2}{\@tabacckludge'o}
426 <all, t1, ly1>\DeclareUnicodeCharacter{00F3}{\@tabacckludge'o}
427 <all, t1, ly1>\DeclareUnicodeCharacter{00F4}{\~o}
428 <all, t1, ly1>\DeclareUnicodeCharacter{00F5}{\~o}
429 <all, t1, ly1>\DeclareUnicodeCharacter{00F6}{\~o}
430 <all, ts1>\DeclareUnicodeCharacter{00F7}{\textdiv}
431 <all, t1, ot1, ly1, lcy>\DeclareUnicodeCharacter{00F8}{\o}
432 <all, t1, ly1>\DeclareUnicodeCharacter{00F9}{\@tabacckludge'u}

```

```

433 <all, t1, ly1>\DeclareUnicodeCharacter{00FA}{\@tabacckludge'u}
434 <all, t1, ly1>\DeclareUnicodeCharacter{00FB}{\~u}
435 <all, t1, ly1>\DeclareUnicodeCharacter{00FC}{\"u}
436 <all, t1, ly1>\DeclareUnicodeCharacter{00FD}{\@tabacckludge'y}
437 <all, t1, ly1>\DeclareUnicodeCharacter{00FE}{\th}
438 <all, t1, ly1>\DeclareUnicodeCharacter{00FF}{\"y}
439 <all, t1>\DeclareUnicodeCharacter{0100}{\@tabacckludge=A}
440 <all, t1>\DeclareUnicodeCharacter{0101}{\@tabacckludge=a}
441 <all, t1>\DeclareUnicodeCharacter{0102}{\u A}
442 <all, t1>\DeclareUnicodeCharacter{0103}{\u a}
443 <all, t1>\DeclareUnicodeCharacter{0104}{\k A}
444 <all, t1>\DeclareUnicodeCharacter{0105}{\k a}
445 <all, t1>\DeclareUnicodeCharacter{0106}{\@tabacckludge'C}
446 <all, t1>\DeclareUnicodeCharacter{0107}{\@tabacckludge'c}
447 <all, t1>\DeclareUnicodeCharacter{0108}{\^C}
448 <all, t1>\DeclareUnicodeCharacter{0109}{\^c}
449 <all, t1>\DeclareUnicodeCharacter{010A}{\ .C}
450 <all, t1>\DeclareUnicodeCharacter{010B}{\ .c}
451 <all, t1>\DeclareUnicodeCharacter{010C}{\v C}
452 <all, t1>\DeclareUnicodeCharacter{010D}{\v c}
453 <all, t1>\DeclareUnicodeCharacter{010E}{\v D}
454 <all, t1>\DeclareUnicodeCharacter{010F}{\v d}
455 <all, t1>\DeclareUnicodeCharacter{0110}{\DJ}
456 <all, t1>\DeclareUnicodeCharacter{0111}{\dj}
457 <all, t1>\DeclareUnicodeCharacter{0112}{\@tabacckludge=E}
458 <all, t1>\DeclareUnicodeCharacter{0113}{\@tabacckludge=e}
459 <all, t1>\DeclareUnicodeCharacter{0114}{\u E}
460 <all, t1>\DeclareUnicodeCharacter{0115}{\u e}
461 <all, t1>\DeclareUnicodeCharacter{0116}{\ .E}
462 <all, t1>\DeclareUnicodeCharacter{0117}{\ .e}
463 <all, t1>\DeclareUnicodeCharacter{0118}{\k E}
464 <all, t1>\DeclareUnicodeCharacter{0119}{\k e}
465 <all, t1>\DeclareUnicodeCharacter{011A}{\v E}
466 <all, t1>\DeclareUnicodeCharacter{011B}{\v e}
467 <all, t1>\DeclareUnicodeCharacter{011C}{\^G}
468 <all, t1>\DeclareUnicodeCharacter{011D}{\^g}
469 <all, t1>\DeclareUnicodeCharacter{011E}{\u G}
470 <all, t1>\DeclareUnicodeCharacter{011F}{\u g}
471 <all, t1>\DeclareUnicodeCharacter{0120}{\ .G}
472 <all, t1>\DeclareUnicodeCharacter{0121}{\ .g}
473 <all, t1>\DeclareUnicodeCharacter{0122}{\c G}
474 <all, t1>\DeclareUnicodeCharacter{0123}{\c g}
475 <all, t1>\DeclareUnicodeCharacter{0124}{\^H}
476 <all, t1>\DeclareUnicodeCharacter{0125}{\^h}
477 <all, t1>\DeclareUnicodeCharacter{0128}{\^I}
478 <all, t1>\DeclareUnicodeCharacter{0129}{\^~i}
479 <all, t1>\DeclareUnicodeCharacter{012A}{\@tabacckludge=I}
480 <all, t1>\DeclareUnicodeCharacter{012B}{\@tabacckludge=i}
481 <all, t1>\DeclareUnicodeCharacter{012C}{\u I}
482 <all, t1>\DeclareUnicodeCharacter{012D}{\u i}
483 <all, t1>\DeclareUnicodeCharacter{012E}{\k I}

484 <all, t1>\DeclareUnicodeCharacter{012F}{\k i}
485 <all, t1>\DeclareUnicodeCharacter{0130}{\ .I}
486 <all, t2c, t2b, t2a, t1, ot2, ot1, ly1, lcy>\DeclareUnicodeCharacter{0131}{\i}

```

```

487 <all,t1>\DeclareUnicodeCharacter{0132}{\IJ}
488 <all,t1>\DeclareUnicodeCharacter{0133}{\ij}
489 <all,t1>\DeclareUnicodeCharacter{0134}{\^J}
490 <all,t1>\DeclareUnicodeCharacter{0135}{\^j}
491 <all,t1>\DeclareUnicodeCharacter{0136}{\c K}
492 <all,t1>\DeclareUnicodeCharacter{0137}{\c k}
493 <all,t1>\DeclareUnicodeCharacter{0139}{\@tabacckludge'L}
494 <all,t1>\DeclareUnicodeCharacter{013A}{\@tabacckludge'l}
495 <all,t1>\DeclareUnicodeCharacter{013B}{\c L}
496 <all,t1>\DeclareUnicodeCharacter{013C}{\c l}
497 <all,t1>\DeclareUnicodeCharacter{013D}{\v L}
498 <all,t1>\DeclareUnicodeCharacter{013E}{\v l}
499 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0141}{\L}
500 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0142}{\l}
501 <all,t1>\DeclareUnicodeCharacter{0143}{\@tabacckludge'N}
502 <all,t1>\DeclareUnicodeCharacter{0144}{\@tabacckludge'n}
503 <all,t1>\DeclareUnicodeCharacter{0145}{\c N}
504 <all,t1>\DeclareUnicodeCharacter{0146}{\c n}
505 <all,t1>\DeclareUnicodeCharacter{0147}{\v N}
506 <all,t1>\DeclareUnicodeCharacter{0148}{\v n}
507 <all,t1>\DeclareUnicodeCharacter{014A}{\NG}
508 <all,t1>\DeclareUnicodeCharacter{014B}{\ng}
509 <all,t1>\DeclareUnicodeCharacter{014C}{\@tabacckludge=0}
510 <all,t1>\DeclareUnicodeCharacter{014D}{\@tabacckludge=o}
511 <all,t1>\DeclareUnicodeCharacter{014E}{\u O}
512 <all,t1>\DeclareUnicodeCharacter{014F}{\u o}
513 <all,t1>\DeclareUnicodeCharacter{0150}{\H O}
514 <all,t1>\DeclareUnicodeCharacter{0151}{\H o}
515 <all,t1,ot1,ly1,lcy>\DeclareUnicodeCharacter{0152}{\OE}
516 <all,t1,ot1,ly1,lcy>\DeclareUnicodeCharacter{0153}{\oe}
517 <all,t1>\DeclareUnicodeCharacter{0154}{\@tabacckludge'R}
518 <all,t1>\DeclareUnicodeCharacter{0155}{\@tabacckludge'r}
519 <all,t1>\DeclareUnicodeCharacter{0156}{\c R}
520 <all,t1>\DeclareUnicodeCharacter{0157}{\c r}
521 <all,t1>\DeclareUnicodeCharacter{0158}{\v R}
522 <all,t1>\DeclareUnicodeCharacter{0159}{\v r}
523 <all,t1>\DeclareUnicodeCharacter{015A}{\@tabacckludge'S}
524 <all,t1>\DeclareUnicodeCharacter{015B}{\@tabacckludge's}
525 <all,t1>\DeclareUnicodeCharacter{015C}{\^S}
526 <all,t1>\DeclareUnicodeCharacter{015D}{\^s}
527 <all,t1>\DeclareUnicodeCharacter{015E}{\c S}
528 <all,t1>\DeclareUnicodeCharacter{015F}{\c s}
529 <all,t1,ly1>\DeclareUnicodeCharacter{0160}{\v S}
530 <all,t1,ly1>\DeclareUnicodeCharacter{0161}{\v s}
531 <all,t1>\DeclareUnicodeCharacter{0162}{\c T}
532 <all,t1>\DeclareUnicodeCharacter{0163}{\c t}
533 <all,t1>\DeclareUnicodeCharacter{0164}{\v T}
534 <all,t1>\DeclareUnicodeCharacter{0165}{\v t}
535 <all,t1>\DeclareUnicodeCharacter{0168}{\^U}
536 <all,t1>\DeclareUnicodeCharacter{0169}{\^u}
537 <all,t1>\DeclareUnicodeCharacter{016A}{\@tabacckludge=U}
538 <all,t1>\DeclareUnicodeCharacter{016B}{\@tabacckludge=u}
539 <all,t1>\DeclareUnicodeCharacter{016C}{\u U}
540 <all,t1>\DeclareUnicodeCharacter{016D}{\u u}

```

```

541 <all,t1>\DeclareUnicodeCharacter{016E}{\r U}
542 <all,t1>\DeclareUnicodeCharacter{016F}{\r u}
543 <all,t1>\DeclareUnicodeCharacter{0170}{\H U}
544 <all,t1>\DeclareUnicodeCharacter{0171}{\H u}
545 <all,t1>\DeclareUnicodeCharacter{0172}{\k U}
546 <all,t1>\DeclareUnicodeCharacter{0173}{\k u}

547 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0174}{\^W}
548 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0175}{\^w}
549 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0176}{\^Y}
550 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0177}{\^y}
551 <all,t1,ly1>\DeclareUnicodeCharacter{0178}{\"Y}
552 <all,t1>\DeclareUnicodeCharacter{0179}{\@tabacckludge'Z}
553 <all,t1>\DeclareUnicodeCharacter{017A}{\@tabacckludge'z}
554 <all,t1>\DeclareUnicodeCharacter{017B}{\Z}
555 <all,t1>\DeclareUnicodeCharacter{017C}{\z}
556 <all,t1,ly1>\DeclareUnicodeCharacter{017D}{\v Z}
557 <all,t1,ly1>\DeclareUnicodeCharacter{017E}{\v z}
558 <all,ts1,ly1>\DeclareUnicodeCharacter{0192}{\textflorin}

559 <all,t1>\DeclareUnicodeCharacter{01C4}{D\ v Z}
560 <all,t1>\DeclareUnicodeCharacter{01C5}{D\ v z}
561 <all,t1>\DeclareUnicodeCharacter{01C6}{d\ v z}
562 <all,t1>\DeclareUnicodeCharacter{01C7}{LJ}
563 <all,t1>\DeclareUnicodeCharacter{01C8}{Lj}
564 <all,t1>\DeclareUnicodeCharacter{01C9}{lj}
565 <all,t1>\DeclareUnicodeCharacter{01CA}{NJ}
566 <all,t1>\DeclareUnicodeCharacter{01CB}{Nj}
567 <all,t1>\DeclareUnicodeCharacter{01CC}{nj}

568 <all,t1>\DeclareUnicodeCharacter{01CD}{\v A}
569 <all,t1>\DeclareUnicodeCharacter{01CE}{\v a}
570 <all,t1>\DeclareUnicodeCharacter{01CF}{\v I}
571 <all,t1>\DeclareUnicodeCharacter{01D0}{\v \i}
572 <all,t1>\DeclareUnicodeCharacter{01D1}{\v O}
573 <all,t1>\DeclareUnicodeCharacter{01D2}{\v o}
574 <all,t1>\DeclareUnicodeCharacter{01D3}{\v U}
575 <all,t1>\DeclareUnicodeCharacter{01D4}{\v u}
576 <all,t1>\DeclareUnicodeCharacter{01E2}{\@tabacckludge=\AE}
577 <all,t1>\DeclareUnicodeCharacter{01E3}{\@tabacckludge=\ae}
578 <all,t1>\DeclareUnicodeCharacter{01E6}{\v G}
579 <all,t1>\DeclareUnicodeCharacter{01E7}{\v g}
580 <all,t1>\DeclareUnicodeCharacter{01E8}{\v K}
581 <all,t1>\DeclareUnicodeCharacter{01E9}{\v k}
582 <all,t1>\DeclareUnicodeCharacter{01EA}{\k O}
583 <all,t1>\DeclareUnicodeCharacter{01EB}{\k o}
584 <all,t1>\DeclareUnicodeCharacter{01F0}{\v\j}
585 <all,t1>\DeclareUnicodeCharacter{01F4}{\@tabacckludge'G}
586 <all,t1>\DeclareUnicodeCharacter{01F5}{\@tabacckludge'g}

587 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0218}{\textcommabelow S}
588 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{0219}{\textcommabelow s}
589 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{021A}{\textcommabelow T}
590 <all,t1,ot1,ly1>\DeclareUnicodeCharacter{021B}{\textcommabelow t}

591 <all,t1>\DeclareUnicodeCharacter{0232}{\@tabacckludge=Y}
592 <all,t1>\DeclareUnicodeCharacter{0233}{\@tabacckludge=y}

```

```

593 <all, t2c, t2b, t2a, t1, ot2, ot1, ly1, lcy>\DeclareUnicodeCharacter{0237}{\j}
594 <all, ly1, utf8>\DeclareUnicodeCharacter{02C6}{\textasciicircum}
595 <all, ts1>\DeclareUnicodeCharacter{02C7}{\textasciicaron}
596 <all, ly1, utf8>\DeclareUnicodeCharacter{02DC}{\textasciitilde}
597 <all, ts1>\DeclareUnicodeCharacter{02D8}{\textasciibreve}
598 <all, t1>\DeclareUnicodeCharacter{02D9}{\.\{}}
599 <all, t1>\DeclareUnicodeCharacter{02DB}{\k{}}
600 <all, ts1>\DeclareUnicodeCharacter{02DD}{\textacutedbl}

```

The Cyrillic code points have been recently checked (2007) and extended and corrected by Matthias Noe (a9931078@unet.univie.ac.at) — thanks.

```

601 <*all, x2, t2c, t2b, t2a, ot2, lcy>
602 \DeclareUnicodeCharacter{0400}{\@tabacckludge'\CYRE}
603 </all, x2, t2c, t2b, t2a, ot2, lcy>
604 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{0401}{\CYRYO}
605 <all, x2, t2a, ot2>\DeclareUnicodeCharacter{0402}{\CYRDJE}
606 <*all, x2, t2c, t2b, t2a, ot2, lcy>
607 \DeclareUnicodeCharacter{0403}{\@tabacckludge'\CYRG}
608 </all, x2, t2c, t2b, t2a, ot2, lcy>
609 <all, x2, t2a, ot2, lcy>\DeclareUnicodeCharacter{0404}{\CYRIE}
610 <all, x2, t2c, t2b, t2a, ot2>\DeclareUnicodeCharacter{0405}{\CYRDZE}
611 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{0406}{\CYRII}
612 <all, x2, t2a, lcy>\DeclareUnicodeCharacter{0407}{\CYRYI}
613 <all, x2, t2c, t2b, t2a, ot2>\DeclareUnicodeCharacter{0408}{\CYRJE}
614 <all, x2, t2b, t2a, ot2>\DeclareUnicodeCharacter{0409}{\CYRLJE}
615 <all, x2, t2b, t2a, ot2>\DeclareUnicodeCharacter{040A}{\CYRNJE}
616 <all, x2, t2a, ot2>\DeclareUnicodeCharacter{040B}{\CYRTSHE}
617 <*all, x2, t2c, t2b, t2a, ot2, lcy>
618 \DeclareUnicodeCharacter{040C}{\@tabacckludge'\CYRK}
619 \DeclareUnicodeCharacter{040D}{\@tabacckludge'\CYRI}
620 </all, x2, t2c, t2b, t2a, ot2, lcy>
621 <all, x2, t2b, t2a, lcy>\DeclareUnicodeCharacter{040E}{\CYRUSHRT}
622 <all, x2, t2c, t2a, ot2>\DeclareUnicodeCharacter{040F}{\CYRDZHE}
623 <*all, x2, t2c, t2b, t2a, ot2, lcy>
624 \DeclareUnicodeCharacter{0410}{\CYRA}
625 \DeclareUnicodeCharacter{0411}{\CYRB}
626 \DeclareUnicodeCharacter{0412}{\CYRV}
627 \DeclareUnicodeCharacter{0413}{\CYRG}
628 \DeclareUnicodeCharacter{0414}{\CYRD}
629 \DeclareUnicodeCharacter{0415}{\CYRE}
630 \DeclareUnicodeCharacter{0416}{\CYRZH}
631 \DeclareUnicodeCharacter{0417}{\CYRZ}
632 \DeclareUnicodeCharacter{0418}{\CYRI}
633 \DeclareUnicodeCharacter{0419}{\CYRISHRT}
634 \DeclareUnicodeCharacter{041A}{\CYRK}
635 \DeclareUnicodeCharacter{041B}{\CYRL}
636 \DeclareUnicodeCharacter{041C}{\CYRM}
637 \DeclareUnicodeCharacter{041D}{\CYRN}
638 \DeclareUnicodeCharacter{041E}{\CYRO}
639 \DeclareUnicodeCharacter{041F}{\CYRP}
640 \DeclareUnicodeCharacter{0420}{\CYRR}
641 \DeclareUnicodeCharacter{0421}{\CYRS}
642 \DeclareUnicodeCharacter{0422}{\CYRT}
643 \DeclareUnicodeCharacter{0423}{\CYRU}

```

```

644 \DeclareUnicodeCharacter{0424}{\CYRF}
645 \DeclareUnicodeCharacter{0425}{\CYRH}
646 \DeclareUnicodeCharacter{0426}{\CYRC}
647 \DeclareUnicodeCharacter{0427}{\CYRCH}
648 \DeclareUnicodeCharacter{0428}{\CYRSH}
649 \DeclareUnicodeCharacter{0429}{\CYRSHCH}
650 \DeclareUnicodeCharacter{042A}{\CYRHRDSN}
651 \DeclareUnicodeCharacter{042B}{\CYRERY}
652 \DeclareUnicodeCharacter{042C}{\CYRSFTSN}
653 \DeclareUnicodeCharacter{042D}{\CYREREV}
654 \DeclareUnicodeCharacter{042E}{\CYRYU}
655 \DeclareUnicodeCharacter{042F}{\CYRYA}
656 \DeclareUnicodeCharacter{0430}{\cyra}
657 \DeclareUnicodeCharacter{0431}{\cyrb}
658 \DeclareUnicodeCharacter{0432}{\cyrv}
659 \DeclareUnicodeCharacter{0433}{\cyrg}
660 \DeclareUnicodeCharacter{0434}{\cyrd}
661 \DeclareUnicodeCharacter{0435}{\cyre}
662 \DeclareUnicodeCharacter{0436}{\cyrzh}
663 \DeclareUnicodeCharacter{0437}{\cyrz}
664 \DeclareUnicodeCharacter{0438}{\cyri}
665 \DeclareUnicodeCharacter{0439}{\cyrishrt}
666 \DeclareUnicodeCharacter{043A}{\cyrk}
667 \DeclareUnicodeCharacter{043B}{\cyr1}
668 \DeclareUnicodeCharacter{043C}{\cyrm}
669 \DeclareUnicodeCharacter{043D}{\cyrn}
670 \DeclareUnicodeCharacter{043E}{\cyro}
671 \DeclareUnicodeCharacter{043F}{\cyrp}
672 \DeclareUnicodeCharacter{0440}{\cyrr}
673 \DeclareUnicodeCharacter{0441}{\cyrs}
674 \DeclareUnicodeCharacter{0442}{\cyrt}
675 \DeclareUnicodeCharacter{0443}{\cyrU}
676 \DeclareUnicodeCharacter{0444}{\cyrf}
677 \DeclareUnicodeCharacter{0445}{\cyrh}
678 \DeclareUnicodeCharacter{0446}{\cyrc}
679 \DeclareUnicodeCharacter{0447}{\cyrch}
680 \DeclareUnicodeCharacter{0448}{\cyrsh}
681 \DeclareUnicodeCharacter{0449}{\cyrshch}
682 \DeclareUnicodeCharacter{044A}{\cyrhrdsn}
683 \DeclareUnicodeCharacter{044B}{\cyrery}
684 \DeclareUnicodeCharacter{044C}{\cyrstsn}
685 \DeclareUnicodeCharacter{044D}{\cyrrrev}
686 \DeclareUnicodeCharacter{044E}{\cyrU}
687 \DeclareUnicodeCharacter{044F}{\cyrya}
688 \DeclareUnicodeCharacter{0450}{\@tabacckludge'\cyre}
689 \DeclareUnicodeCharacter{0451}{\cyryo}
690 </all, x2, t2c, t2b, t2a, ot2, lcy>
691 <all, x2, t2a, ot2>\DeclareUnicodeCharacter{0452}{\cyrdje}
692 <*all, x2, t2c, t2b, t2a, ot2, lcy>
693 \DeclareUnicodeCharacter{0453}{\@tabacckludge'\cyrg}
694 </all, x2, t2c, t2b, t2a, ot2, lcy>
695 <all, x2, t2a, ot2, lcy>\DeclareUnicodeCharacter{0454}{\cyrie}
696 <all, x2, t2c, t2b, t2a, ot2>\DeclareUnicodeCharacter{0455}{\cyrdze}
697 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{0456}{\cyrii}

```

```

698 <all, x2, t2a, lcy>\DeclareUnicodeCharacter{0457}{\cyryi}
699 <all, x2, t2c, t2b, t2a, ot2>\DeclareUnicodeCharacter{0458}{\cyrje}
700 <all, x2, t2b, t2a, ot2>\DeclareUnicodeCharacter{0459}{\cyrlje}
701 <all, x2, t2b, t2a, ot2>\DeclareUnicodeCharacter{045A}{\cyrnje}
702 <all, x2, t2a, ot2>\DeclareUnicodeCharacter{045B}{\cyrtshe}
703 <*all, x2, t2c, t2b, t2a, ot2, lcy>
704 \DeclareUnicodeCharacter{045C}{\@tabacckludge'\cyrk}
705 \DeclareUnicodeCharacter{045D}{\@tabacckludge'\cyri}
706 </all, x2, t2c, t2b, t2a, ot2, lcy>
707 <all, x2, t2b, t2a, lcy>\DeclareUnicodeCharacter{045E}{\cyrushrt}
708 <all, x2, t2c, t2a, ot2>\DeclareUnicodeCharacter{045F}{\cyrzhe}
709 <all, x2, ot2>\DeclareUnicodeCharacter{0462}{\CYRYAT}
710 <all, x2, ot2>\DeclareUnicodeCharacter{0463}{\cyryat}
711 <all, x2>\DeclareUnicodeCharacter{046A}{\CYRBYUS}
712 <all, x2>\DeclareUnicodeCharacter{046B}{\cyrbyus}

```

The next two declarations are questionable, the encoding definition should probably contain \CYROTLD and \cyrotld. Or alternatively, if the characters in the X2 encodings are really meant to represent the historical characters in Ux0472 and Ux0473 (they look like them) then they would need to change instead.

However, their looks are probably a font designers decision and the next two mappings are wrong or rather the names in OT2 should change for consistency.

On the other hand the names \CYROTLD are somewhat questionable as the Unicode standard only describes “Cyrillic barred O” while TLD refers to a tilde (which is more less what the “Cyrillic FITA looks according to the Unicode book).

```

713 <all, ot2>\DeclareUnicodeCharacter{0472}{\CYRFITA}
714 <all, ot2>\DeclareUnicodeCharacter{0473}{\cyrfita}

715 <all, x2, ot2>\DeclareUnicodeCharacter{0474}{\CYRIZH}
716 <all, x2, ot2>\DeclareUnicodeCharacter{0475}{\cyrizh}

```

While the double grave accent seems to exist in X2, T2A, T2B and T2C encoding, the letter izhitsa exists only in X2 and OT2. Therefore, izhitsa with double grave seems to be possible only using X2.

```

717 <all, x2>\DeclareUnicodeCharacter{0476}{\C\CYRIZH}
718 <all, x2>\DeclareUnicodeCharacter{0477}{\C\Cyrizh}

719 <all, t2c>\DeclareUnicodeCharacter{048C}{\CYRSEMISFTSN}
720 <all, t2c>\DeclareUnicodeCharacter{048D}{\cyrsemisftsn}
721 <all, t2c>\DeclareUnicodeCharacter{048E}{\CYRRTICK}
722 <all, t2c>\DeclareUnicodeCharacter{048F}{\cyrrtick}
723 <all, x2, t2a, lcy>\DeclareUnicodeCharacter{0490}{\CYRGUP}
724 <all, x2, t2a, lcy>\DeclareUnicodeCharacter{0491}{\cyrgup}
725 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{0492}{\CYRGHCRS}
726 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{0493}{\cyrghcrs}
727 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{0494}{\CYRGHK}
728 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{0495}{\cyrghk}
729 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{0496}{\CYRZHDSC}
730 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{0497}{\cyrzhdsc}
731 <all, x2, t2a>\DeclareUnicodeCharacter{0498}{\CYRZDSC}
732 <all, x2, t2a>\DeclareUnicodeCharacter{0499}{\cyrzdsc}
733 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{049A}{\CYRKDSC}
734 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{049B}{\cyrkdsc}
735 <all, x2, t2a>\DeclareUnicodeCharacter{049C}{\CYRKVCRS}
736 <all, x2, t2a>\DeclareUnicodeCharacter{049D}{\cyrkvcrs}

```

```

737 <all, x2, t2c>\DeclareUnicodeCharacter{049E}{\CYRKHCRS}
738 <all, x2, t2c>\DeclareUnicodeCharacter{049F}{\cyrkhrs}
739 <all, x2, t2a>\DeclareUnicodeCharacter{04A0}{\CYRKBEAK}
740 <all, x2, t2a>\DeclareUnicodeCharacter{04A1}{\cyrkbeak}
741 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04A2}{\CYRNDSC}
742 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04A3}{\cyrndsc}
743 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{04A4}{\CYRNG}
744 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{04A5}{\cyrng}
745 <all, x2, t2c>\DeclareUnicodeCharacter{04A6}{\CYRPHK}
746 <all, x2, t2c>\DeclareUnicodeCharacter{04A7}{\cyrphk}
747 <all, x2, t2c>\DeclareUnicodeCharacter{04A8}{\CYRABHHA}
748 <all, x2, t2c>\DeclareUnicodeCharacter{04A9}{\cyrabhha}
749 <all, x2, t2a>\DeclareUnicodeCharacter{04AA}{\CYRSDSC}
750 <all, x2, t2a>\DeclareUnicodeCharacter{04AB}{\cyrsdsc}
751 <all, x2, t2c>\DeclareUnicodeCharacter{04AC}{\CYRTDSC}
752 <all, x2, t2c>\DeclareUnicodeCharacter{04AD}{\cyrtdsc}
753 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{04AE}{\CYRY}
754 <all, x2, t2b, t2a>\DeclareUnicodeCharacter{04AF}{\cyry}
755 <all, x2, t2a>\DeclareUnicodeCharacter{04B0}{\CYRYHCRS}
756 <all, x2, t2a>\DeclareUnicodeCharacter{04B1}{\cyrhcrs}
757 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04B2}{\CYRHDSC}
758 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04B3}{\cyrhdsc}
759 <all, x2, t2c>\DeclareUnicodeCharacter{04B4}{\CYRTETSE}
760 <all, x2, t2c>\DeclareUnicodeCharacter{04B5}{\cyrtetse}
761 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04B6}{\CYRCHRDSC}
762 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04B7}{\cyrchrdsc}
763 <all, x2, t2a>\DeclareUnicodeCharacter{04B8}{\CYRCHVCRS}
764 <all, x2, t2a>\DeclareUnicodeCharacter{04B9}{\cyrchvcrs}
765 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04BA}{\CYRSHHA}
766 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04BB}{\cyrshha}
767 <all, x2, t2c>\DeclareUnicodeCharacter{04BC}{\CYRABHCH}
768 <all, x2, t2c>\DeclareUnicodeCharacter{04BD}{\cyrabhch}
769 <all, x2, t2c>\DeclareUnicodeCharacter{04BE}{\CYRABHCHDSC}
770 <all, x2, t2c>\DeclareUnicodeCharacter{04BF}{\cyrabhchdsc}

```

The character \CYRpalochka is not defined by OT2 and LCY. However it is looking identical to \CYRII and the Unicode standard explicitly refers to that (and to Latin I). So perhaps those encodings could get an alias? On the other hand, why are there two distinct slots in the T2 encodings even though they are so pressed for space? Perhaps they don't always look alike.

```

771 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04C0}{\CYRpalochka}
772 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04C1}{\U\CYRZH}
773 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04C2}{\U\cyrzh}
774 <all, x2, t2b>\DeclareUnicodeCharacter{04C3}{\CYRKHK}
775 <all, x2, t2b>\DeclareUnicodeCharacter{04C4}{\cyrkhk}

```

According to the Unicode standard Ux04C5 should be an L with “tail” not with descender (which also exists as Ux04A2) but it looks as if the char names do not make this distinction. Should they?

```

776 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{04C5}{\CYRLDSC}
777 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{04C6}{\cyrldsc}

778 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{04C7}{\CYRNHK}
779 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{04C8}{\cyrnhk}

```



```

780 <all, x2, t2b>\DeclareUnicodeCharacter{04CB}{\CYRCHLDSC}
781 <all, x2, t2b>\DeclareUnicodeCharacter{04CC}{\cyrchldsc}

```

According to the Unicode standard Ux04CD should be an M with “tail” not with descender. However this time there is no M with descender in the Unicode standard.

```

782 <all, x2, t2c>\DeclareUnicodeCharacter{04CD}{\CYRMDSC}
783 <all, x2, t2c>\DeclareUnicodeCharacter{04CE}{\cyrmdsc}

784 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04D0}{\U\CYRA}
785 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04D1}{\U\cyra}
786 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04D2}{\"\CYRA}
787 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04D3}{\"\cyra}
788 <all, x2, t2a>\DeclareUnicodeCharacter{04D4}{\CYRAE}
789 <all, x2, t2a>\DeclareUnicodeCharacter{04D5}{\cyrae}
790 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04D6}{\U\CYRE}
791 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04D7}{\U\cyre}
792 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04D8}{\CYRSCHWA}
793 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04D9}{\cyrswa}
794 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04DA}{\"\CYRSCHWA}
795 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04DB}{\"\cyrswa}
796 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04DC}{\"\CYRZH}
797 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04DD}{\"\cyrzh}
798 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04DE}{\"\CYRZ}
799 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04DF}{\"\cyrz}
800 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{04E0}{\CYRABHDZE}
801 <all, x2, t2c, t2b>\DeclareUnicodeCharacter{04E1}{\cyrabhdze}
802 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04E2}{\@tabacckludge=\CYRI}
803 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04E3}{\@tabacckludge=\cyri}
804 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04E4}{\"\CYRI}
805 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04E5}{\"\cyri}
806 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04E6}{\"\CYRO}
807 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04E7}{\"\cyro}
808 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04E8}{\CYROTLD}
809 <all, x2, t2c, t2b, t2a>\DeclareUnicodeCharacter{04E9}{\cyrotld}
810 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04EC}{\"\CYREREV}
811 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04ED}{\"\cyrerev}
812 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04EE}{\@tabacckludge=\CYRU}
813 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04EF}{\@tabacckludge=\cyru}
814 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F0}{\"\CYRU}
815 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F1}{\"\cyru}
816 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F2}{\H\CYRU}
817 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F3}{\H\cyru}
818 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F4}{\"\CYRCH}
819 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F5}{\"\cyrch}
820 <all, x2, t2b>\DeclareUnicodeCharacter{04F6}{\CYRGDSC}
821 <all, x2, t2b>\DeclareUnicodeCharacter{04F7}{\cyrghsc}
822 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F8}{\"\CYRERY}
823 <all, x2, t2c, t2b, t2a, ot2, lcy>\DeclareUnicodeCharacter{04F9}{\"\cyrery}
824 <all, t2b>\DeclareUnicodeCharacter{04FA}{\CYRGDSCHCRS}
825 <all, t2b>\DeclareUnicodeCharacter{04FB}{\cyrghschcrs}
826 <all, x2, t2b>\DeclareUnicodeCharacter{04FC}{\CYRHHK}
827 <all, x2, t2b>\DeclareUnicodeCharacter{04FD}{\cyrhkh}
828 <all, t2b>\DeclareUnicodeCharacter{04FE}{\CYRHHCRS}
829 <all, t2b>\DeclareUnicodeCharacter{04FF}{\cyrhhcrs}

```

```

830 <all, ts1>\DeclareUnicodeCharacter{0E3F}{\textbaht}
831 <all, t1>\DeclareUnicodeCharacter{1E02}{\ .B}
832 <all, t1>\DeclareUnicodeCharacter{1E03}{\ .b}

833 <all, t1>\DeclareUnicodeCharacter{1E0D}{\d d}
834 <all, t1>\DeclareUnicodeCharacter{1E1E}{\ .F}
835 <all, t1>\DeclareUnicodeCharacter{1E1F}{\ .f}
836 <all, t1>\DeclareUnicodeCharacter{1E25}{\d h}
837 <all, t1>\DeclareUnicodeCharacter{1E30}{\@tabacckludge'K}
838 <all, t1>\DeclareUnicodeCharacter{1E31}{\@tabacckludge'k}
839 <all, t1>\DeclareUnicodeCharacter{1E37}{\d l}
840 <all, t1>\DeclareUnicodeCharacter{1E8E}{\ .Y}
841 <all, t1>\DeclareUnicodeCharacter{1E8F}{\ .y}
842 <all, t1>\DeclareUnicodeCharacter{1E43}{\d m}
843 <all, t1>\DeclareUnicodeCharacter{1E45}{\ .n}
844 <all, t1>\DeclareUnicodeCharacter{1E47}{\d n}
845 <all, t1>\DeclareUnicodeCharacter{1E5B}{\d r}
846 <all, t1>\DeclareUnicodeCharacter{1E63}{\d s}
847 <all, t1>\DeclareUnicodeCharacter{1E6D}{\d t}
848 <all, t1>\DeclareUnicodeCharacter{1E90}{\ ^Z}
849 <all, t1>\DeclareUnicodeCharacter{1E91}{\ ^z}
850 <all, t1>\DeclareUnicodeCharacter{1E9E}{\SS}
851 <all, t1>\DeclareUnicodeCharacter{1EF2}{\@tabacckludge'Y}
852 <all, t1>\DeclareUnicodeCharacter{1EF3}{\@tabacckludge'y}

853 <all, x2, t2c, t2b, t2a, t1, utf8>\DeclareUnicodeCharacter{200C}{\textcompwordmark}

854 <all, t1>\DeclareUnicodeCharacter{2010}{\ -}
855 <all, t1>\DeclareUnicodeCharacter{2011}{\mbox{-}}

```

U+2012 should be the width of a digit, endash is OK in many fonts including cm.

```

856 <all, t1>\DeclareUnicodeCharacter{2012}{\textendash}
857 <*all, x2, t2c, t2b, t2a, t1, ot2, ot1, ly1, lcy>
858 \DeclareUnicodeCharacter{2013}{\textendash}
859 \DeclareUnicodeCharacter{2014}{\textemdash}

```

U+2015 is Horizontal bar

```

860 <all, t1>\DeclareUnicodeCharacter{2015}{\textemdash}
861 </all, x2, t2c, t2b, t2a, t1, ot2, ot1, ly1, lcy>
862 <all, ts1>\DeclareUnicodeCharacter{2016}{\textbardbl}
863 <*all, x2, t2c, t2b, t2a, t1, ot2, ot1, lcy>
864 \DeclareUnicodeCharacter{2018}{\textquoteleft}
865 \DeclareUnicodeCharacter{2019}{\textquoteright}
866 </all, x2, t2c, t2b, t2a, t1, ot2, ot1, lcy>
867 <all, t1>\DeclareUnicodeCharacter{201A}{\quotesinglbase}
868 <*all, x2, t2c, t2b, t2a, t1, ot2, ot1, ly1, lcy>
869 \DeclareUnicodeCharacter{201C}{\textquotedblleft}
870 \DeclareUnicodeCharacter{201D}{\textquotedblright}
871 </all, x2, t2c, t2b, t2a, t1, ot2, ot1, ly1, lcy>
872 <all, x2, t2c, t2b, t2a, t1, lcy>\DeclareUnicodeCharacter{201E}{\quotedblbase}
873 <all, ts1, oms, ly1>\DeclareUnicodeCharacter{2020}{\textdagger}
874 <all, ts1, oms, ly1>\DeclareUnicodeCharacter{2021}{\textdaggerdbl}
875 <all, ts1, oms, ly1>\DeclareUnicodeCharacter{2022}{\textbullet}
876 <all, ly1, utf8>\DeclareUnicodeCharacter{2026}{\textellipsis}
877 <*all, x2, ts1, t2c, t2b, t2a, t1, ly1>
878 \DeclareUnicodeCharacter{2030}{\textperthousand}
879 </all, x2, ts1, t2c, t2b, t2a, t1, ly1>

```

```

880 <all,x2,ts1,t2c,t2b,t2a,t1>
881 \DeclareUnicodeCharacter{2031}{\textpertenthousand}
882 </all,x2,ts1,t2c,t2b,t2a,t1>
883 <all,t1,ly1>\DeclareUnicodeCharacter{2039}{\guilsinglleft}
884 <all,t1,ly1>\DeclareUnicodeCharacter{203A}{\guilsinglright}
885 <all,ts1>\DeclareUnicodeCharacter{203B}{\textreferencemark}
886 <all,ts1>\DeclareUnicodeCharacter{203D}{\textinterrobang}
887 <all,ts1>\DeclareUnicodeCharacter{2044}{\textfractionsolidus}
888 <all,ts1>\DeclareUnicodeCharacter{204E}{\textasteriskcentered}
889 <all,ts1>\DeclareUnicodeCharacter{2052}{\textdiscount}
890 <all,ts1>\DeclareUnicodeCharacter{20A1}{\textcolonmonetary}
891 <all,ts1>\DeclareUnicodeCharacter{20A4}{\textlira}
892 <all,ts1>\DeclareUnicodeCharacter{20A6}{\textnaira}
893 <all,ts1>\DeclareUnicodeCharacter{20A9}{\textwon}
894 <all,ts1>\DeclareUnicodeCharacter{20AB}{\textdong}
895 <all,ts1>\DeclareUnicodeCharacter{20AC}{\texteuro}
896 <all,ts1>\DeclareUnicodeCharacter{20B1}{\textpeso}
897 <all,ts1>\DeclareUnicodeCharacter{2103}{\textcelsius}
898 <all,x2,ts1,t2c,t2b,t2a,ot2,lcyl>\DeclareUnicodeCharacter{2116}{\textnumero}
899 <all,ts1>\DeclareUnicodeCharacter{2117}{\textcircledP}
900 <all,ts1>\DeclareUnicodeCharacter{211E}{\textrecipe}
901 <all,ts1>\DeclareUnicodeCharacter{2120}{\textservicemark}
902 <all,ts1,ly1,utf8>\DeclareUnicodeCharacter{2122}{\texttrademark}
903 <all,ts1>\DeclareUnicodeCharacter{2126}{\textohm}
904 <all,ts1>\DeclareUnicodeCharacter{2127}{\textmho}
905 <all,ts1>\DeclareUnicodeCharacter{212E}{\textestimated}
906 <all,ts1>\DeclareUnicodeCharacter{2190}{\textleftarrow}
907 <all,ts1>\DeclareUnicodeCharacter{2191}{\textuparrow}
908 <all,ts1>\DeclareUnicodeCharacter{2192}{\textrightarrow}
909 <all,ts1>\DeclareUnicodeCharacter{2193}{\textdownarrow}

910 <all,x2,ts1,t2c,t2b,t2a>\DeclareUnicodeCharacter{2329}{\textlangle}
911 <all,x2,ts1,t2c,t2b,t2a>\DeclareUnicodeCharacter{3008}{\textlangle}
912 <all,x2,ts1,t2c,t2b,t2a>\DeclareUnicodeCharacter{232A}{\textrangle}
913 <all,x2,ts1,t2c,t2b,t2a>\DeclareUnicodeCharacter{3009}{\textrangle}
914 <all,ts1>\DeclareUnicodeCharacter{2422}{\textblank}
915 <all,x2,t2c,t2b,t2a,t1,utf8>\DeclareUnicodeCharacter{2423}{\textvisiblespace}
916 <all,ts1>\DeclareUnicodeCharacter{25E6}{\textopenbullet}
917 <all,ts1>\DeclareUnicodeCharacter{25EF}{\textbigcircle}
918 <all,ts1>\DeclareUnicodeCharacter{266A}{\textmusicalnote}

919 <all,x2,ts1,t2c,t2b,t2a>\DeclareUnicodeCharacter{27E8}{\textlangle}
920 <all,x2,ts1,t2c,t2b,t2a>\DeclareUnicodeCharacter{27E9}{\textrangle}
921 <all,t1>\DeclareUnicodeCharacter{1E20}{\@tabacckludge=G}
922 <all,t1>\DeclareUnicodeCharacter{1E21}{\@tabacckludge=g}

```

When doing cut-and-paste from other documents f-ligatures might show up as Unicode characters. We translate them back to individual characters so that they get accepted. If supported by the font (which is normally the case) they are then reconstructed as ligatures so they come out as desired. Otherwise they will come out as individual characters which is fine too.

```

923 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB00}{ff} % ff
924 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB01}{fi} % fi
925 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB02}{fl} % fl
926 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB03}{ffi} % ffi

```

```

927 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB04}{ffl} % ffl
928 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB05}{st} % st -- this is the long s (not
929 <all,t1,ot1,ly1,t2a,t2b,t2c>\DeclareUnicodeCharacter{FB06}{st} % st
930 <all,ts1,utf8>\DeclareUnicodeCharacter{FEFF}{\ifhmode\nobreak\fi}

```

3.3 Notes

The following inputs are inconsistent with the 8-bit inputenc files since they will always only produce the ‘text character’. This is an area where inputenc is notoriously confused.

```

%<all,ts1,t1,ot1,ly1>\DeclareUnicodeCharacter{00A3}{\textsterling}
%<*all,x2,ts1,t2c,t2b,t2a,oms,ly1>
\DeclareUnicodeCharacter{00A7}{\textsection}
%</all,x2,ts1,t2c,t2b,t2a,oms,ly1>
%<all,ts1,utf8>\DeclareUnicodeCharacter{00A9}{\textcopyright}
%<all,ts1>\DeclareUnicodeCharacter{00B1}{\textpm}
%<all,ts1,oms,ly1>\DeclareUnicodeCharacter{00B6}{\textparagraph}
%<all,ts1,oms,ly1>\DeclareUnicodeCharacter{2020}{\textdagger}
%<all,ts1,oms,ly1>\DeclareUnicodeCharacter{2021}{\textdaggerdbl}
%<all,ly1,utf8>\DeclareUnicodeCharacter{2026}{\textellipsis}

```

The following definitions are in an encoding file but have no direct equivalent in Unicode, or they simply do not make sense in that context (or we have not yet found anything or ...:-). For example, the non-combining accent characters are certainly available somewhere but these are not equivalent to a $\text{\TeX}$ accent command.

```

\DeclareTextSymbol{\j}{OT1}{17}
\DeclareTextSymbol{\SS}{T1}{223}
\DeclareTextSymbol{\textcompwordmark}{T1}{23}

\DeclareTextAccent{"}{OT1}{127}
\DeclareTextAccent{'}{OT1}{19}
\DeclareTextAccent{\.}{OT1}{95}
\DeclareTextAccent{\=}{OT1}{22}
\DeclareTextAccent{\H}{OT1}{125}
\DeclareTextAccent{\^}{OT1}{94}
\DeclareTextAccent{\'}{OT1}{18}
\DeclareTextAccent{\r}{OT1}{23}
\DeclareTextAccent{\u}{OT1}{21}
\DeclareTextAccent{\v}{OT1}{20}
\DeclareTextAccent{\~}{OT1}{126}
\DeclareTextCommand{\b}{OT1}[1]
\DeclareTextCommand{\c}{OT1}[1]
\DeclareTextCommand{\d}{OT1}[1]
\DeclareTextCommand{\k}{T1}[1]

```

3.4 Mappings for OT1 glyphs

This is even more incomplete as again it covers only the single glyphs from OT1 plus some that have been explicitly defined for this encoding. Everything that is

provided in T1, and that could be provided as composite glyphs via OT1, could and probably should be set up as well. Which leaves the many things that are not provided in T1 but can be provided in OT1 (and in T1) by composite glyphs.

Stuff not mapped (note that \j (j) is not equivalent to any Unicode character):

```
\DeclareTextSymbol{\j}{OT1}{17}
\DeclareTextAccent{"}{OT1}{127}
\DeclareTextAccent{\'}{OT1}{19}
\DeclareTextAccent{\.}{OT1}{95}
\DeclareTextAccent{\=}{OT1}{22}
\DeclareTextAccent{\~}{OT1}{94}
\DeclareTextAccent{\'}{OT1}{18}
\DeclareTextAccent{\~}{OT1}{126}
\DeclareTextAccent{\H}{OT1}{125}
\DeclareTextAccent{\u}{OT1}{21}
\DeclareTextAccent{\v}{OT1}{20}
\DeclareTextAccent{\r}{OT1}{23}
\DeclareTextCommand{\b}{OT1}[1]
\DeclareTextCommand{\c}{OT1}[1]
\DeclareTextCommand{\d}{OT1}[1]
```

3.5 Mappings for OMS glyphs

Characters like \textbackslash are not mapped as they are (primarily) only in the lower 127 and the code here only sets up mappings for UTF-8 characters that are at least 2 octets long.

```
\DeclareTextSymbol{\textbackslash}{OMS}{110}      % "6E
\DeclareTextSymbol{\textbar}{OMS}{106}            % "6A
\DeclareTextSymbol{\textbraceleft}{OMS}{102}      % "66
\DeclareTextSymbol{\textbraceright}{OMS}{103}     % "67
```

But the following (and some others) might actually lurk in Unicode somewhere...

```
\DeclareTextSymbol{\textasteriskcentered}{OMS}{3} % "03
\DeclareTextCommand{\textcircled}{OMS}
```

3.6 Mappings for TS1 glyphs

Exercise for somebody else.

3.7 Mappings for latex.ltx glyphs

There is also a collection of characters already set up in the kernel, one way or the other. Since these do not clearly relate to any particular font encoding they are mapped when the utf8 support is first set up.

Also there are a number of \providecommands in the various input encoding files which may or may not go into this part.

```
931 <utf8>
932 % This space is intentionally empty ...
933 </utf8>
```

3.8 Old utf8.def file as a temp fix for pTeX and friends

```

934 <*utf8-2018>
935 \ProvidesFile{utf8.def}
936 [2018/10/05 v1.2f UTF-8 support for inputenc]
937 \makeatletter
938 \catcode'\ \saved@space@catcode
939 \long\def\UTFviii@two@octets#1#2{\expandafter
940   \UTFviii@defined\csname u8:#1\string#2\endcsname}
941 \long\def\UTFviii@three@octets#1#2#3{\expandafter
942   \UTFviii@defined\csname u8:#1\string#2\string#3\endcsname}
943 \long\def\UTFviii@four@octets#1#2#3#4{\expandafter
944   \UTFviii@defined\csname u8:#1\string#2\string#3\string#4\endcsname}
945 \def\UTFviii@defined#1{%
946   \ifx#1\relax
947     \if\relax\expandafter\UTFviii@checkseq\string#1\relax\relax
948       \UTFviii@undefined@err{#1}%
949     \else
950       \PackageError{inputenc}{Invalid UTF-8 byte sequence}%
951       \UTFviii@invalid@help
952     \fi
953   \else\expandafter
954     #1%
955   \fi
956 }
957 \def\UTFviii@invalid@err#1{%
958   \PackageError{inputenc}{Invalid UTF-8 byte "\UTFviii@hexnumber{‘#1}}}%
959   \UTFviii@invalid@help}
960 \def\UTFviii@invalid@help{%
961   The document does not appear to be in UTF-8 encoding.\MessageBreak
962   Try adding \noexpand\UseRawInputEncoding as the first line of the file.\MessageBreak
963   or specify an encoding such as \noexpand\usepackage[latin1]{inputenc}.\MessageBreak
964   in the document preamble.\MessageBreak
965   Alternatively, save the file in UTF-8 using your editor or another tool}
966 \def\UTFviii@undefined@err#1{%
967   \PackageError{inputenc}{Unicode character \expandafter
968     \UTFviii@splitcsname\string#1\relax
969     \MessageBreak
970     not set up for use with LaTeX}%
971   {You may provide a definition with\MessageBreak
972     \noexpand\DeclareUnicodeCharacter}%
973 }
974 \def\UTFviii@checkseq#1:#2#3{%
975   \ifnum‘#2<"80 %
976     \ifx\relax#3\else1\fi
977   \else
978     \ifnum‘#2<"C0 %
979       1 %
980     \else
981       \expandafter\expandafter\expandafter\UTFviii@check@continue
982       \expandafter\expandafter\expandafter#3%
983     \fi
984   \fi}
985 \def\UTFviii@check@continue#1{%

```

```

986 \ifx\relax#1%
987 \else
988 \ifnum'#1<"80 1\else\ifnum'#1>"BF 1\fi\fi
989 \expandafter\UTFviii@check@continue
990 \fi
991 }
992 \begingroup
993 \catcode'\~13
994 \catcode'\ "12
995 \def\UTFviii@loop{%
996 \uccode'\~\count@
997 \uppercase\expandafter{\UTFviii@tmp}%
998 \advance\count@<\@one
999 \ifnum\count@<@\tempcnta
1000 \expandafter\UTFviii@loop
1001 \fi}
1002 \def\UTFviii@tmp{\xdef~{\noexpand\UTFviii@undefined@err{:~\string~}}
1003 \count@"1
1004 \@tempcnta9
1005 \UTFviii@loop
1006 \count@"11
1007 \@tempcnta12
1008 \UTFviii@loop
1009 \count@"14
1010 \@tempcnta32
1011 \UTFviii@loop
1012 \count@"80
1013 \@tempcnta"C2
1014 \def\UTFviii@tmp{\xdef~{\noexpand\UTFviii@invalid@err\string~}}
1015 \UTFviii@loop
1016 \count@"C2
1017 \@tempcnta"E0
1018 \def\UTFviii@tmp{\xdef~{\noexpand\UTFviii@two@octets\string~}}
1019 \UTFviii@loop
1020 \count@"E0
1021 \@tempcnta"F0
1022 \def\UTFviii@tmp{\xdef~{\noexpand\UTFviii@three@octets\string~}}
1023 \UTFviii@loop
1024 \count@"F0
1025 \@tempcnta"F5
1026 \def\UTFviii@tmp{\xdef~{\noexpand\UTFviii@four@octets\string~}}
1027 \UTFviii@loop
1028 \count@"F5
1029 \@tempcnta"100
1030 \def\UTFviii@tmp{\xdef~{\noexpand\UTFviii@invalid@err\string~}}
1031 \UTFviii@loop
1032 \endgroup
1033 \@inpenctest
1034 \ifx\@begindocumenthook\@undefined
1035 \makeatother
1036 \endinput \fi
1037 \begingroup
1038 \catcode'\ "=12
1039 \catcode'\ <=12

```

```

1040 \catcode'\.=12
1041 \catcode'\,=12
1042 \catcode'\;=12
1043 \catcode'\!=12
1044 \catcode'\~=13
1045 \gdef\DeclareUnicodeCharacter#1#2{%
1046   \count@"#1\relax
1047   \wlog{ \space\space defining Unicode char U+#1 (decimal \the\count@)}%
1048   \begingroup
1049     \parse@XML@charref
1050     \def\UTFviii@two@octets##1##2{\csname u8:##1\string##2\endcsname}%
1051     \def\UTFviii@three@octets##1##2##3{\csname u8:##1%
1052                                     \string##2\string##3\endcsname}%
1053     \def\UTFviii@four@octets##1##2##3##4{\csname u8:##1%
1054                                     \string##2\string##3\string##4\endcsname}%
1055     \expandafter\expandafter\expandafter
1056     \expandafter\expandafter\expandafter
1057     \expandafter
1058     \gdef\UTFviii@tmp{\IeC{#2}}%
1059   \endgroup
1060 }
1061 \gdef\parse@XML@charref{%
1062   \ifnum\count@<"A0\relax
1063     \ifnum\catcode\count@=13
1064       \uccode'\~=count@\uppercase{\def\UTFviii@tmp{\@empty\@empty~}}%
1065     \else
1066       \PackageError{inputenc}{Cannot define non-active Unicode char value < 00A0}%
1067       \@eha
1068       \def\UTFviii@tmp{\UTFviii@tmp}%
1069     \fi
1070   \else\ifnum\count@<"800\relax
1071     \parse@UTFviii@a,%
1072     \parse@UTFviii@b C\UTFviii@two@octets.,%
1073   \else\ifnum\count@<"10000\relax
1074     \parse@UTFviii@a;%
1075     \parse@UTFviii@a,%
1076     \parse@UTFviii@b E\UTFviii@three@octets.{,;}%
1077   \else
1078     \ifnum\count@>"10FFFF\relax
1079       \PackageError{inputenc}%
1080         {\UTFviii@hexnumber\count@\space too large for Unicode}%
1081         {Values between 0 and 10FFFF are permitted}%
1082     \fi
1083     \parse@UTFviii@a;%
1084     \parse@UTFviii@a,%
1085     \parse@UTFviii@a!%
1086     \parse@UTFviii@b F\UTFviii@four@octets.{!,;}%
1087   \fi
1088   \fi
1089   \fi
1090 }
1091 \gdef\parse@UTFviii@a#1{%
1092   \@tempcnta\count@
1093   \divide\count@ 64

```



```

1094 \count@tempcntb\count@
1095 \multiply\count@ 64
1096 \advance\count@tempcnta-\count@
1097 \advance\count@tempcnta 128
1098 \uccode' #1\count@tempcnta
1099 \count@\count@tempcntb}
1100 \gdef\parse@UTFviii@b#1#2#3#4{%
1101 \advance\count@ "#10\relax
1102 \uccode'#3\count@
1103 \uppercase{\gdef\UTFviii@tmp{#2#3#4}}
1104 \ifx\numexpr\undefined
1105 \gdef\decode@UTFviii#1{0}
1106 \else
1107 \gdef\decode@UTFviii#1\relax{%
1108 \expandafter\UTFviii@cleanup
1109 \the\numexpr\dec@de@UTFviii#1\relax}})\@empty}
1110 \gdef\UTFviii@cleanup#1#2\@empty{#1}
1111 \gdef\dec@de@UTFviii#1{%
1112 \ifx\relax#1%
1113 \else
1114 \ifnum'#1>"EF
1115 ((((' #1-"F0)%
1116 \else
1117 \ifnum'#1>"DF
1118 ((((' #1-"E0)%
1119 \else
1120 \ifnum'#1>"BF
1121 ((((' #1-"C0)%
1122 \else
1123 \ifnum'#1>"7F
1124 )*64+(' #1-"80)%
1125 \else
1126 +' #1 %
1127 \fi
1128 \fi
1129 \fi
1130 \fi
1131 \expandafter\dec@de@UTFviii
1132 \fi}
1133 \fi
1134 \ifx\numexpr\undefined
1135 \global\let\UTFviii@hexnumber\@firstofone
1136 \global\UTFviii@hexdigit\hexnumber@
1137 \else
1138 \gdef\UTFviii@hexnumber#1{%
1139 \ifnum#1>15 %
1140 \expandafter\UTFviii@hexnumber\expandafter{\the\numexpr(#1-8)/16\relax}%
1141 \fi
1142 \UTFviii@hexdigit{\numexpr#1\ifnum#1>0-((#1-8)/16)*16\fi\relax}%
1143 }
1144 \gdef\UTFviii@hexdigit#1{\ifcase\numexpr#1\relax
1145 0\or1\or2\or3\or4\or5\or6\or7\or8\or9\or
1146 A\or B\or C\or D\or E\or F\fi}
1147 \fi

```

```

1148 \gdef\UTFviii@hexcodepoint#1{U+%
1149 \ifnum#1<16 0\fi
1150 \ifnum#1<256 0\fi
1151 \ifnum#1<4096 0\fi
1152 \UTFviii@hexnumber{#1}%
1153 }%
1154 \gdef\UTFviii@splitcsname#1:#2\relax{%
1155 #2 (\expandafter\UTFviii@hexcodepoint\expandafter{%
1156 \the\numexpr\decode@UTFviii#2\relax})%
1157 }
1158 \endgroup
1159 \@onlypreamble\DeclareUnicodeCharacter
1160 \@onlypreamble\parse@XML@charref
1161 \@onlypreamble\parse@UTFviii@a
1162 \@onlypreamble\parse@UTFviii@b
1163 \begingroup
1164 \def\cdp@elt#1#2#3#4{%
1165 \wlog{Now handling font encoding #1 ...}%
1166 \lowercase{%
1167 \InputIfFileExists{#1enc.dfu}}%
1168 {\wlog{... processing UTF-8 mapping file for font %
1169 encoding #1}%
1170 \catcode'\ 9\relax}%
1171 {\wlog{... no UTF-8 mapping file for font encoding #1}}%
1172 }
1173 \cdp@list
1174 \endgroup
1175 \def\DeclareFontEncoding@#1#2#3{%
1176 \expandafter
1177 \ifx\csname T@#1\endcsname\relax
1178 \def\cdp@elt{\noexpand\cdp@elt}%
1179 \xdef\cdp@list{\cdp@list\cdp@elt{#1}%
1180 {\default@family}{\default@series}%
1181 {\default@shape}}%
1182 \expandafter\let\csname#1-cmd\endcsname\@changed@cmd
1183 \begingroup
1184 \wlog{Now handling font encoding #1 ...}%
1185 \lowercase{%
1186 \InputIfFileExists{#1enc.dfu}}%
1187 {\wlog{... processing UTF-8 mapping file for font %
1188 encoding #1}}%
1189 {\wlog{... no UTF-8 mapping file for font encoding #1}}%
1190 \endgroup
1191 \else
1192 \@font@info{Redefining font encoding #1}%
1193 \fi
1194 \global\@namedef{T@#1}{#2}%
1195 \global\@namedef{M@#1}{\default@M#3}%
1196 \xdef\LastDeclaredEncoding{#1}%
1197 }
1198 \DeclareUnicodeCharacter{00A9}{\textcopyright}
1199 \DeclareUnicodeCharacter{00AA}{\textordfeminine}
1200 \DeclareUnicodeCharacter{00AE}{\textregistered}
1201 \DeclareUnicodeCharacter{00BA}{\textordmasculine}

```

```

1202 \DeclareUnicodeCharacter{02C6}{\textasciicircum}
1203 \DeclareUnicodeCharacter{02DC}{\textasciitilde}
1204 \DeclareUnicodeCharacter{200C}{\textcompwordmark}
1205 \DeclareUnicodeCharacter{2026}{\textellipsis}
1206 \DeclareUnicodeCharacter{2122}{\texttrademark}
1207 \DeclareUnicodeCharacter{2423}{\textvisiblespace}
1208 \DeclareUnicodeCharacter{FEFF}{\ifhmode\nobreak\fi}
1209 \endinput
1210 </utf8-2018>

```

4 A test document

Here is a very small test document which may or may not survive if the current document is transferred from one place to the other.

```

1211 <*test>
1212 \documentclass{article}
1213
1214 \usepackage[latin1,utf8]{inputenc}
1215 \usepackage[T1]{fontenc}
1216 \usepackage{trace}
1217
1218 \scrollmode % to run past the error below
1219
1220 \begin{document}
1221
1222 German umlauts in UTF-8: ^^c3^^a4^^c3^^b6^^c3^^bc %% äü
1223
1224 \inputencoding{latin1} % switch to latin1
1225
1226 German umlauts in UTF-8 but read by latin1 (and will produce one
1227 error since \verb=\textcurrency= is not provided):
1228 ^^c3^^a4^^c3^^b6^^c3^^bc
1229
1230 \inputencoding{utf8} % switch back to utf8
1231
1232 German umlauts in UTF-8: ^^c3^^a4^^c3^^b6^^c3^^bc
1233
1234
1235 Some codes that should produce errors as nothing is set up
1236 for them: ^^c3F ^^e1^^a4^^b6
1237
1238 And some that are not legal utf8 sequences: ^^c3X ^^e1XY
1239
1240 \showoutput
1241 \tracingstats=2
1242 \stop
1243 </test>

```