

Secretariat of CPMI (cpmi@bis.org)
Secretariat of IOSCO (upi@iosco.org)

Paris, Monday 29 August 2016

CPMI - IOSCO: Second consultative report on Harmonisation of the Unique Product Identifier

Dear Sirs,

Please find below our contribution to the consultation.

Question 1: Do you believe that the data elements within each asset class described above are appropriate? Why or why not? If there are additional subcategories that you believe should be included for one or more asset classes, please describe them and discuss why you believe they should be included.

Data elements described are appropriate.

Question 2: Do you believe generally that the value "Other" is required in certain data elements? If so, which ones and why?

The Value Other is required, in case of innovative data elements.

Question 3: For an OTC derivative product based on a custom basket of securities or assets, please provide your view of the optimal means of representing that OTC derivative product. Do you believe that it is practical to include all of the underlying securities or assets and their risk weights in the UPI reference data? If not, how do you believe that the elements of the custom basket and their risk weights should be reported to a TR?

It is not practical to include all of the underlying securities or asset in the UPI reference data. A reference of the "nature" or "type" of the elements of the custom basket and their average or consolidated risk weights could be reported to a TR.

Question 4: How should underlying assets and reference entities be represented in the UPI data library? Would LEIs be suitable, at least for corporate reference entities? Why or why not? Are there suitable identifiers for indices? If not, is it feasible to use an existing identifier such as an ISIN code for them?

It is not appropriate to use LEIs for corporate reference entities for confidentiality reasons.

Question 5: Do you envisage any obstacles to including the source of the identifier for the underlier as part of the reference data element for the underlier? Please explain and justify.

It is not appropriate to include the source of the identifier as part of the reference data element for the underlier, as sources of identifier may not be unique, would require public dissemination. This information would require additional work on standardization. This information could be requested on an ad-hoc basis in case of investigation by regulators.

Question 6: Could there be issues related to including proprietary benchmarks and indices in publicly available reference data or publicly disseminated UPIs? Please elaborate on any issues, such as licensing, that may exist.

Question 7: What are the arguments for and against the use of a dummy UPI code or an intelligent UPI code, or having both types of code coexisting?

Whatever the selection (either dummy or intelligent), it needs to be consistent across all jurisdictions.

Question 8: Do you agree that a well-articulated UPI reference data library could support interoperability between dummy UPI codes and intelligent UPI codes? Why or why not? What steps could be taken with the UPI reference data to facilitate supporting both types of UPI code?

Whatever the selection (either dummy or intelligent), it needs to be consistent across all jurisdictions. We are of the opinion that intelligent code would be a better option, as it would facilitate aggregation (by asset class, by product type...).

Question 9: What are the minimum and maximum lengths (in terms of number of characters) that you believe the industry could accommodate for a UPI code system? How does this vary between dummy and intelligent codes? What do you believe is the optimal number of characters, and why?

Minimum length could be six characters. There need to be a maximum in case dummy codes are selected. 20 characters could be workable solution.

Question 10: For intelligent codes, how should the information be encoded? Are there existing models for this? How much adaptation would existing models require in order to meet the needs described in this consultation?

For intelligent codes, the encoding should enable aggregation by asset type, by product type, by nature of underlying.

A solution whereby, the first two characters would be devoted to asset type (FX, IR, CO, EQ...), the following two to products (SW, FW, OP,...); the following x to any required additional information for market risk assessment.....

Question 11: Do you believe that UPI codes should have an inherent means of validation? For example, should UPI codes include a check digit? Why or why not? Does this vary between dummy and intelligent codes and/or depend on the encoding method used in an intelligent code?

UPI codes would require an inherent means of validation. This could be easier to achieve with intelligent code (format+ check digit validation).

Question 12: Another means of having a simple, partial validation for a UPI code would be for all UPI codes to be of uniform length: thus, any code that was not of the required length could be recognised as prima facie invalid. Do you believe that all UPI codes should be of uniform length? Why or why not? Or are optimal UPI codes of one asset class likely to be longer or shorter than optimal UPI codes for other asset classes? If so, do you believe that extra dummy characters should be inserted into the shorter codes to make them of the uniform length? Why or why not?

In any case, UPI codes should be of maximum length, so as to be integrated into systems. For intelligent codes, there is no specific reason to favor a uniform length as the structure of the code would enable code checking. For dummy code, uniform length (at least per asset class) could facilitate prima facie validation.

Question 13: For an intelligent UPI code, how should underlying the asset(s) or reference entity (entities) be represented within the UPI code? Would it be preferable for the part of the UPI code that represents the underlying asset(s) or reference entity (entities) to be dummy while the rest of the code is intelligent? Why or why not?

It could be of interest to add reference to the underlying asset (s) or entity (ies) in the PI code, but for confidentiality reasons, represented as dummy.

Question 14: Should the UPI code system avoid using Roman letters? Why or why not? Are there particular jurisdictions whose computer systems cannot accommodate Roman letters? UPI code should use Roman letters, at least when intelligent code is selected, which makes the code easily controlled and understood as far as asset class or product is concerned.

Question 15: Would it be preferable for the UPI code system to use only Roman letters, only Indo-Arabic numerals, or a combination of the two? Why? If Roman letters are included in the UPI code system, should they avoid being case-sensitive? If the UPI code system uses both Roman letters and Indo-Arabic numerals, should the system not disallow particular characters that could be mistaken for each other (the lower-case letter "l" and the number "1", the digit "0" and the upper-case letter "O" etc)

UPI code should include both Roman letters for the intelligent part of the code, and Indo-Arabic numerals for the dummy part. Roman letter should avoid being case-sensitive. There is no need to disallow specific characters as Roman letters and Indo-Arabic numerals would be used in different "location" and could therefore not be mistaken one for the other. It is suitable to disallow specific characters such as \$,&,!,% which may lead to mistakes.



Valérie SAINSAULIEU
Présidente commission AFTE
"Conformité et contrôle interne"