

Comments

on CPMI/IOSCO's Second consultative report
Harmonisation of the Unique Product Identifier (UPI)

Register of Interest Representatives
Identification number in the register: 52646912360-95

Contact:

Oliver Wirnhier

Division Manager

Telephone: +49 30 1663-3320

Fax: +49 30 1663-3399

Email: oliver.wirnhier@bdb.de

Berlin, 30 September 2016

The **German Banking Industry Committee (GBIC)** is the joint committee operated by the central associations of the German banking industry. These associations are the Bundesverband der Deutschen Volksbanken und Raiffeisenbanken (BVR), for the cooperative banks, the Bundesverband deutscher Banken (BdB), for the private commercial banks, the Bundesverband Öffentlicher Banken Deutschlands (VÖB), for the public banks, the Deutscher Sparkassen- und Giroverband (DSGV), for the savings banks finance group, and the Verband deutscher Pfandbriefbanken (vdp), for the Pfandbrief banks. Collectively, they represent approximately 1,700 banks.

Coordinator:
Association of German Banks
Burgstraße 28 | 10178 Berlin | Germany
Telephone: +49 30 1663-0
Telefax: +49 30 1663-1399
www.die-deutsche-kreditwirtschaft.de

GBIC comments on harmonisation of the Unique Product Identifier (UPI)

The German Banking Industry Committee (GBIC) expressly welcomes the global (jurisdiction-agnostic) approach of the system proposed in CPMI/IOSCO's second consultative report. The precise hierarchy and open-source concept support, in our view, a widespread proliferation and rapid as possible introduction of a Unique Product Identifier (UPI). An UPI based on the proposed principles should significantly mitigate the data aggregation problem that still exists.

However, we deem it essential that, notwithstanding CPMI/IOSCO's mandate to enable the aggregation of data by authorities, the work of the Harmonisation Group is aligned with other regulatory activities in the field of derivatives categorisation; the most prominent example being the current work of the Association of National Numbering Agencies (ANNA) with a view to designating ISIN for financial instruments, including OTC derivatives, for reporting purposes under MiFID and MiFIR.

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.

Firstly, we would suggest separating the information on the underlying asset and the contract type into two different categories, because (as especially the asset class "Commodities" illustrates) transactions can have the same underlying asset and different contract types or share the same contract type for different underliers.

Apart from that, we believe that certain derivatives products might benefit from an enhancement of the reference data. For some types of interest-rate derivatives, the tenor may apply as a distinctive feature of an OTC derivatives product. For instance, Market-Agreed Coupon (MAC) swaps, being interest-rate contracts with pre-defined market-agreed terms, are currently set quarterly for six currencies and 12 tenors. CUSIP and ISIN numbers are assigned for each MAC swap contract.

Secondly, we wonder how quanto derivatives, i.e. derivatives in which the underlying is denominated in one currency, but the instrument itself is settled in another currency at some rate, can be mirrored in CPMI/IOSCO's reference data. For instance, an IRS denominated in USD with a floating leg making reference to a GBP Libor index should not be characterised merely by the underlier Libor, since the contract as such does not make reference to the USD Libor. The mapping of non-deliverable swaps appears questionable, too. Could you please provide an explanation of how to distinguish a non-deliverable KRW/USD cross-currency swap (USD Libor float leg vs KRW fixed leg) whose settlement currency is USD from a plain fixed vs float USD Libor IR swap. Finally, the exact depiction of swaps with underliers that vary over the contract's tenor appears to be difficult under the current reference data. Given a contract that for its first 5 years refers to the 3M Libor and switches to 6M Libor from year 6, neither the field "Single or multiple tenor" nor the field "Underlying rate index tenor period multiplier" seems to be able to reflect this specificity.

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

A reference data system should allow the value "Other" for the following reasons:

GBIC comments on harmonisation of the Unique Product Identifier (UPI)

- It may not be possible to provide a system of UPI reference data for each existing type of derivative instrument. Some products lack the necessary degree of standardisation to warrant a stand-alone classification.
- The availability of the value "Other" would help to achieve a level of granularity which should keep the number of product groups that contain only a single or a limited number of transactions to a minimum.
- The possession of an "Other" bucket allows users to classify products that do not fit precisely into pre-defined classification values. Otherwise, market participants may tend to categorise products into buckets that are not entirely accurate for the sole reason that an "Other" bucket is not available.

Such an "Other" category should be allowed for any data element, as otherwise some trades might not obtain an UPI at all. In case the number of trades captured as "Other" increases to a significant level, the generation of a new UPI (based on an analysis of the trade details in the reporting set) might become useful later.

Since the treatment, criteria, and process for the "Other" bucket should be part of the governance oversight for the UPI, we would recommend a governance process which enables the product identifier to evolve in line with market needs and undergo major revision updates on a regular basis. According to the UPI adaptability principle, OTC derivative products originally falling into the "Other" category may be recategorised as its product type becomes more common and is traded more widely.

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?

We believe that an approach aimed at identifying each element of an underlying basket must be regarded as inappropriate. Information on whether a product's underlier is a "custom basket" or an "index" appears to be sufficient. This is because the composition of a basket may change over time and baskets that have been traded on exchange do not need to be decomposed and reflect the components, as the information will be defined in the static data from the exchanges under their unique identifier for the basket. As far as the representation of reference debtors is concerned, we would suggest allowing ISO 3166 country codes and ISO 3166-2 in order to identify contracts on municipalities.

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?

Regarding underlying assets, we would advocate the use of ISIN. The LEI should be the means of choice to represent reference entities of OTC products since, by virtue of ISO 17442, it allows for the necessary bijective identification of any legal entity. Furthermore, due to regulatory requirements such as EMIR and MiFID, the LEI is by now the most prominent and widespread identifier for legal persons.

GBIC comments on harmonisation of the Unique Product Identifier (UPI)

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.

We believe that referring to the source of an underlier would create much too granular information, as there may be more than one provider for one and the same underlier. The identification of the underlier itself should suffice in this context.

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.

We believe the main issues would be any licensing issues concerning the public dissemination of proprietary benchmarks and indices. It is also unclear whether including these proprietary benchmarks and indices when reporting to the TRs would breach any licensing terms of the provider.

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?

We believe that the rationale as to whether the UPI is designed as a dummy code or an intelligent code should be consistent with the architecture of other identifiers, e.g. the UTI. Furthermore, such decision should at least, inter alia, take into consideration the code's operability from a user's perspective.

That said, we would opt for a UPI being conceived as a partially intelligent code, since it is a much less complex approach as opposed to a fully intelligent code and any additional transaction data will be available via the trade repository. In contrast, a fully intelligent code might prove to be unmanageable, bearing in mind that the UPI would often have to be processed manually by reporting entities. The number of data elements involved (e.g. 14 elements for asset class rates) and its variations within one data element (abbreviations of a necessary length of 3-4 digits) would require an UPI code with a length of approx. 50 digits to become entirely "readable". An only partially intelligent code, on the other hand, might reduce erroneous data entries due to the possibility of self-monitoring the input process.

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?

There should be only one type of UPI – dummy, intelligent or partially intelligent – but not different types at the same time, as we cannot see any advantages accompanied with a hybrid approach.

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?

Given the number of UPI codes necessary to reflect the relevant reference data, we would set the length of the identifier at between a minimum of 5 characters and a maximum of 20 characters.

GBIC comments on harmonisation of the Unique Product Identifier (UPI)

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?

As we advocate a partially intelligent code, we would suggest implementing a rather simple and stylised approach as described in footnote 11.

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?

From a proportionality standpoint, the UPI encoding scheme should refrain from including check digits. We cannot see that the theoretical benefits of a check digit would make up for its cost of implementation. Furthermore, a partially intelligent UPI that encodes the main characteristics intelligently (such as asset class) offers the possibility of quick checks and thus reduces the risk of failure.

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?

Uniform length of UPI codes is preferable, as it enables a basic plausibility check. Filler digits do not present any problems.

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?

Such an approach appears preferable in order to achieve a UPI which is (partially) intelligent and satisfies users' requirements in terms of operability.

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?

No, since only number codes that have weaknesses regarding the intelligence carried by a single character would remain as an alternative. An alphabetic UPI would deliver much better legibility. In addition, we are not aware of any jurisdictions whose computer systems are unable to process Roman letters.

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

GBIC comments on harmonisation of the Unique Product Identifier (UPI)

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).

We would welcome an approach aimed at eliminating characters from the UPI that can be mistaken for others. The inclusion of both "O" (oh) and "0" (zero) characters in the LEI has shown that there is significant potential for confusion when the information is keyed into systems manually. The letters "O" (oh) and "I" (I) should be eliminated to remove the confusion with "l" (lower-case ell) and digits. The digits "0" (zero) and 1 (one) should be kept for dates and internal reference numbers.
