

SET Statement Considered Harmful

Bartosz Jabłoński, yabwon/Warsaw University of Technology

ABSTRACT

The **SET** statement is the primary way SAS® programmers access the data (I could bet \$5, it is over 90% of cases). Since data processing in SAS is I/O oriented the majority of optimization can be done by reducing I/O, i.e., the proper use of the **SET** statement¹ in the code. The article is dedicated to beginner SAS programmers and considered education focused. In the paper we present a group of examples showing how (mis)use of the **SET** statement can be (and is) harmful for data processing and how the grim situation can be fixed.

Table of contents

WHEN SET DOES NOT MEAN "ALL SET"	1	"DOING MORE BY DOING LESS"	8
THE DATA	2	"DOING BY NOT DOING"	9
THE LIST	3	"NO RAW MACRO LOOPS", PART 1.	10
"TROLLING IS A ART"	3	"NO RAW MACRO LOOPS", PART 2.	11
"TWO BIRDS WITH ONE STONE"	4	"#HASH TABLE FOR HELP"	14
"IT MAKES MY BLOOD BOIL"	4	"NO RAW MACRO LOOPS", PART 2. - REVISITED	15
"CUTTING, SLASHING, AND SHREDDING"	5	"ONE STEP TO RULE THEM ALL"	16
"THE FINAL (BACKWARD) COUNTDOWN"	6	"I HAVE SEEN THIS BEFORE"	17
"DISENCHANTING"	6	CONCLUSION	20
"DISENCHANTING" AGAIN	7	REFERENCES	20
"DISENCHANTING" A BIT MORE	7	Appendix A - code coloring guide	21
"NICE PIECE OF ABSOLUTELY USELESS STEP"	8	INDEX	21

WHEN SET DOES NOT MEAN "ALL SET"

*Toutes choses sont dites déjà; mais comme personne n'écoute, il faut toujours recommencer.*²

André Gide

Contrary to the famous article by Edgar Dijkstra (see [Dijkstra 1968]), in this text we are not going to discourage SAS programmers from using the **SET** statement. In fact, discouraging use of the **SET** statement could be rather hard (though not impossible) to achieve. In this article, we want bring beginner SAS programmer's attention to the following *good programming practice*, shared by Aster and Seidman in their book, to "*consider every SET statement¹ with suspicion*" (see [Aster & Seidman 1996]).

However strange as it might sound, a SAS data set (not a variable, not an array) is a fundamental data structure in SAS programming. All data communication is done through SAS data sets. Whenever we are placing data into the Program Data Vector (PDV) or pushing data off of it, we are using SAS data sets as a transport medium. Data sets by design are stored on disk drives and require input and output (I/O) operations to move data into and from memory. Even though modern disk drives went through huge technological transformation gaining a lot of speed, still I/O operations are one of the slowest (or rather most time-consuming) part of code execution. Thus, reducing the number of I/O operations seems to be the most natural way of improving program efficiency. If reading a data set one time takes some time, doing it twice doubles the time, etc. If one says: "*my data are small and read/write in fraction of a second*", we postulate that even small data, when read thousands of times, can become annoyingly "big".

¹Of course also statements like: **MERGE**, **INPUT**, **DATA=**, etc.

²See "Conclusion" section for explanation

Of course, working with the SET statement is not the only way to "interact" with the data; reading in from or writing out to external text files, or processing through library engines, also fit into I/O operations realm, and because of that, can be considered potentially "harmful."

The "SET" statement in the article's title is a representative of a wider class of statements also containing: the MERGE, INPUT, DATA, DATA=, OUT=, or any other I/O related statement. In fact, the title could be: *"Use of too many input/output statements can be potentially harmful for a program performance"* (but you must admit, it would not be "it"). In fact, in the Dijkstra's article, it is not that the GOTO³ statement is evil – it is the lack of structured programming that causes harm. Similarly, here, it is not the SET statement itself that is at fault. The thing we want to highlight is a general (so often haunting inexperienced SAS programmers) problem of writing code that in one or another way inefficiently uses I/O operations.

Being a beginner SAS programmer is not an easy job; that is why we want to help you in the journey by showing some of the most common I/O pitfalls and how to avoid them. To be clear, this article should *not* be considered exhaustive (frankly, we do not even dare to claim so). There is a rather high probability you will have to search through more literature to make yourself a better SAS programmer, but we hope you will be able to use this article as a convenient starting point or, at least, a "handy" support.

THE DATA

Let's assume we are starting with the following data set and text file:

```

code: data
1 options DLcreateDIR;
2 libname source "%sysfunc(pathname(WORK))/have";
3 filename source "%sysfunc(pathname(WORK))/have/have.txt";
4
5 data source.have;
6   call streaminit(42);
7   file source;
8   do grp = 1 to 3333 by 2, 2 to 3333 by 2;
9     length id $ 8;
10    do id = "A","B","C","D","E","F","G","H","I","J","K","L","M",
11          "N","O","P","Q","R","S","T","U","V","W","X","Y","Z";
12      do number = 1 to rand("integer",17,42);
13        obs+1;
14        output;
15        put grp id number obs;
16      end;
17    end;
18  end;
19 run;

```

The data set HAVE located in the library named SOURCE, is rather moderate in size (somewhere around 80MB in volume). HAVE has the following variables: grp, id, number, and obs. Table 1 displays part of the data set. The LOG summarizes it as:

```

the log - observations and variables
1 NOTE: The data set SOURCE.HAVE has 2'555'488 observations and 4 variables.

```

The have.txt text file created contains the same data and is approximately 45MB in volume.

³Yes, the SAS language also has GOTO statements, both in the 4GL and the macro language.

Table 1: Data set SOURCE.HAVE, example

grp	id	number	obs	grp	id	number	obs	grp	id	number	obs	grp	id	number	obs
1	A	1	1	1	L	21	373	1	X	1	717	3	I	28	1090
1	A	30	30	1	M	1	374	1	X	18	734	3	J	1	1091
1	B	1	31	1	M	22	395	1	Y	1	735	3	J	24	1114
1	B	36	66	1	N	1	396	1	Y	26	760	3	K	1	1115
1	C	1	67	1	N	32	427	1	Z	1	761	3	K	27	1141
1	C	20	86	1	O	1	428	1	Z	37	797	3	L	1	1142
1	D	1	87	1	O	37	464	3	A	1	798	3	L	35	1176
1	D	36	122	1	P	1	465	3	A	20	817	3	M	1	1177
1	E	1	123	1	P	32	496	3	B	1	818	3	M	30	1206
1	E	35	157	1	Q	1	497	3	B	37	854	3	N	1	1207
1	F	1	158	1	Q	42	538	3	C	1	855	3	N	40	1246
1	F	40	197	1	R	1	539	3	C	37	891	3	O	1	1247
1	G	1	198	1	R	35	573	3	D	1	892	3	O	37	1283
1	G	34	231	1	S	1	574	3	D	17	908	3	P	1	1284
1	H	1	232	1	S	38	611	3	E	1	909	3	P	28	1311
1	H	18	249	1	T	1	612	3	E	42	950	3	Q	1	1312
1	I	1	250	1	T	37	648	3	F	1	951	3	Q	33	1344
1	I	42	291	1	U	1	649	3	F	40	990	3	R	1	1345
1	J	1	292	1	U	24	672	3	G	1	991	3	R	25	1369
1	J	22	313	1	V	1	673	3	G	33	1023	3	S	1	1370
1	K	1	314	1	V	17	689	3	H	1	1024	3	S	37	1406
1	K	39	352	1	W	1	690	3	H	39	1062	3	T	1	1407
1	L	1	353	1	W	27	716	3	I	1	1063	3	T	24	1430

THE LIST

Examples presented here have (more or less) the following form: first there is a short introduction, then (usually) two parallel snippets of code are compared, and some comments are shared. To be 100% clear, those examples and techniques used should not be perceived as "silver bullets", they are rather "80/20 rules" – i.e., *usually* working best practices. It might happen that your specific programming setup could be so exotic that you will basically have to use extra [SET](#), [MERGE](#), or similar statement, but that's life...

"TROLLING IS A ART"

Finding the following snippet in code you read you can be almost certain someone is playing a "troll."

```

code: the existing situation
1  /* a troll */
2  data source.have;
3  set source.have;
4  run;

```

```

code: a smarter approach
1  /* solution */
2
3  /* a rare case of no-code/low-code */
4  /* beating programming */

```

The DATA step is not only wasting I/O operations, your time, and space in the code file, but also can be potentially dangerous because overwriting may destroy potential indexes associated with the data set. And in case someone is really trying to get rid of an index, there are much better ways to do it, like [PROC DATASETS](#), which we are going to talk about in a moment.

In this case we replace one data reading (I, input) and one data write (O, output) with zero, so the I/O savings can be symbolically presented as: $2I/O \rightarrow 0I/O$, which gives an infinite (∞) gain.

"TWO BIRDS WITH ONE STONE"

When you need only a vertical subset (i.e., only selected variables) of your data for some sort of BY-group processing (i.e., sorting is required) you can do both in a single processing step.

code: the existing situation

```

1  /* two birds... */
2  data work.have;
3    SET source.have;
4    drop obs;
5  run;
6
7  proc sort data=work.have;
8    by id grp;
9  run;
```

code: a smarter approach

```

1  /* ...one stone */
2  proc sort
3    data=source.have(drop=obs)
4    out=work.have
5  ;
6    by id grp;
7  run;
```

Instead of copying the sub-selected data and then sorting, we can do sorting directly on a subset of data. With this approach we replace two data reading (SET statement and DATA= option) and two data write (DATA statement and implicit output of the PROC SORT) with just two (DATA= and OUT= options), and the I/O reduction in this case is $4I/O \rightarrow 2I/O$.

[Note] This is a good moment to state a disclaimer. We are aware that such a simplified approach (i.e., presenting $xI/O \rightarrow yI/O$ reduction purely based on the number of SAS statements) and, especially, its precision might be debated, but the aim of the article is not to provide exact (up to every byte) I/O estimations. Our idea is to point out some programming styles which can result in better (or worse) performance of our programs.

"IT MAKES MY BLOOD BOIL"

I had a chance to see the following snippet multiple times when reviewing a coworker's code, and it always made (and still makes) my blood boil. It is almost like the "troll" but worse...

code: the existing situation

```

1  /* "blood boiler" */
2  data work.have;
3    SET work.have;
4    format number ROMAN12.;
5  run;
```

code: a smarter approach

```

1  /* "cooler" */
2  proc datasets lib=work noprint;
3    modify have;
4    format number ROMAN12.;
5    run;
6  quit;
```

Instead of data read and write just to modify metadata, i.e., format assigning to the number variable, use of the PROC DATASETS makes much more sense. With the PROC DATASETS we can achieve our goal with just a fraction of the original I/O operations sacrifice. The bottom line is: $2I/O \rightarrow 0.1I/O$. We cannot write $0I/O$ because the PROC DATASETS has to touch metadata, what involves some I/O operations, but definitely much less than a "full" data step.

The Polish language idiom for "it makes my blood boil" is: "nóż się w kieszeni otwiera"⁴ what literally translates to "a knife opens in one's pocket", and this reminds me that a very good introduction (and more) to PROC DATASETS can be found in [Raithel 2010 & 2018].

⁴See: https://en.wiktionary.org/wiki/n%C3%B3%C5%BC_si%C4%99_w_kieszeni_otwiera

"CUTTING, SLASHING, AND SHREDDING"

Sometimes we have to split data into separate data sets according to values of a variable, or according to a bunch of different splitting rules.

To split by values, people often involve macro-loops, but usually not the way it is "the way" for macro loops. Both snippets below, the left and the right one, can be imagined as a result of executing %DO-LOOP(s) over values: 1, 2, and 3.

code: the existing situation

```
1  /* silly */
2  data work.group_1;
3    SET source.have;
4    where grp=1;
5  run;
6  data work.group_2;
7    SET source.have;
8    where grp=2;
9  run;
10 data work.group_3;
11   SET source.have;
12   where grp=3;
13 run;
```

code: a smarter approach

```
1  /* smarter */
2  data work.group_1
3      work.group_2
4      work.group_3;
5    SET source.have;
6    where grp in (1 2 3);
7    select(grp);
8      when(1) output work.group_1;
9      when(2) output work.group_2;
10     when(3) output work.group_3;
11     otherwise;
12 end;
13 run;
```

Though both achieve the same final effect, the processing efficiency differs. The left one just wraps a 4GL code snippet of the whole DATA step inside the loop's body, like this:

code: 1 loop

```
1  %do i = 1 %to 3;
2    data work.group_&i.;
3      SET source.have;
4      where grp=&i.;
5    run;
6  %end;
```

When we think twice, and realize (or rather remind ourselves) that macro language is not a 4GL-code generator but, it is a *whatever-text-you-want* generator(!), we can write it the "right" way, and easily reduce not needed I/O operations with three macro loops:

code: 3 loops

```
1  data %do i = 1 %to 3;
2      work.group_&i.
3      %end;;
4  SET source.have;
5  where grp in (%do i = 1 %to 3; &i. %end;);
6  select(grp);
7      %do i = 1 %to 3;
8          when(&i.) output work.group_&i.;
9      %end;
10     otherwise;
11 end;
12 run;
```

We get the same three data writes, but only one data read!

The trick with splitting data in a single read works well also for the "different splitting rules" case, even if result data sets have "overlapping" observations. In such cases, instead SELECT statement, a "good-old" IF-THEN conditional logic does the job:

code: the existing situation

```

1 data work.group_1;
2   SET source.have;
3   where grp in (1 2 3331 3332);
4 run;
5
6 data work.group_2;
7   SET source.have;
8   where "A"<id<"C" or "X"<id<"Z";
9 run;

```

code: a smarter approach

```

1 data work.group_1 work.group_2;
2   SET source.have;
3
4   if grp in (1 2 3331 3332) then
5     output work.group_1;
6   if "A"<id<"C" or "X"<id<"Z" then
7     output work.group_2;
8 run;

```

In the scenarios we have considered, reduction went, respectively, from 6I/O to 4I/O, and from 4I/O to 3I/O, but in the general case it goes: $2n\text{I/O} \rightarrow n+1\text{I/O}$, where n is the number of "groups/rules" we have to split the data.

"THE FINAL (BACKWARD) COUNTDOWN"

Sometimes we may be faced with a task that requires reversing our data order, i.e. reading a data set from bottom to top. In such a situation, instead of creating an artificial ordering variable n and resorting our newly created data by it, it is much better (and cheaper) to use the [POINT=](#) option in the [SET](#) statement.

code: the existing situation

```

1 /* sloppy */
2 data work.have;
3   SET source.have;
4   n + 1;
5 run;
6
7 proc sort data=work.have
8           out=work.have(drop=n);
9   by descending n;
10 run;

```

code: a smarter approach

```

1 /* smarter */
2 data work.have;
3   do point = nobs to 1 by -1;
4     SET source.have point=point
5                               nobs=nobs;
6   output;
7   end;
8 stop;
9 run;

```

Explicit [DO-LOOP](#) reads the data backward and gives the following reduction: 4I/O $\rightarrow$ 2I/O.

"DISENCHANTING"

Sometimes people may think that reading data from an external file with help of [INFILE](#) and [INPUT](#) statements requires "special" treatment. In fact there is nothing special in such DATA steps.

code: the existing situation

```

1 data work.have;
2   infile source;
3   INPUT grp id $ number obs;
4 run;
5
6 data work.have2;
7   SET work.have;
8   if grp NE 1;
9   number2 = number + 1000;
10 run;

```

code: a smarter approach

```

1 data work.have2;
2   infile source;
3   INPUT grp id $ number obs;
4   if grp NE 1;
5   number2 = number + 1000;
6 run;

```

There is no need to split our processing in two steps. We can read data directly from an external file, filter them out, derive new variables, or do calculations. Again: 4I/O $\rightarrow$ 2I/O.

"DISENCHANTING" AGAIN

The aim of this tip is also to "demystify", this time the `MERGE` statement. For this example let's assume that we were provided with those two data sets: `work.data1` (2457997 observations) and `work.data2` (2554654 observations). Our goal is, after merging those two data sets over variable `obs`, to aggregate variables `number2` and `number3`

code: 1st data set

```
1 data work.data1;
2   SET source.have;
3   where id ne "B";
4   number2=number*number;
5 run;
```

code: 2nd data set

```
1 data work.data2;
2   SET source.have;
3   where grp ne 17;
4   number3=number + 17;
5 run;
```

One could think that merging two data sets has to be executed exclusively, so that in the next step we could filter, summarize or use options like for example `END=`. Nothing could be further from the truth than that. The `MERGE` statement accepts: `END=`, `NOBS=`, or `INDSNAME=` options, same way the `SET` statement does. Furthermore, there is nothing wrong in not generating a data set (i.e. `data _null_;`) from a DATA step that merges data.

code: the existing situation

```
1 /* wandering around */
2 data work.interimStep;
3   MERGE work.data1 work.data2;
4   by obs;
5 run;
6 data OneObs;
7   set work.interimStep END=_E_;
8   sum + (number2 + number3);
9   if _E_;
10  put sum;
11 run;
```

code: a smarter approach

```
1 /* straight to the point */
2 data _null_;
3   MERGE work.data1 work.data2 END=_E_;
4   by obs;
5   sum + (number2 + number3);
6   if _E_;
7   put sum;
8 run;
```

This way we gain some I/O savings too, $4.01\text{I/O} \rightarrow 2\text{I/O}$ (the 0.01 is for the `OneObs`).

"DISENCHANTING" A BIT MORE

We have taken already away a bit of "holiness" from the `MERGE` statement, and we are not going to stop there. After playing a bit with SQL, a beginner SAS user might think that similarly to SQL's join the `MERGE` allows only two tables. And that is not true. The `MERGE` statement allows us to combine data from as many data sets as we want, of course as long as the merging is done over the same variable.

code: the existing situation

```
1 /* I/O waste */
2 data work.step1;
3   MERGE source.have work.data1;
4   by obs;
5 run;
6
7 data work.final;
8   MERGE work.step1 work.data2;
9   by obs;
10 run;
```

code: a smarter approach

```
1 /* smarter */
2 data work.final;
3   MERGE source.have
4         work.data1
5         work.data2;
6   by obs;
7 run;
```

So, instead of executing several steps with two data sets each we can simply do one with all data sets merged at once. This way we again gain some more I/O savings, $6\text{I/O} \rightarrow 4\text{I/O}$, and in general case $3n\text{I/O} \rightarrow n+2\text{I/O}$, in this case n is the number of data sets to merge, not counting the "main" data set.

"NICE PIECE OF ABSOLUTELY USELESS STEP"

Sometimes when we are focused on deriving next and next variables we may fall into a trap called by Aster and Seidman in [Aster & Seidman 1996] a "stream of consciousness" programming: one thought (derivation) one DATA step. Instead of producing a bunch of intermediate DATA steps, and through those DATA steps more data sets, and eventually re-merging those results (fortunately now we know only one MERGE data step is enough!) we can sometimes reduce I/O operations and combine processing into one data step.

code: the existing situation

```

1  /* stream of consciousness */
2  data work.step1;
3      SET source.have;
4      where number > 20;
5      number2=number*number;
6      keep obs number2;
7  run;
8  data work.step2;
9      SET source.have;
10     where id NE "B";
11     number3=number+17;
12     keep obs number3;
13 run;
14 /*data work.step3;
15     SET source.have;
16     where id > "X";
17     number4=number**42;
18     keep obs number4;
19 run;*/
20 data work.final;
21     MERGE source.have
22           work.step1
23           work.step2
24           /*work.step3*/
25     ;
26     by obs;
27 run;
```

code: a smarter approach

```

1  /* smart */
2  data work.final;
3      SET source.have;
4
5      if number > 20 then
6          number2=number*number;
7
8      if id NE "B" then
9          number3=number+17;
10
11     /*
12     if id > "X" then
13         number4=number**42;
14     */
15 run;
```

Sometimes such code has one more advantage: if we have to comment out a part of the code (e.g., a variable derivation or a conditional processing) we only need to comment out in one place of the code. Savings here are $8I/O \rightarrow 2I/O$, but the I/O reduction as a function of the derivations number n would have the following form $2+3nI/O \rightarrow 2I/O$.

"DOING MORE BY DOING LESS"

For this example, let's assume that we were provided with those two data sets: `work.one` (842279 observations) and `work.two` (2457997 observations). This time our goal is to create three data sets: *a*) the first with only those OBS values from `WORK.ONE` that does not show up in `WORK.TWO`, *b*) the second that has observations with a reversed relation, and *c*) the third that contains OBS values in the intersection of `WORK.ONE` and `WORK.TWO`.

code: 1st data set

```

1  data work.one;
2      set source.have;
3      where number > 20;
4  run;
```

code: 2nd data set

```

1  data work.two;
2      infile source;
3      INPUT grp id $ number obs;
4      if id ne "B";
5  run;
```

Classic PROC SQL approach can be replaced with shorter and clearer MERGE data step. Of course, this example quietly assumes that `WORK.ONE` and `WORK.TWO` are sorted by values of OBS.

code: the existing situation

```

1  /* "classic" */
2  proc sql;
3    create table only_in_one as
4      select one.*
5      FROM work.one
6      left join
7        work.two
8      on one.obs = two.obs
9      where two.obs is missing;
10   create table only_in_two as
11     select two.*
12     FROM work.one
13     right join
14       work.two
15     on one.obs = two.obs
16     where one.obs is missing;
17   create table one_and_two as
18     select one.*
19     FROM work.one
20     inner join
21       work.two
22     on one.obs = two.obs;
23 quit;

```

code: a smarter approach

```

1  /* smarter */
2  data
3    only_in_one
4    only_in_two
5    one_and_two;
6
7    MERGE work.one(in=o1)
8          work.two(in=t2);
9    by obs;
10
11   select;
12     when(    o1 and not t2)
13       output only_in_one;
14     when(not o1 and    t2)
15       output only_in_two;
16     when(    o1 and    t2)
17       output one_and_two;
18     otherwise;
19   end;
20 run;

```

The reduction here, 9I/O→5I/O, may not look so spectacular, but again we can see the superiority of the **MERGE** statement skills (of course, as long as we are staying in the world of 1-to-1 and 1-to-*N* merges).

"DOING BY NOT DOING"

Imagine you have two (in practice more) data sets (for simplicity, with excluding data) produced by the following snippet:

code: data

```

1  data work.one work.two;
2    set source.have;
3    if id NE "C" then output work.one;
4                      else output work.two;
5 run;

```

For example, one data set contains data for placebo subjects and the other for treatment. Or maybe the first data set has data from January, the second from February, the third from March, etc.

Our task is to concatenate data (i.e. to stack data from the second data sets behind data from the first, and so on) so we can read them sequentially. There is a non-zero probability that a DATA step with sequential **SET** statement reading both data sets and overwriting the first one would be the first choice. But for such a job **PROC APPEND**, moving only data from the second data set could work better (giving reduction 3I/O→1I/O, as the **BASE=** data set is not touched)

code: the existing situation

```

1  /* silly */
2  data work.one;
3    SET work.one work.two;
4  run;

```

code: a smarter approach

```

1  /* smarter */
2  proc append base=work.one
3             data=work.two;
4  run;

```

But since we are going to read the data in the subsequent step, maybe we do not have to move data at all?

By creating SAS view we can reduce I/O operations of the concatenation process to zero and materialize it only when we really need it, just like in this snippet:

```
code: a smarter approach
1  /* possibly even smarter */
2  data work.onetwo / view=work.onetwo;
3      SET work.one work.two;
4  run;
```

"NO RAW MACRO LOOPS", PART 1.

The macro language is a very practical way to write reusable code. Wrapping a snippet in a macro and have it for every call seems like a good idea, but sometimes, especially with the wrong application of %DO-LOOPS, it may slow your processing down (we have seen examples of this earlier). In his blog posts Rick Wicklin wrote about doing things "the BY way" multiple times (see [Wicklin 2012] or [Wicklin 2017]). The following examples will present work driven by a spirit similar to the one described by Rick, i.e., reducing unnecessary macro loops.

The first example is pretty easy to follow. Let's assume we have to calculate percentiles for our data in groups over variable id. A naive (but not so uncommon) approach would be to write a macro that extracts the list of id values and then execute a %DO-LOOP over the list. The loop, in each iteration, subsets the data set, then runs the PROC UNIVARIATE to calculate percentiles, and eventually use the PROC APPEND to combine results. In this, or any other, approach the PROC UNIVARIATE seems to be rather unavoidable, especially if we consider how very robust it is for calculating percentiles. But all those sub-setting DATA steps and all that appending can be avoided.

```
code: the existing situation
1  %macro slow();
2      proc sql;
3          select distinct id
4              into :id_list separated by " "
5              FROM source.have;
6          %let n = &SQLObs.;
7      quit;
8      %do i = 1 %to &n.;
9          %let id = %scan(&id_list.,&i.);
10         data subset;
11             SET source.have;
12             where id = "&id.";
13         run;
14         proc univariate data=subset;
15             var number;
16             output out=p
17                 pctlpre=P pctlpts=0 to 100;
18         run;
19         proc append base=pctl1 data=p;
20         run;
21     %end;
22 %mend slow;
23 %slow()
```

```
code: a smarter approach
1  /* smarter */
2  proc univariate data=source.have;
3      class id;
4      var number;
5      output out=pctl12
6          pctlpre=p pctlpts=0 to 100;
7  run;
```

Instead of macro looping we can do processing in groups, right? You could argue: "But to do the BY-group processing we still have to sort the data by the grouping variable". Well, that is true - but fortunately PROC UNIVARIATE provides the CLASS statement that saves our day. It is important to remember that the CLASS statement can specify at most two variables to be used to group the data into classification levels!

The I/O reduction here has the following form: $2+5n \text{ I/O} \rightarrow 2 \text{ I/O}$, the 2 on the left side is for SQL and for the first creation of `pctl1` data set. There is one more advantage of the [CLASS](#) approach, it not only reduces I/O operations, but also gives us values of the `id` variable in the final `pctl2` data set.

"NO RAW MACRO LOOPS", PART 2.

The second example is a bit more complicated. It is based on a true story, though any similarities...

Assuming we have an input data set have grouped by variable `grp`, though not necessarily sorted, a dictionary `dict` of separate parameters for each value of `grp` variable, and we want to execute a parameterized calculation of a new variable, let's say something like the following:

$$x = \text{function}(\alpha_{grp} \times var_1, \beta_{grp} \times var_2, \gamma_{grp} \times var_3) + Constant_{grp};$$

where α , β , γ , and *Constant* are group dependent, i.e., their values changes depending on the `grp` variable values (that is why we are picking those values from the dictionary). And the `function(...)` can be considered as a group of conditional logic expressions composed of SAS statements and SAS functions (including [PROC FCMP](#) user defined functions too).

In our example, we are going to use the [RAND\(\)](#) function and the following snippet:

```
code: formula to evaluate
1 x = rand("Distribution", ParA, ParB) + Threshold;
```

The [RAND\(\)](#) function accepts various number of parameters - depending on the distribution used - so it perfectly fits into our need for a "pretender-of-some-business-logic", by being semantically simple enough and syntactically complicated enough at the same time. The dictionary is a SAS data set `source.dict` with the following (long) structure, and for brevity we have it prepared only for selected values of `id` variable from the `source.have`. Here is the code for the dictionary:

```
code: dictionary
1 data source.dict;
2   array distr[9] $ 4 ("BERN" "CAUC" "EXPO" "F" "GAMM" "INTE" "LOGI" "T" "UNIF");
3   array Npar[9] $ 4 ("1" "0" "1" "2" "1" "2" "2" "1" "2" );
4   array ParA[9] $ 4 ("0.5" " " "1" "3" "42" "0" "0" "1" "0" );
5   array ParB[9] $ 4 ( " " " " " " "4" " " "10" "1" " " "1" );
6
7   do i = 1 to 9;
8     length id $ 8;
9
10    id = char(distr[i],1);
11    par="Distribution";
12    val=distr[i]; output;
13    par="Npar";
14    val=Npar[i]; output;
15    par="ParA";
16    val=ParA[i];
17    if val NE "" then output;
18    par="ParB";
19    val=ParB[i];
20    if val NE "" then output;
21    par="Threshold";
22    val = put(i/10,3.1); output;
23  end;
24  keep id par val;
25 run;
```

and here is the output of the dictionary displayed in two columns tabular form:

Table 2: Data set SOURCE.DICT

id	par	val	id	par	val
B	Distribution	BERN	I	Distribution	INTE
B	Npar	1	I	Npar	2
B	ParA	0.5	I	ParA	0
B	Threshold	0.1	I	ParB	10
C	Distribution	CAUC	I	Threshold	0.6
C	Npar	0	L	Distribution	LOGI
C	Threshold	0.2	L	Npar	2
E	Distribution	EXPO	L	ParA	0
E	Npar	1	L	ParB	1
E	ParA	1	L	Threshold	0.7
E	Threshold	0.3	T	Distribution	T
F	Distribution	F	T	Npar	1
F	Npar	2	T	ParA	1
F	ParA	3	T	Threshold	0.8
F	ParB	4	U	Distribution	UNIF
F	Threshold	0.4	U	Npar	2
G	Distribution	GAMM	U	ParA	0
G	Npar	1	U	ParB	1
G	ParA	42	U	Threshold	0.9
G	Threshold	0.5			

The solution I witnessed (left column below) was based on the following logic:

- (1) start the process and create a list of unique values of the `id` variable from the dictionary [lines 2-8],
- (2) loop over the list and for *each* value from the list [lines 10,11,28]:
 - (a) sub-set the dictionary for given `id` value and produce a bunch of macro variables named like parameters and with their corresponding values as values [lines 12-16],
 - (b) sub-set the have data set for given `id` value and for every given set of macro variables execute calculation [lines 17-24],
 - (c) stack the result by appending it to previous iteration's result [lines 25-27]
- (3) end the process [line 29].

The number of I/O operations here is $5nI/O$, where n is the number of the `id` variable values (9 in our case), plus one I/O for SQL and one for the first execution of the `PROC APPEND`. So, in the example here the number of inputs and outputs ends to be $47I/O$.

code: the existing situation

```

1  /* slow */
2  %macro slow2();
3  proc sql;
4      select distinct id
5      into :id_list separated by " "
6      FROM source.dict;
7      %let n = &SQLobs.;
8  quit;
9
10 %do i = 1 %to &n.;
11 %let id = %scan(&id_list.,&i.);
12 data dict;
13     SET source.dict;
14     where id = "&id.";
15     call symputX(par,val,"L");
16 run;
17 data tmp;
18     SET source.have;
19     where id = "&id.";
20     x = rand("&Distribution."
21     %if &Npar.>0 %then %do;,&ParA.%end;
22     %if &Npar.>1 %then %do;,&ParB.%end;
23     ) + &Threshold.;
24 run;
25 proc append base=work.result1
26             data=tmp;
27 run;
28 %end;
29 %mend slow2;
30
31 options Mprint;
32 %slow2()

```

code: a smarter approach

```

1  /* smarter */
2  proc transpose data=source.dict
3              out=dict(drop=_:);
4      by id;
5      id par;
6      var val;
7  run;
8
9  PROC SQL;
10     create table work.result2 as
11     select h.*,
12           case
13             when d.Npar="2" then
14               rand(Distribution
15                   ,input(ParA, best.)
16                   ,input(ParB, best.))
17             when d.Npar="1" then
18               rand(Distribution
19                   ,input(ParA, best.))
20             else
21               rand(Distribution)
22             end + input(Threshold, best.) as x
23     from
24         source.have as h
25         inner join
26         dict as d
27         on h.id = d.id;
28 QUIT;

```

An alternative to the first approach is based only on 5I/O, the first two are for the **PROC TRANSPOSE** which transforms the dictionary from long into wide format (see table printout below). The other three are for the **PROC SQL** where the **CASE-WHEN-END** conditional expression combined with several calls to the **INPUT()** function does the heavy lifting. Remembering that n is the number of the `id` variable values we eventually end up with $2+5nI/O \rightarrow 5I/O$ reduction.

Table 3: Data set with dictionary after transposition

id	Distribution	Npar	ParA	ParB	Threshold	id	Distribution	Npar	ParA	ParB	Threshold
B	BERN	1	0.5		0.1	I	INTE	2	0	10	0.6
C	CAUC	0			0.2	L	LOGI	2	0	1	0.7
E	EXPO	1	1		0.3	T	T	1	1		0.8
F	F	2	3	4	0.4	U	UNIF	2	0	1	0.9
G	GAMM	1	42		0.5						

"#HASH TABLE FOR HELP"

Classical problem of re-merging summary statistics back with the original data (often seen as a note after **PROC SQL** execution) usually is done by a) sorting data by grouping variables, b) doing some **BY**-group processing to calculate statistics, and c) merging created summary statistics data set with the source data.

code: the existing situation	code: a smarter approach
1 /* standard */	1 /* with hash tables */
2 proc sort data=source.have	2 data work.want;
3 out=work.have;	3 dcl hash S(ordered:"A");
4 by id;	4 S.defineKey("id");
5 run;	5 S.defineData("id","maxN","minN");
6 data work.aggr;	6 S.defineDone();
7 SET work.have;	7 declare hiter I("S");
8 by id;	8
9 if first.id then	9 dcl hash D(multidata:"Y",ordered:"A");
10 do;	10 D.defineKey("id");
11 maxN=number;	11 D.defineData("grp","id","number","obs");
12 minN=number;	12 D.defineDone();
13 end;	13
14 maxN = maxN max number;	14 do until(_E_);
15 minN = minN min number;	15 SET source.have END=_E_;
16 if last.id then	16 D.add();
17 do;	17
18 range = maxN-minN;	18 shift = S.find();
19 output;	19 maxN = max(maxN,number);
20 end;	20 minN = min(minN,number);
21 keep id range;	21 S.replace();
22 retain maxN minN;	22 end;
23 run;	23
24 data work.want;	24 do while(0=I.next());
25 merge work.have work.aggr;	25 do while(D.do_over()=0);
26 by id;	26 shift = number/(maxN-minN);
27 shift = number/range;	27 output;
28 drop range;	28 end;
29 run;	29 end;
	30 stop;
	31 drop maxN minN;
	32 run;

Use of hash tables can give us pretty nice I/O operations reduction if we use them both for storing summary statistics and the source data. Like in the example above hash table S stores minimum and maximum for **id** groups, and hash table D stores all data from the input data set. Use of the explicit **DO-UNTIL** loop allows us to read-in source data and calculate statistics at the same time with only one data pass. Then in a double **DO-WHILE** loop we combine data and the aggregate to produce final result. Original seven I/O operations reduce to two giving: $7I/O \rightarrow 2I/O$. Of course, we realize that reduction of I/O operations here has a price of higher memory consumption. But often that price is worth paying, especially that hash tables allow us to reduce or stop I/O intensive data sorting. The best hash tables source of knowledge is [Dorfman & Henderson 2018] book.

"NO RAW MACRO LOOPS", PART 2. - REVISITED

If consider our just gained hash tables experience and we add some "good old" arrays to it, we can go even one step further with I/O operations reduction.

```

code: hash tables approach
1 data work.result3;
2   if 1=_N_ then do;
3       declare hash D();                declare hash I();
4       D.defineKey("id", "par");          I.defineKey("id");
5       D.defineData("val");
6       D.defineDone();                    I.defineDone();
7       do until(_E_);
8           SET source.dict end=_E_; /* one data reading for 2 hash tables */
9           if D.add() then stop;      /* "quality stop" if dict has doubles */
10          I.replace();
11      end;
12  end;
13
14  SET source.have;
15  by ID notsorted;
16  if 0=I.check(); /* to get "inner join" effect */
17
18  if first.ID then do; /* get parameters */
19      array Dict[*] $ Distribution Npar ParA ParB Threshold; retain;
20      call missing(of Dict[*]);
21      do j=1 to Dim(Dict);
22          if 0=D.find(key=id,key=vname(Dict[j])) then Dict[j] = val;
23      end;
24  end;
25  select(Npar);
26      when("2") x = rand(Distribution,input(ParA, best.),input(ParB, best.));
27      when("1") x = rand(Distribution,input(ParA, best.));
28      otherwise x = rand(Distribution);
29  end;
30  x = x + input(Threshold, best.);
31
32  drop par val j Distribution Npar ParA ParB Threshold;
33 run;

```

We end up with $2+5nI/O \rightarrow 3I/O$ instead of $5I/O$. But this example cannot be left without a comment. There is a fair chance that the **PROC TRANSPOSE** + **PROC SQL**'s approach, though it has two more I/O operations, will win here (in terms of "wall clock" time) especially when the dictionary is small. This is because **PROC SQL** is designed smart and if it figures out the dictionary table is small enough to fit in-memory, it will use so called "hash join" method to combine data faster. If you run **PROC SQL** from our example with (undocumented) option: `__METHOD`⁵, it will print something similar to this:

```

the log - SQL and __method
1 NOTE: SQL execution methods chosen are:
2   sqxcrta /* Create table */
3   sqxjhsh /* Hash join operation */
4   sqxsrc( SOURCE.HAVE(alias = H) ) /* Source rows from table */
5   sqxfil /* Rows filtration */
6   sqxsrc( WORK.DICT(alias = D) )

```

⁵See [Lavery 2005] to learn the `__method` secrets

Of course, if the program logic requires data step solutions, e.g., arrays, this one is clearly the winner. The example also shows that: a single data read can populate multiple hash tables, what is an extremely convenient trick. Further more we did not have to sort the `source.have` data set to do the table join. The bottom line here is that, in general, hash tables give us a bunch of useful I/O-saving features.

"ONE STEP TO RULE THEM ALL"

Sometimes there is a need to create a "template" data set that contains all possible combinations of categorical variables. A "standard" `PROC SQL` approach, of creating a few data sets with unique values, crowned by SQL's Cartesian product of those, is often the choice. But, in such a case, the much more I/O-efficient is a data step.

code: SQL 1	code: hash table
<pre> 1 proc sql; 2 create table 3 work.sql_dist_number as 4 select 5 distinct number 6 from 7 source.have; 8 9 create table 10 work.sql_dist_id as 11 select 12 distinct id 13 from 14 source.have; 15 16 create table 17 work.sql_dist_grp as 18 select 19 distinct grp 20 from 21 source.have; 22 23 create table 24 all_possible_crosses_S as 25 select * 26 from 27 work.sql_dist_number, 28 work.sql_dist_id, 29 work.sql_dist_grp; 30 quit; </pre>	<pre> 1 data all_possible_crosses_H; 2 declare hash H1(ordered:"A"); 3 H1.defineKey("number"); 4 H1.defineDone(); 5 declare hiter i1("H1"); 6 7 declare hash H2(ordered:"A"); 8 H2.defineKey("id"); 9 H2.defineDone(); 10 declare hiter i2("H2"); 11 12 declare hash H3(ordered:"A"); 13 H3.defineKey("grp"); 14 H3.defineDone(); 15 declare hiter i3("H3"); 16 17 do until(_E_); 18 set source.have(keep=number id grp) 19 end=_E_; 20 H1.ref(); 21 H2.ref(); 22 H3.ref(); 23 end; 24 25 do while(0=i1.next()); 26 do while(0=i2.next()); 27 do while(0=i3.next()); 28 output; 29 end; 30 end; 31 end; 32 stop; 33 run; </pre>

Use of hash tables to store unique values allows us to traverse data only once and the result is produced directly. The I/O reduction, in the example here, goes from 10I/O (read, write, and again read - for each variable, plus one for result data set) to just 2I/O (one for reading `source.have` and the second for writing the result data set). In general case, having n variables selected, it goes $1+3n\text{I/O} \rightarrow 2\text{I/O}$.

An inquisitive reader may doubt the efficiency, and suspect we are pulling the wool over one's eyes, and even support the doubts with the following (shorter) SQL query that seems to reduce part of those I/O operations:

```

code: SQL 2
1 proc sql _method;
2   create table all_possible_cross_S2 as
3   select *
4   from (select distinct number from source.have)
5        , (select distinct id      from source.have)
6        , (select distinct grp     from source.have)
7   ;
8 quit;

```

but, when the execution is closely inspected with the [_METHOD](#) option, we can see the `source.have` data set is read 3 times, and the Cartesian join (`sqxjsl`) needs some utility storage too. We have at least four evident I/O operations (including write of the result data set) and some unknown number of "under the hood" utility data wranglings.

```

the log - SQL and _method
1 NOTE: SQL execution methods chosen are:
2   sqxcrt
3   sqxjsl /* Step loop join (Cartesian) */
4   sqxunqs
5   sqxsrc( SOURCE.HAVE )
6   sqxjsl
7   sqxunqs /* Select unique values */
8   sqxsrc( SOURCE.HAVE )
9   sqxunqs
10  sqxsrc( SOURCE.HAVE )
11 NOTE: SAS threaded sort was used.

```

"I HAVE SEEN THIS BEFORE"

Yet another example of an "I-have-seen-this-before" situation is the case where we have to add (combine) some new data into our base table, furthermore:

- (1) the data are provided in separate data sets (no Excels, CSVs, etc., just for simplicity),
- (2) the data we want to combine must be merged over different variables, and
- (3) the final data set should maintain the original observations order.

For simplicity, data sets with additional data have only two variables each, one for "joining key" and the other for data, e.g., `number` and `number_text`, etc.

Usually, two approaches are used to solve the task. The first one combines merging and sorting intertwined together, the second approach, as one can expect, is SQL's left joining.

The "merge and sort" solution starts with ordering all additional data sets by their merging variables. Next, the base data set is sorted by selected variable and then merged with additional data. The process is repeated for all variables, and eventually the final data set is resorted to restore the original observations order. In our example setup the process involves 23 I/O. The alternative, with [PROC SQL](#), seems to be less I/O intense, at least at the first glance. Unfortunately, the [_METHOD](#) option reveals it is not just a '94 "Four Reads and a Write" comedy but it is (again) much more complicated. The LOG shows that before joining, all data sets are resorted, and since the merge join is used, the temporary data have to be stored in utility files. Of course the final re-sorting takes its toll too.

code: Merge and sort

```

1 proc sort
2   data = number_data
3   out = number_data_sort;
4   by number;
5 run;
6 proc sort
7   data = id_data
8   out = id_data_sort;
9   by id;
10 run;
11 proc sort
12   data = grp_data
13   out = grp_data_sort;
14   by grp;
15 run;
16 proc sort
17   data = source.have
18   out = have_by_grp;
19   by grp;
20 run;
21 data have_1;
22   merge
23     have_by_grp
24     grp_data_sort;
25   by grp;
26 run;
27 proc sort
28   data = have_1
29   out = have_by_id;
30   by id;
31 run;
32 data have_2;
33   merge
34     have_by_id
35     id_data_sort;
36   by id;
37 run;
38 proc sort
39   data = have_2
40   out = have_by_number;
41   by number;
42 run;
43 data have_3;
44   merge
45     have_by_number
46     number_data_sort;
47   by number;
48 run;
49 proc sort
50   data = have_3
51   out = combined_DS;
52   by obs;
53 run;

```

code: SQL

```

1 proc sql _method;
2   create table combined_SQL as
3   select
4     h.*,
5     n.number_text,
6     i.id_text,
7     g.grp_text
8   from
9     source.have as h
10  left join
11    number_data as n
12    on h.number=n.number
13
14  left join
15    id_data as i
16    on h.id=i.id
17
18  left join
19    grp_data as g
20    on h.grp=g.grp
21
22  order by obs;
23 quit;

```

The LOG unveils it all. Even if we assume that the `sqxsrc` and `sqxsort` utilize I/O operations together we can count something around eleven I/O in the process.

```

_____ the log - SQL and _method _____
1 NOTE: SQL execution methods chosen are:
2   sqxcrt
3     sqxsort
4       sqxjm
5         sqxsort
6           sqxsrc( WORK.GRP_DATA(alias = G) )
7         sqxsort
8           sqxjm
9         sqxsort
10          sqxsrc( WORK.ID_DATA(alias = I) )
11        sqxsort
12          sqxjm /* Merge join operation */
13        sqxsort
14          sqxsrc( WORK.NUMBER_DATA(alias = N) )
15        sqxsort /* Sort operation */
16          sqxsrc( SOURCE.HAVE(alias = H) )
17 NOTE: SAS threaded sort was used.

```

This time, with a little help from our hash table friends, the "optimized" solution proposed here uses only 5I/O: three data reads to populate hash tables: H1, H2, and H3 with data, one run over the `source.have` base data set, and eventually for writing the final `combined_hash` data set, with no need for any re-sorting at all (which also reduces Operating System Memory cost by 15 times!). The I/O costs reduction goes 23I/O→5I/O. And just to be clear, code in line number 2 does not read any data during the execution phase. The `SET` statement is there just for metadata extraction in the compilation phase and is used just to minimize variables metadata interaction (simple `LENGTH` would work too).

```

_____ code: hash tables approach _____
1 data combined_hash;
2   if 0 then set work.number_data work.id_data work.grp_data;
3
4   declare hash H1(dataset:"work.number_data");
5   H1.defineKey("number");
6   H1.defineData("number_text");
7   H1.defineDone();
8
9   declare hash H2(dataset:"work.id_data");
10  H2.defineKey("id");
11  H2.defineData("id_text");
12  H2.defineDone();
13
14  declare hash H3(dataset:"work.grp_data");
15  H3.defineKey("grp");
16  H3.defineData("grp_text");
17  H3.defineDone();
18
19  do until(_E_);
20    set source.have
21      end=_E_;
22    if H1.find() then number_text="";
23    if H2.find() then id_text="";
24    if H3.find() then grp_text="";
25    output;
26  end;
27 stop;
28 run;

```

CONCLUSION

In this article we tried to bring beginner SAS programmers' attention to mastering I/O operations efficiency by "pruning" redundant [SET](#), [MERGE](#), [INPUT](#), or any other I/O related statements. Undoubtedly, we can state that SAS is a very syntax rich and flexible language, and it takes quite some effort to master it. Even tough examples and snippets presented here may seem to be "repetitions of obvious things", like the 1947 Nobel Prize in Literature winner, French, André Gide⁶ said:

"Toutes choses sont dites déjà; mais comme personne n'écoute, il faut toujours recommencer."

("Everything has been said before, but since nobody listens we have to keep going back and beginning all over again.")

From time to time, new, inexperienced SAS programmers show up in the SAS community seeking ways to make their code better, thus we strongly believe that some of those "seems to be obvious" programming rules should be presented again and again. Writing good and efficient code is a non-trivial task. The article, by use of the inevitable part of the education process - repetition, tries to help to get to the stage when "we know how to write good code" easier, without need to first wander around through meanders and off-roads of ugly or (probably more often) inefficient code. We hope that the topics presented in the article will help to make life of "junior" SAS programmers easier and allow them to shorten the "from bad to good code" path they are walking. Hopefully, programmers with more experience will also find the paper useful.

Thus, to make your program efficient, let's repeat after **[Aster & Seidman 1996]** one more time: *"consider every [SET](#) ([MERGE](#), [INPUT](#), and any other I/O related) statement with suspicion!"*

The End

REFERENCES

- [Dijkstra 1968] Edsger Dijkstra, "Go To Statement Considered Harmful", Communications of the ACM. 11 (3): 147-148, 1968,
<https://doi.org/10.1145%2F362929.362947>,
<https://homepages.cwi.nl/~storm/teaching/reader/Dijkstra68.pdf>
- [Aster & Seidman 1996] Rick Aster, Rhena Seidman, "Professional SAS Programming Secrets", McGraw-Hill Inc., US; 2nd Updated Edition, 1996
- [Lavery 2005] Russ Lavery, "The SQL Optimizer Project: _Method and _Tree in SAS 9.1", SUGI 30 Proceedings, 2005,
<https://support.sas.com/resources/papers/proceedings/proceedings/sugi30/101-30.pdf>
- [Raithel 2010 & 2018] Michael A. Raithel, "PROC DATASETS; The Swiss Army Knife of SAS Procedures", WUSS Proceedings, 2018, https://www.lexjansen.com/wuss/2018/144_Final_Paper_PDF.pdf
SGF Proceedings, 2010, <https://support.sas.com/resources/papers/proceedings10/138-2010.pdf>
- [Wicklin 2012] Rick Wicklin, The DO-LOOP blog, 2012 "Simulation in SAS: The slow way or the BY way",
<https://blogs.sas.com/content/iml/2012/07/18/simulation-in-sas-the-slow-way-or-the-by-way.html>
- [Wicklin 2017] Rick Wicklin, The DO-LOOP blog, 2017 "An easy way to run thousands of regressions in SAS",
<https://blogs.sas.com/content/iml/2017/02/13/run-1000-regressions.html>
- [Dorfman & Henderson 2018] Paul M. Dorfman, Don Henderson, "Data Management Solutions Using SAS Hash Table Operations: A Business Intelligence Case Study", SAS Institute Press, 2018

ACKNOWLEDGMENTS

We would like to thank Filip Kulon, Troy Martin Hughes, and Louise Hadden! We are utterly grateful for their help in "polishing" this text.

⁶Quote taken from: https://en.wikiquote.org/wiki/Andr%C3%A9_Gide

CONTACT INFORMATION

Your comments and questions are valued and encouraged!

Contact Bart at one of the following e-mail addresses:

yabwon@gmail.com or bartosz.jablonski@pw.edu.pl

or via the following LinkedIn profile: www.linkedin.com/in/yabwon or at the communities.sas.com by mentioning @yabwon.

SAS and all other SAS Institute Inc. product or service names are registered trademarks or trademarks of SAS Institute Inc. in the USA and other countries. ® indicates USA registration.

Other brand and product names are trademarks of their respective companies.

Appendix A - code coloring guide

The best experience for reading this article is in color and the following convention is used:

- The code snippets use the following coloring convention:

```
code: is surrounded by a black frame
1 In general we use black ink for the code but:
2 - for reading clarity we sometimes mark code in orange ink,
3 - and comments pertaining to code are in a bluish ink for easier reading.
```

- The LOG uses the following coloring convention:

```
the log - is surrounded by a blueish frame
1 The source code and general log text are blueish.
2 Log NOTES are green.
3 Log WARNINGS are violet.
4 Log ERRORS are red.
5 Log text generated by the user is purple.
```

INDEX

function	INDSNAME=, 7	BY, 4, 10, 14
INPUT(), 13	NOBS=, 7	CASE-WHEN-END, 13
RAND(), 11	OUT=, 2, 4	CLASS, 10, 11
I/O(input/output), 1–20	POINT=, 6	DATA, 2, 4
macro-statement	procedure	DO-LOOP, 6
%DO-LOOP, 5, 10	APPEND, 9, 10, 12	DO-UNTIL, 14
%IF-THEN-ELSE, 13	DATASETS, 3, 4	DO-WHILE, 14
%LET, 10	FCMP, 11	GOTO, 2
option	SORT, 4	IF-THEN, 5
_METHOD, 15, 17	SQL, 8, 13–17	INFILE, 6
BASE=, 9	TRANSPOSE, 13, 15	INPUT, 1, 2, 6, 20
DATA=, 1, 2, 4	UNIVARIATE, 10	LENGTH, 19
END=, 7	statement	MERGE, 1–3, 7–9, 20
		SELECT, 5
		SET, 1–4, 6, 7, 9, 19, 20