

MULTIPLE INDICATOR MONITORING (MIM) DATA INSTRUCTIONS GUIDE



Timothy A. Burton, Fisheries Biologist/Hydrologist
Retired USDI Bureau of Land Management/USDA Forest Service
Boise, Idaho

Steven J. Smith, Riparian Ecologist/Rangeland Management Specialist
Retired USDI Bureau of Land Management
Prineville, Oregon

Mark A. Gonzalez, Riparian/Wetland Ecologist (Soils)
USDI Bureau of Land Management
Prineville, Oregon

2024
(version 3, revised April 4, 2025)

Table of Contents

MULTIPLE INDICATOR MONITORING (MIM) DATA INSTRUCTIONS GUIDE.....	1
I. INTRODUCTION.....	6
II. DATA ENTRY AND ANALYSIS MODULES.....	6
Which module to use	6
A. DATA ENTRY MODULE	7
DATA ENTRY MODULE - WORKBOOK TABS:	8
1. "INSTRUCTIONS" TAB	9
2. "HEADER" TAB	10
3. "DMA" TAB	17
4. "SUBSTR" TAB	21
5. "THAL" TAB	22
6. "COMMENTS" TAB.....	23
7. "GRAPHS" TAB	24
8. "CODES" TAB.....	25
9. "CALCS" TAB	25
10. "PLANTS" TAB	25
11. "KEYSP" TAB.....	27
12. "SPATIAL" TAB	28
B. DATA ANALYSIS MODULE	30
DATA ANALYSIS MODULE - WORKBOOK TABS:	32
1. "INSTRUCTIONS" TAB	32
2. "HEADER" TAB	35
3. "PHOTOS" TAB	35
4. "DMA" TAB	36
5. "SUBSTR" TAB	37
6. "THAL" TAB	37
7. "COMMENTS" TAB.....	37
8. "DATA SUMMARY" TAB	38
9. "STHT" TAB	41

10.	"SPATIAL" TAB	44
11.	"CORREL" TAB:	51
12.	"GRAPHS" TAB	51
13.	"CODES" TAB.....	53
14.	"EXPORT" TAB.....	54
15.	"CALCS" TAB.....	54
16.	"PLANTS" TAB	54
17.	"KEYSP" TAB.....	55
18.	"PFC" TAB.....	56
19.	"BOOT" TAB:	56
	Metrics Calculated in the Data Analysis Module	57
	Uploading and Correcting Historic Data in the Data Analysis Module	61
C.	STATISTICAL ANALYSIS MODULE - EXCEL TABS:.....	64
1.	"INSTRUCTIONS" TAB:	64
2.	"GETDATA" TAB:	64
3.	"COMP" TAB:	64
4.	"NORMAL PLOTS" TAB:.....	65
5.	"HISTOGRAMS" TAB:	66
6.	"SPATIAL" TAB:	67
7.	"SHTERM" TAB:.....	67
8.	"CHANNEL" TAB:.....	68
9.	"VEG" TAB:.....	69
10.	"TTEST" TAB:	70
11.	"MANNW" TAB:	71
12.	"CHISQ" TAB:	72
13.	"REPORTS" TAB:.....	74
	REFERENCES.....	74
III.	TESTING PRECISION AND DETECTING CHANGE	75
A.	INTRODUCTION.....	75
B.	TESTING THE PROTOCOL.....	75
1.	Precision	76
2.	The confidence interval	77
3.	The coefficient of variation.....	79

4.	Margin of error (ME):	80
5.	Testing Observer variation	81
6.	Field testing the MIM protocol.....	82
7.	Testing sample size versus margin of error	85
8.	Estimating sample size.....	89
9.	Testing observer variation	92
10.	Displaying results.....	93
APPENDIX A – SPATIAL AUTOCORRELATION ASSESSMENT OF THE MULTIPLE INDICATOR MONITORING (MIM) PROTOCOL		95
Methods.....		95
Results.....		101
Summary		102
Recommendations		110
Conclusions		113
REFERENCES		122
APPENDIX B - OBTAINING THE CONFIDENCE INTERVAL FOR NON-NORMALLY DISTRIBUTED DATA USING BOOTSTRAPPING.....		123
METHODS.....		124
FINDINGS.....		125
CONCLUSIONS.....		127
APPENDIX C – BLANK DATA FORMS.....		130
APPENDIX D – MIM DATA ANALYSIS EXAMPLES.....		136
PART 1: Non-normal distribution.....		136
PART 2: Spatial autocorrelation.		139
STEP-BY-STEP DATA ANALYSIS PROCEDURE		144
APPENDIX E – DMA MODIFICATION CALIBRATION.....		146
APPENDIX F. Simplified explanations of statistics used in the Data Analysis Module.....		151
A.	BOOTSTRAP ANALYSIS	151
B.	Spatial Autocorrelation	153
	Step-by-Step Explanation	155
	All of this is done for you in the Data Analysis Module – Spatial tab.....	156
	Example	156
	Why This Matters	157

C. Advantages and disadvantages of using the 95% confidence interval for observer variation..... 158

References 160

VERSION 3: April 2025 This is the third version of the 2024 publication. Updates will be made from time-to-time and will be described here. This version has added Appendix F to discuss statistical analysis in more detail.

VERSION 2: January 2025 This is the second version of the 2024 publication. This version has an expanded section on stubble height analysis, bootstrapping statistics, and spatial autocorrelation.

I. INTRODUCTION

This Multiple Indicator Monitoring Data Instructions Guide is designed to be used as a companion to the Multiple Indicator Monitoring (MIM) of Stream Channels and Streamside Vegetation Technical Reference (1737-23 version 2). During drafting of the updated MIM technical reference, the authors decided to move the content on data entry, analysis, interpretation, summary metrics, statistics, etc. from the 2011 TR and place it into a separate companion document – this Data Instructions Guide. The Data Instructions Guide is intended to provide users with the information necessary to collect and correct data in preparation for analysis, interpretation, and evaluation. This guide is an online resource that will be frequently updated; thereby enabling the authors to make timely modifications to data-related instructions, data processing, analysis tools, studies, etc. This also allows the MIM technical reference to be a streamlined field document focused on stream reach/riparian complex stratification, DMA selection, and instructions for collecting data on the 10 MIM indicators.

II. DATA ENTRY AND ANALYSIS MODULES

Several data modules have been developed using a Microsoft Office Excel spreadsheet format. The data modules facilitate data entry in the field, data correction and analysis in the office, data storage for future uses, and statistical analyses for various applications including planning, reporting, and management. The data modules are often updated to facilitate minor refinements and improvements identified by the users, and additions are often made to facilitate new information useful to MIM analyses, but the basic format and function of the data modules remains constant and has been so through the years.

Which module to use

There are several data modules applicable to the MIM protocol. Data Entry Modules facilitate data entry in the field. Data Analysis Modules are for correcting and analyzing field data and for data storage. The Statistical Analysis Module is for comparing the data from more than one

designated monitoring area (DMA) as well as from more than one monitoring period at a single DMA to estimate conditions and trend statistically and can also be used to help determine if current management is leading to (i.e., trending toward) desired conditions or achieving riparian management objectives.

A. DATA ENTRY MODULE

There are four Data Entry Modules. The first three listed in figure 1 below, are to collect all 10 MIM indicators on various kinds of field devices. The fourth is collect data for just the short-term indicators of use (i.e. ungulate use), thus the name “livestock use” is part of the name to uniquely identify them. The basic “Data Entry Module 2024” or the “Data Entry Module Livestock Use 2024” are the classic versions that have been used in the past and have not changed. The new modules “Drop-down BOXES” and “for IPAD” are designed for specific field devices. The “Drop-down BOXES” version is most useful for devices in which the screen works best with a finger-sized stylus and where the normal dropdowns within data cells are too small to make it easy to select items from the drop-down list. In these modules, EXCEL combo boxes have been added to make it easier to select items because the boxes and drop-down arrows are much larger. The “for IPAD” version is designed for those devices that do not support EXCEL macros, such as IPAD or other IOS-based tablets. These devices contain no macros, but they do support drop-down lists (but not combo boxes which operate with macros).

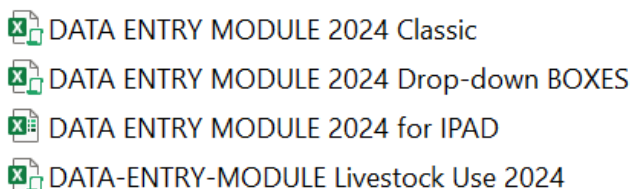


Figure 1. Data Entry Modules

The macro-supported versions (.xlsm) can be used for entering data using a device that supports but does not require “touch screen” input. In all versions, individual cells can be expanded manually using Excel’s expansion control in the lower right-hand corner of the screen. Drop-down lists are provided to minimize typing and allow the user to select items in predefined lists. Drop-down lists include the plant lists and key species lists that have been identified for the DMA, as well as for other indicators selected for monitoring. This makes manual data entry much more convenient than in the Data Analysis Module. **This is why the Data Analysis Module should not be used for entering data in the field (as some users have suggested doing in the past).** Also, the Data Entry Module includes a macro for automatically

populating the species lists from prior field data collection events at the same DMA. The user migrates to a previous Data Analysis file and selects it. The module then uploads the plant and key species lists directly into the Data Entry Module. In the versions that do not support macros, the user must enter the plant lists manually. A good way to do that is to simply copy the plant list from the previous Data Analysis Module (select the range of plant codes and press ctrl c) and then paste the contents into columns M, N, and O of the “PLANTS” tab (press ctrl V).

The **Data Entry Module - Livestock Use** is used specifically for entering short-term monitoring indicators (stubble height, streambank alteration, and woody riparian species use), plus one long-term indicator - streambank stability and cover, which is often reflective of the effects of the short-term indicators.

Each Data Entry Module has an “Instructions” tab, as shown in figure 2.

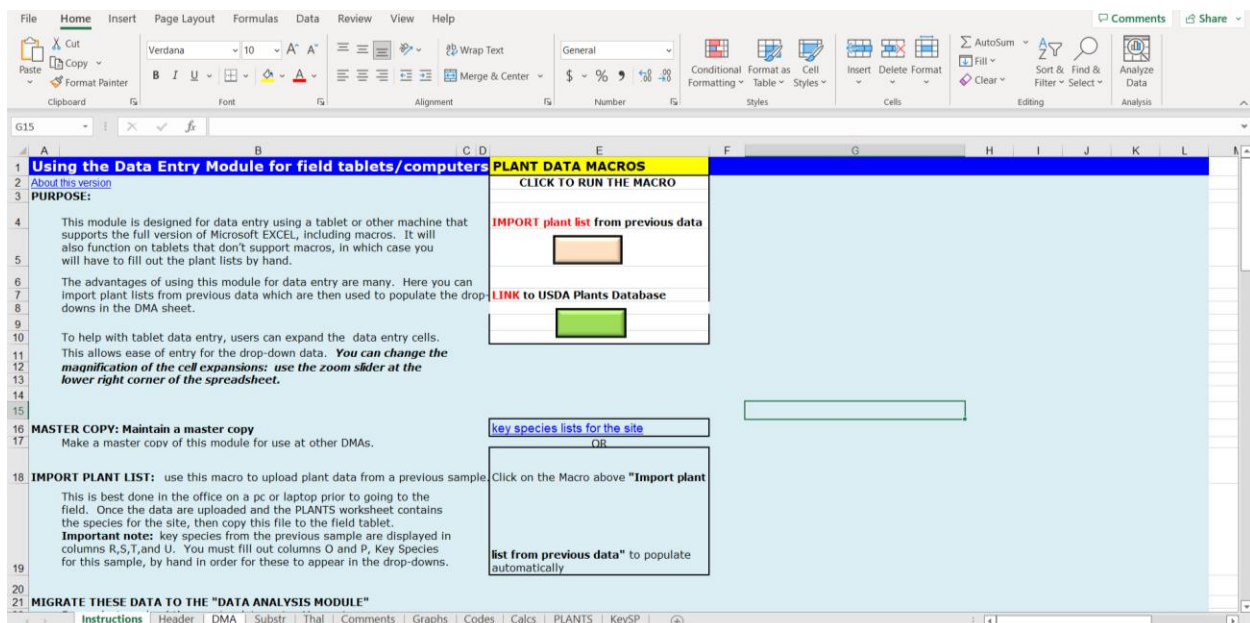


Figure 2. “Instructions” tab for the Data Entry module showing the macro buttons. These buttons are absent in the version that does not support macros.

In addition to basic instructions for operating the module, this tab contains macros that allow the user to import plant lists from previously collected data and to access the USDA PLANTS Database online. Access to the latter requires that the user is connected to the internet. Note the tabs at the bottom of the graphic, all of which are described in the following section.

DATA ENTRY MODULE - WORKBOOK TABS:

The Data Entry Module includes 12 separate worksheets, each designed (1) to organize the collection of field data, (2) to provide basic information used to evaluate data, or (3) to run

programmed macros to automate data migration and data analysis. The 12 worksheets of the Data Entry Module are described below.

1. “INSTRUCTIONS” TAB.

The instructions worksheet has basic instructions for entering data and running programmed macros in the Data Entry Module. Note that some institutional servers may disable macros. If there is a message banner near the top of the Data Entry Module that indicates macros have been disabled, follow the instructions beginning in row 23 to enable macros.

Many of the instructions found on the Instructions worksheet are repeated in this section, but they are covered in greater detail here. Also, the Instructions worksheet has two macros and a couple of links to useful information on the Instructions worksheet. The “IMPORT plant list from previous data” macro (see Instructions worksheet, column E), is used when an existing DMA is being re-monitored. By importing a plant list from an earlier reading of the DMA, users can prepopulate a list of plants previously observed in the DMA. This expedites the prerequisite of creating a plant list or reconnaissance of plants found along the greenline (see the MIM TR – Section 4.1 Systematic Procedures, Step 1. Develop a List of Plant Species (Burton et al. 2024). It also automatically populates columns M and Q-T in the PLANTS worksheet, which is used to create drop down lists for data entry in the DMA worksheet. Finally, importing plant data from a previous site visit permits the users to readily identify key species for stubble height and/or woody riparian species use.

The “Spatial Analysis” macro is used after monitoring data has been collected from about 10 to 20 sample points. This macro evaluates if there is spatial autocorrelation between data collected at adjacent sample points. Results of the spatial autocorrelation analysis and suggestions for how to proceed with the remainder of the data collection are discussed on the “Spatial” worksheet.

The Instructions worksheet, column E, includes a hot link (labelled “key species lists for the site”) to columns M through T in the PLANTS worksheet. This is the section where users make decisions regarding key species.

Also, in column E there is a link to the USDA-NRCS PLANTS database. This is a web resource with information on plant species and plant genera, including the official plant symbols that are used in the MIM Data Analysis Module. Access to the USDA-NRSC PLANTS database requires a connection to the internet. Whenever adding a plant to any part of the PLANTS worksheet, the

official USDA-NRSC PLANTS symbol must be used. The process of adding plant symbols is described in detail in “10. PLANTS” Tab below.

2. “HEADER” TAB

The “Header” tab includes descriptive information about the designated monitoring area (DMA). Much of the information entered in the first 12 rows of the Header worksheet is self-explanatory and required, e.g., the location (allotment, Forest/District, latitude and longitude (or UTM coordinates)), the observers, and the DMA identification, stream name and date of data collection.

Desired minimum sample size: The sample size estimator is described on the bottom of this tab. The desired minimum sample size for any variable that produces a mean or proportion can be estimated while collecting field data (Bartlett II et al. 2001). This is provided for in the “Header” and “DMA” tabs of the Data Entry Module. The equation for estimating the sample size estimate is:

$$n = (Z_{\alpha})^2 * s^2 / (\beta)^2$$

Where:

n = The sample size estimate.

Z_{α} = The standard normal coefficient from the table below.

s = The standard deviation.

β = The desired precision level (margin of error) expressed as half of the maximum acceptable confidence interval width (as a percentage of the mean or proportion).

This margin of error is derived from the sampling distribution of the repeat samples collected by different teams of observers (see Chapter III) and approximates the precision of the method or metric. The sampling distribution of a mean or proportion is generated by repeated sampling from the same DMA. This forms a distribution of different means, and this distribution has its own mean and variance from which the desired precision level is derived.

Standard Normal Coefficients:

Confidence level	Alpha (α) level	(Z_{α})
80%	0.20	1.28
85%	0.15	1.49
90%	0.10	1.64
95%	0.05	1.96
99%	0.01	2.58

The desired precision (β) level (or margin of error) provided in the Data Entry Module (on the “Header” tab) uses the margin of error (ME) for individual metrics as summarized on Table 9 in Chapter III of this Guide. This margin of error (ME) varies according to the metric value. That is, as the value increases, so does the ME. A relative margin of error (RME) was derived from the field test data and determined by the ratio of the ME to the metric value. This RME is therefore used to derive the ME or β values used in the module and displayed on the top of the “Header” tab as shown in figure 3.

A	B	C	D	E	F	G	H	I	J
1	HEADER FORM		For sample size: *	GGW	Stubble Ht	Substrate	Woody Use	Alteration	
2	Allotment:		Margin of error:	0.45	0.83	0.15	0.5%	6.6%	
3	Forest/District:		Confidence level:	95%	*See below for an explanation of this table.				
4	RD/FO:								
5	Observer(s):		Random number for starting pace set						
6	Plot spacing (m)	3.75	3	(Enter data in any open cell to re-calculate)			[click HERE to convert feet to meters]		
7	Starting Distance (m)	2	7.5	(sample point interval)					
8	DESIGNATED MONITORING AREA:			Downstream Marker		Upstream Marker		Reference Marker	
9	DMA ID	PASTURE NAME	STREAM	DATE	Latitude	Longitude	Latitude	Longitude	Latitude
10									
11	DMA NAME and/or Description:			(mm/dd/yyyy)	UTM Northing	UTM Easting	UTM Northing	UTM Easting	UTM zone:
12									
13	*Are hydrophytic woody plants supposed to be present at this site (y/n)?								
14	*Are there any hydrophytic woody plants present (y/n)?								
15	*Are all age classes of hydrophytic woody plants present (y/n)?								
16									
17	Units used to record Stubble Height (In - Inches, CM = Centimeters):			In		Slope Class**	Substrate Class**	Predicted from "Substr"	
18									
19	* - Required for calculating Ecological Status (see "Codes" worksheet (column A for instructions))			6th Field HUC:					# NUM!
20	Slope class: less than .5%, .5 to 2%, 2 to 4%, >4%, and >10%								

Figure 3. The Header form in the “Header” tab showing the basic descriptive information for the DMA. Also displayed here is the information used to compute the desired minimum sample size. The margin of error is shown for the 5 indicators for which sample size estimates are made. The default confidence level is 95%. The user has the ability to modify this level here if prior to sampling it was determined that a lower level would be used in the survey. This level and the margin of error are NOT modified after data have been collected at the DMA.

Using the above equation, sample sizes are estimated for the site using the desired precision derived from the RME (or ratio of the field-tested ME to the metric value) with 95% as the precision level. The Z_{α} values in the above table are applied to the "confidence level," and the standard deviation from the data collected in the sample. The default Z_{α} value for 95% confidence level is 1.96. Thus, while samples are being collected at the DMA and the standard deviation is changing with each sample, the sample size estimate that meets the desired precision level or margin of error is calculated and displayed on row 4 of the “DMA” tab, as shown in Figure 7.

If the estimated sample size on the “DMA” tab exceeds the number of samples collected, a lower confidence level would reduce that number on the “DMA” tab, **but using this method to adjust samples collected in the field is not recommended.** The purpose is to allow the user to assess sample size at the desired precision (β), not to change the confidence level. If the

observer changes the value in cell E3 to a new confidence level the change in sample size needed will be displayed in row 4 of the "DMA" tab. If the user wants to collect more samples at the default 95% confidence level to match the desired precision level and narrow the confidence interval, then the length of the DMA must be expanded to accommodate the additional samples. Re-sampling within the existing DMA is not recommended as it may result in spatial autocorrelation.

As stated by Elzinga et al. (1998, page 94):

"If you are faced with a monitoring situation where there is a lot of variability between sampling units (despite all of your sampling design efforts to lower this variability) and the components of your sampling objective lead to a recommended sample size of more sampling units than you can afford to sample, then you need to reassess the monitoring study. Is it reasonable to make changes to some components of the sampling objective? For target/threshold types of management objectives, this may mean lowering the level of confidence or decreasing the precision of the estimate (i.e., increasing the confidence interval width) or both. "

Standard deviation squared in the above equation is the variance. As described by Elzinga et al. (1998, Chapter 7), sample size formulas assume that the population approximately fits a normal probability distribution. For streambank alteration and woody riparian species use data, the samples are usually not normally distributed. They tend to be positively skewed (more low values than moderate or high values). Calculated variances, and therefore margins of error (and confidence intervals) may be underestimated in such circumstances. One method to correct this underestimation is to use random re-sampling methods (Elzinga et al. 1998, Chapter 11). Bootstrapping is a statistical procedure that applies random resamples to a dataset to create many simulated samples, allowing calculation of standard errors, and to construct confidence intervals from non-normally distributed data (Johnston and Faulkner 2020). As described in Appendix B, for purposes of estimating the variance for streambank alteration and woody riparian species use, bootstrapping was used to create 1000 re-samples from 50 MIM DMAs. This comparison of the 95% confidence interval of field samples to bootstrapped samples was examined by regression to estimate the amount of underestimation in the margin of error. Regression coefficients in this exercise were: streambank alteration ($r = .96$, $se = .004$), and woody riparian species use ($r = .88$, $se = .94$). These represent a reasonable adjustment to the calculated confidence interval and thus the sample size estimator. Since the standard deviation is proportional to the margin of error, the adjustment in confidence interval was applied to the

sample size estimator for streambank alteration and woody riparian species use. Fortunately, the amount of underestimation in the margin of error was found to be minor so that this sample size adjustment is not excessive. The author's test data showed that there was less than 1% difference between field sampled and bootstrapped standard deviations.

Calculation of ecological status: The MIM Header worksheet requires the observer to populate three key questions to calculate properly the ecological status of the DMA. These questions are found in rows 13-15 and include: (1) Are hydrophytic woody plants supposed to be present at this site, yes/no? (2) Are there any hydrophytic woody plants present, yes/no? And (3) Are all age classes of hydrophytic woody plants present, yes/no? For the purposes of these questions, hydrophytic plants are those with a wetland indicator status of obligate, facultative wetland, or facultative. The three primary age-classes of hydrophytic woody plants are represented by S (seedling), Y (young), or M (mature). Some general guidance on how to answer these questions follows.

Are hydrophytic woody plants supposed to be present at this site (yes/no)? If the stream at the DMA has a gradient over 0.5% and has water forces adequate to periodically cut banks and deposit bars, it likely should support a hydrophytic woody component and would be answered "yes." If the gradient is less than 0.5% and depositional features are absent, it would be "no." The presence of coarse sand to gravel on the bars, banks, or floodplains favors oxygenation of groundwater and establishment of woody species, which supports a "yes" response. Bars, banks, or floodplains entirely composed of silt and clay can be an impediment to establishment of woody species, which may indicate a "no" response. Groundwater near the surface and anaerobic conditions can also limit the establishment of woody species and should also be considered.

Are there any hydrophytic woody plants present (yes/no)? If any hydrophytic woody plants are present on the DMA, this would be "yes." If none exist, it would be "no."

Are all age classes of hydrophytic woody plants present (yes/no)? If there are seedlings, young, and mature (S, Y, M) hydrophytic woody plants along the DMA, this would be "yes." If one of these three age-classes is absent (or nearly absent) it would be "no." The woody riparian age classes are described in columns L-M of the **Codes** worksheet.

These questions must be answered for the ecological status metric in the Data Analysis Module (Data Summary worksheet) to be populated. The slope class (cell G18) and substrate class (cell H18) must be populated as well. Answering these questions will adjust the

ecological status rating downward if a hydrophytic woody component should be present but isn't, and/or if all three age classes (S, Y, M) are not represented.

For a detailed discussion of these concepts, see Winward (2000, pages 27-28).

DMA selection rationale: Whenever a new MIM DMA is established, the ID team should answer the selection criteria questions found in rows 26-37 of the Header worksheet. These questions determine if the criteria are being met for representative, reference, or critical DMAs. Generally, a representative DMA will meet all the selection criteria, and all 8 questions would have a "yes" or "not applicable," (NA), response. However, if there are mitigating circumstances, one or more of the selection criteria could be "no" and the mitigating factor(s) would be described in the Narrative section (see below). Reference DMAs commonly have many of the same criteria as representative DMAs.

The primary consideration is that the reference DMA is at or near potential or desired conditions and is in a sensitive complex that is of interest to management.

Critical DMAs do not have to meet any of the selection criteria of a representative DMA. The following is an example showing the type of information that is typically included in the DMA selection rationale and description.

22		
23	DMA Selection Rationale	
24		
25	Y, N, or NA	CRITERIA FOR REPRESENTATIVE DMA
26	Y	1. Was the riparian complex selected by an ID Team?
27	Y	2. Is the DMA in a complex that represents and is accessible to the management activity?
28	Y	3. Is the DMA randomly located in the riparian complex most sensitive to management?
29	Y	4. Is the DMA sensitive to disturbance (not armored)?
30	Y	5. Will the DMA site respond to management?
31	NA	6. If stream is over 4% gradient, does it have a well developed floodplain?
32	Y	7. Is the DMA located outside of a livestock concentration area?
33	Y	8. Is the DMA free from the influence of compounding activities?
34		
35	N	Is it a CRITICAL DMA?
36		
37	N	Is it a REFERENCE DMA?
38		
39	NARRATIVE	
40	The DMA is located approximately 65 meters upstream of the boundary fence (a random distance). The DMA is situated in a meadow complex with perennial streamflow, high water table throughout the growing season, and a Nebraska sedge/Beaked sedge dominant plant community. Few woody plants occur on the greenline. This riparian complex was selected by the ID Team for monitoring because it is sensitive to livestock grazing and will respond more rapidly to management changes. In addition, bull trout spawning is likely more prevalent in these gravel-dominated substrates. <u>Other riparian complexes in the pasture are dominated by cobble and boulder</u>	

Figure 4. An example of the DMA selection rationale.

Narrative. As described above, a cell at row 40 is provided to include a narrative about the DMA. The narrative is meant to be open-ended and allows for any inclusion of relevant information. Commonly, when a DMA is set up, it is a good practice to explain what the resource issues or management objectives are for the reach. For example:

This representative DMA was established in accordance with MIM DMA selection criteria and randomly located in riparian complex DBC-5 (700 m in length). The purpose is to use this DMA to assess overall stream/riparian condition in DBC-5 and for livestock management. The DMA was established in an herbaceous community (sedge/rush) within a meadow where livestock tend to gather for water and feed. Historically high levels of trampling and utilization have created low streambank stability and an overly widened channel. The management objective is to increase streambank stability to at least 80% and to decrease the greenline-to-greenline width by 50% within 6 years.”

Other issues addressed in the narrative could include:

- Recent environmental events that might help with the evaluation of data. For example, recent high-magnitude floods might explain a notable change in residual pool depth or coarsening of substrate.
- Mitigating factors that explain why a DMA is considered “representative” even if the site does not meet all the selection criteria.
- General impressions about condition.
- Observations or known grazing history in the current or previous year, for example, “Monitoring this year occurred before the pasture was grazed by livestock; however, there was ample evidence of elk trampling and woody browse along the streambanks.”

Satellite photo or sketch of DMA: A dedicated space is included to add a satellite or aerial drone view or a hand-sketched drawing of the layout of the DMA. The photos or sketches should be annotated to show the top (upstream), and bottom (downstream) ends of the DMA and the location of reference markers or landmarks that can help to relocate the DMA during future monitoring visits. An example photo is shown below.

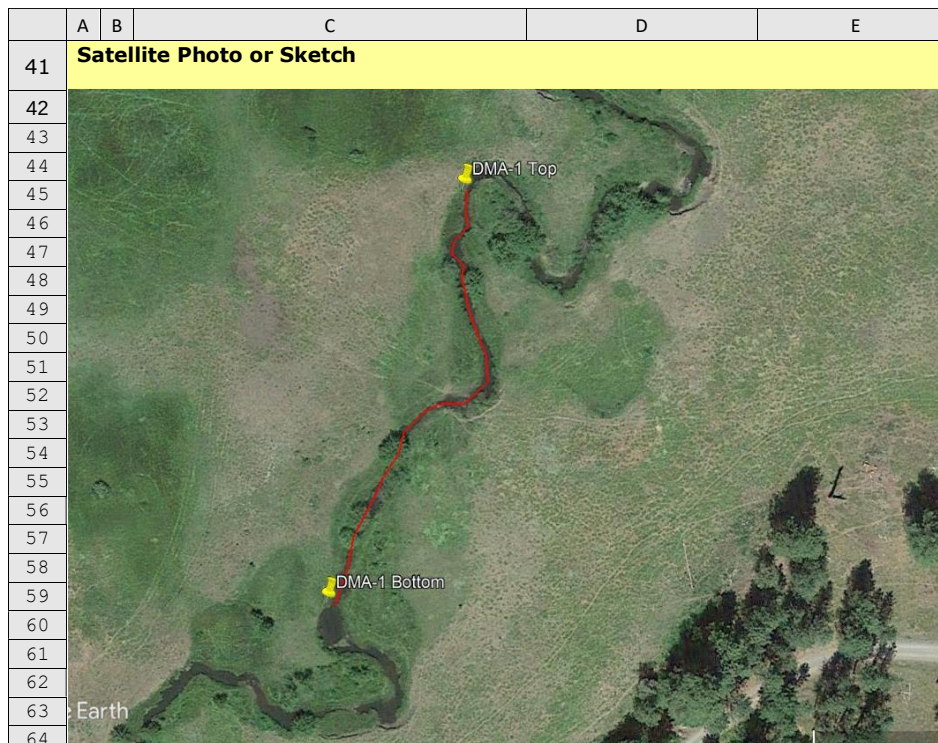


Figure 5. Example of an annotated satellite image of a DMA.

Photo log: A minimum of 4 photos, two at the top of the DMA and two at the bottom of the DMA, is required. Record the file name or path to a digital copy of the photos. The following contains an example of the “jpg” file names (without the extension) for an example showing the code name for the DMA followed by an acronym for the photo position.

	G	H	I	J	K	L
22						
23						
24						
25						
26						
27	?					
28						
29						
30						
31						

Photo log	File Name
Lower Up	SCTA_LU
Lower Across	SCTA_LA
Upper Down	SCTA_UD
Upper Across	SCTA_UA

Figure 6. Example photo log

Random number generator: The first quadrat is selected by a random number generator, which is embedded into the worksheet at cell D6. A new number is generated whenever the cursor is placed in a fillable space and any value is entered. (Hit the “Enter” key on the keyboard after typing a value to generate another random number.) The random number generator is shown in Figure 3 and is located near the top of the “Header” tab.

The random number generator selects a value between 1 and 5. This is because the average step length is 0.76 m, thus 5 steps equal 3.8 m, which is about the same as the standard sampling point spacing of 3.75 m. Using this random number of steps, the observer then proceeds upstream from the lower end of the DMA to locate the first sampling point. That sampling point’s starting distance is then recorded in cell C7 (Starting Distance (m)). The sampling interval used in the survey is also entered in cell C6. The default value, 3.75 m is supplied for convenience, but if there is any deviation from the default spacing, that distance **MUST** be entered in cell C6.

3. "DMA" TAB

The DMA worksheet provides a highly structured form for the entry of raw data for 8 MIM indicators (greenline composition and cover, woody species height class, streambank alteration, streambank stability and cover, stubble height, greenline-to-greenline width (GGW), woody riparian species age class, and woody riparian species use). There is also space provided to enter bankfull width, an optional measurement, which is not a formal MIM indicator.

Column A: Sample Number. Column A is used to show the sampling point, or individual quadrat sample point, within the DMA. All sampling points must be numbered sequentially beginning with “1.” In addition, the sample number is entered ONLY ONCE per sample point. If there is more than one row of data associated with a sample point, the sampling number is listed only with the first row of data and not with any other rows. In the example provided below, sample points 1, 4, 7, 11, and 12 have more than one row of data, but the sample number is only inserted with the first row of data per sample point. In contrast, sample points 2, 3, 5, 6, 8, 9 and 10 have only a single row of data.

	A	B	C	D	E	F	G	H	I	J	K	L
1	Unfreeze titles	Greenline Cover		Woody Height	Streambanks				Stubble Height			Width
2	Freeze titles	Species Rock/Wood	Cover	Woody Species height class	Streambank Alteration	Streambank Stability			Species	Height	Grazed?	GGW
3		(Code)	(%)	1 = < .5 m, 2 = .5 - 1 m, 3 = 1 - 2 m, 4 = 2 - 4 m, 5 = 4 - 8 m, 6 = > 8 m	(0 to 5)	E= Erosional D= Depositional	C= Covered U= Uncovered	F = Fracture, SP = Slump, SF = Slough, E = Eroding, A = Absent	(Code)	Inches	y/n	(meters)
4	Sample #		Composition = 100		N Min= 43	Sample points N = 59	Link to Cover table		N= 15	N Min= 77		N Min= 30
5	1	GUSA2	100	1	0	E	U	A				
6		MURI2	100						MURI2	22.0	7.9	6.0
7	2	MURI2	100		0	E	U	a	MURI2	20.0	5.3	6.0
8	3	rk	100		0	e	c	a			9.8	6.0
9	4	base	50	1	0	e	u	a			8.6	6.0
10		rk	50									
11	5	rk	100		2	e	c	a			8.2	5.0
12	6	rk	100		3	e	c	a			6.2	5.0
13	7	MURI2	50		5	e	c	a	MURI2	10.0	2.5	5.0
14		rk	50									
15	8	rk	100		0	e	u	e			1.1	5.0
16	9	rk	100		1	e	c	e			2.5	3.0
17	10	rk	100		1	e	c	e	AGROS2	4.0	9.2	
18	11	AGROS2	95		0	d	c		juto	10.0	2.4	2.8
19		juto	5									
20	12	AGROS2	33		5	d	c		AGROS2	4.0	9.9	3.0
21		juto	33						juto	15.0		
22		QUTU2	34	1								

Figure 7. Example DMA table with data. Composition in cell C4 is being used to add up plant species composition by sample point to simplify data entry. “N Min =” is displaying the calculated desired minimum sample size. The survey does not stop once this sample size is achieved, it must continue until all sample points on both sides of the stream have been completed.

Columns B and C: Greenline composition and cover. The plant codes for greenline composition are entered into column B and the corresponding percent of relative cover is added to column C. The relative cover must total 100 percent for understory cover and 100 percent for overstory cover. If only understory or overstory is present, then the total cover for the quadrat will be 100 percent; however, if both understory and overstory are present in the quadrat, then total cover will be 200 percent. An addition calculator is included in cell C4 to help the user check that the composition percentages add to either 100 or 200 percent in each quadrat.

Column D: Woody species height class. Woody species height class is required for all woody species listed in column B. **Note that the woody species height class entered in column D must**

be on the same row as the corresponding woody species listed in Column B. Failure to align these entries will generate an error when the data correction macros are run in the Data Analysis Module.

Column E: Streambank alteration. The data entry is straightforward. Enter a number from 0 through 5 for the number of observation lines in the MIM monitoring frame that intersect an alteration. Note that if there are no alterations, enter “0,” do not leave the value blank. The data entered does not have to be in the first row of the sample-point data, but just in any one of the rows associated with a particular sample point.

Columns F through H: Streambank stability and cover. Data entered in columns F through H are used to calculate the percent stable banks and the percent of covered banks. Data entry is straightforward. In column F, the only entry options are E (erosional bank) or D (depositional bank). In column G, the only entry options are C (covered bank) or U (uncovered bank). Column H is addressed only if the streambank is erosional (as shown in column F). If the streambank is depositional (as shown in column F), then Column H is left blank. For erosional streambanks, the entry options in column H are F (fracture), SP (slump), SF (slough), E (eroding), or A (absent) when there is no clear erosional feature.

Note in Figure 7, at cell G4 there is a link to the absolute cover table. As shown below in Figure 8, this table is used to input the percentages of live perennial vegetation, rock, and large wood on the streambank. The amount of bare bank and litter is calculated such that the total of the 4 categories equals 100%.

EO	EP	EQ	ER	ES	ET
Absolute streambank cover Return to DMA table					
	PERCENT COVER				
Sample #	Veg	Rock	Large wood	Calculated Bare/litter	
1	80			20	Determine the absolute cover for each of the cover constituents: (i) Perennial foliar vegetation cover, (ii) embedded rock including bedrock (>15cm or 6 in), and (iii) large wood (>10 cm or 4 in) Record each cover constituent to the nearest 10 percent. If two or more cover types overlap, do not add the overlap amounts and only record the portion of overlapping cover closest to or directly on the ground surface (e.g., record rock on the ground and not the overlapping vegetation cover immediately above the rock). Cover includes all of the above up to .5 meter above the ground surface Do NOT include bare ground, litter, or moss in these estimates. That cover quantity is calculated.
2	60	20		20	
3	30			70	
4	90			10	
5	100			0	
6	50		20	30	
7	50		20	30	
8	50			50	
9	80			20	
10	90			10	
11	80			20	
12	20	50		30	
13	30			70	
14	40			60	
15	20			80	
16	90			10	

Figure 8. Absolute cover table located at columns EO to ET on the “DMA” tab.

Columns I, J, and K: Stubble height. Stubble height data are entered into columns I, J, and K. The plant codes for key species are entered in column I. All key species observed in the quadrat are measured, therefore, there can be more than one measurement per quadrat. If there are no key species in the quadrat, leave it blank. Also, do not enter data for unavailable plants. A key graminoid plant is considered “unavailable” if it is caged (i.e. contained within or beneath a woody tree or shrub) or if it is on a cliff or other surface that cannot be reached by livestock or wild ungulates. The height is entered in column J to the nearest inch or nearest 2 (even) centimeters. Be sure to record on the Header worksheet, cell F17, the units of measurement – inches (In) or centimeters (CM). Inches is the default unit. Indicate whether the plant measured was grazed (y) or ungrazed (n) in column K. This added observation can be helpful if the grazing-use criteria are evaluated by percent utilization instead of height. Also, the grazed/ungrazed information can be used to evaluate data when there appear to be extenuating or unusual circumstances in plant growth. For example, if all the ungrazed plants do not even make a grazing-use criterion of 6 inches, then application of the grazing-use criterion is probably inappropriate at this DMA or at the selected monitoring time. There are known situations where the ungrazed plant heights do not meet grazing-use criterion, such as measurement too early in the growth period, plants stunted by drought, frost, or low soil temperatures, or extended floods that have delayed plant growth. Conversely, noting the potential height of ungrazed plants during optimal or normal conditions permits a reasonable evaluation of the level of herbivory. Substantiating the height of ungrazed plants is critical in this type of evaluation.

Columns L and M: GGW and BFW (optional). Greenline-to-greenline width (GGW) is entered to the nearest tenth of a meter in column L. Though not common, the option to record bankfull width (BFW) to the nearest tenth of a meter is provided in column M. Bankfull width is prone to high observer variability. It is recommended that bankfull width measurements are made by experienced users with direct knowledge of bankfull discharge and preferably in stable stream systems. Bankfull indicators tend to be ambiguous in unstable streams and along incised channels. Additionally, some stream types do not form bankfull indicators.

Columns N through R: Woody Riparian Species Age Class. Woody riparian species are hydrophytic, which includes all woody species with a wetland indicator status of obligate, facultative wetland, or facultative. The plant codes for all woody riparian species observed in a quadrat are entered in column N. The number of seedlings, young, and mature plants of each species in the quadrat are entered into columns O, P, and Q, respectively. However, woody riparian plants that are rhizomatous are not aged. They are treated as single plants and a “1” is entered into the rhizomatous column in the same row with a rhizomatous/clonal woody species in column N.

Columns S and T: Woody Riparian Species Use. The use level on key woody species is entered in columns S and T. The plant codes for the woody riparian key species are entered in column S and the corresponding use class is entered in column T. If there is more than one key species in the quadrat, enter all the species and all the corresponding use levels for each quadrat. Do not enter data for unavailable plants. A woody riparian plant is considered “unavailable” if more than 50% of the current year’s leaders on the plant are above the reach of the browsing animal. For example, for assessing cattle use, **the observer(s) would only consider key woody plants that have most of their current year’s leaders below 1.5 m (5 feet).** Woody plants with over 50 percent of the current year’s leaders above 1.5 m (5 feet) are considered unavailable for cattle.

4. “SUBSTR” TAB

Substr, short for **substrate**, is the worksheet for entry of raw substrate data collected in the field. The method for measuring substrate is provided in Section 6.3.6. Substrate in the MIM technical reference (Burton et al. 2024). The size of each selected particle is recorded. As shown in figure 9, each set of pebbles collected at a specific sample point are associated with a habitat type – pool or riffle, which is recorded in column L. Pools are defined exactly the same as those in the thalweg procedure. Riffles are all other habitat types not meeting the definition of a pool.

	A	B	C	D	E	F	G	H	I	J	K	L	M
1	DATE:		mm/dd/yyyy										
2	Substrate Form			Sample size needed:				531.66		N=	210		
3	Record sizes in millimeters			Did you use a gravelometer (y/n)?:									
4	Sample #	Pebble 1	Pebble 2	Pebble 3	Pebble 4	Pebble 5	Pebble 6	Pebble 7	Pebble 8	Pebble 9	Pebble 10	Habitat - pool(p) riffle(r)	Indicate which pebbles were estimated*
5	2	2	2	2	2	32	16	2	11	11	2	r	
6	4	2	2	2	8	22.6	8	11	16	8	2	p	
7	6	2	2	2	2	2	2	2	2	2	2	r	
8	8	2	2	2	2	16	11	8	2	2	2	r	
9	10	2	8	16	45	2	2	2	2	2	2	r	
10	12	2	2	2	5.6	5.6	16	16	8	8	22.6	p	
11	14	2	2	2	2	16	5.6	8	8	8	8	p	
12	16	2	8	11	45	16	22.6	2	2	2	2	r	
13	18	8	22.6	11	2	64	128	11	128	128	128	r	
14	20	2	11	16	11	5.6	16	22.6	8	8	4	r	
15	22	64	180	90	16	180	128	64	90	90	2	p	5, 6
16	24	2	2	8	16	11	16	22.6	32	22.6	2	r	
17	26	8	8	16	11	8	4	16	11	5.6	2	p	
18	28	5.6	8	64	16	22.6	16	16	22.6	11	8	r	
19	30	2	2	2	22.6	16	2	2	8	8	16	r	
20	32	2	2	2	11	22.6	32	16	22.6	2	5.6	p	
21	34	2	2	16	16	94	16	16	16	4	2	n	

Figure 9. The Substrate form. A habitat type, pool or riffle, is provided in column L for each sample point. If one or more pebbles are estimated rather than measured, the pebble number is recorded in column M.

A summary of the primary particle sizes of interest is included in cells N4 through R6. The percentage fines (cell N5) include the percent of particles that are less than 6 mm in size (b- or intermediate-axis measurement of particle diameter). D₅₀ stands for the 50th percentile particle size or the median particle size. D₁₆ stands for the 16-percentile particle size, or the size that is

larger than 16% of the particles in the sample. D_{84} is the size that is larger than 84% of the particles in the sample.

This worksheet also has a cumulative frequency distribution graph for the substrate data. The details of the graph show the size of the substrate particles as well as the degree of sorting. Well sorted populations have a steep cumulative frequency curve, while poorly sorted populations have a lower angle curve. Coarse populations are shifted to the right on the graph, while populations with an abundance of fine particles are shifted to the left.

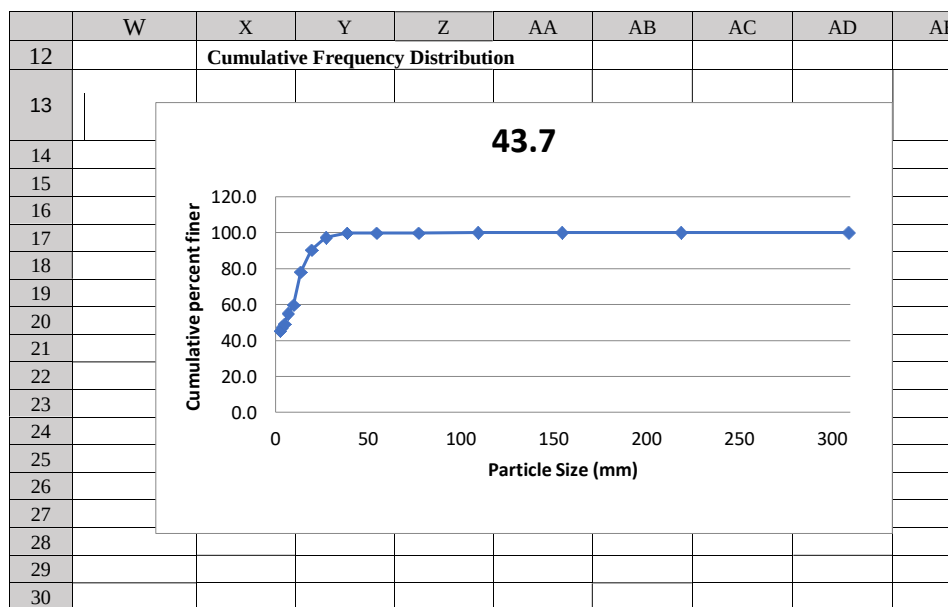


Figure 10. Cumulative frequency graph in the “Substr” tab.

5. “THAL” TAB

The Thal, short for thalweg, worksheet is the form for entry of raw data for residual pool depth and pool frequency. Data entry is made in columns A and B of the worksheet:

- Column A -- The distances between individual pool bottoms and riffle crests (or pool tails) are entered in column A. These distances are recorded to the nearest tenth of a meter (0.1 m).
- Column B – The depth of each riffle crest and pool bottom are entered in column B. These depths are recorded in meters (m) to the nearest hundredth of a meter (0.01 m).
- The entries for riffle crests (r) and pool bottoms (p) are aligned by row with the “r” and “p” entries in column C.

Summary calculations and statistics are provided in columns I through L (figure 11).

	I	J	K	L
1				
2				
3				
4	Total Distance		44.5	Meters
5	Percent Pools:		40%	
6	Number Pools:		3	
7	Pool Frequency:		108	Per Mile
8	Residual Pool Depth		0.14	Meters
9				
10				
11	ALL DEPTHS (m)			
12	MEAN	0.14		
13	ST DEV	0.08		
14	COUNT	7.00		
15	CI	0.06		
16				
17				

Figure 11. Summary metrics for pools in the “Thal” tab.

A similar summary table is provided in columns Q through T for the population of pools with a residual depth of at least 6 centimeters, the so-called **quality pools**. The summary statistics in this table exclude shallow pools with a residual pool depth less than 6 cm.

Columns G and O display a crude estimate of the length of pools. This is not a direct measure of the length of pools from pool tail to pool head; instead, it is a rough approximation of pool length calculated as 2 times the measured length from pool tail to pool bottom. It is this length that is used in the module to estimate Percent Pools.

6. “COMMENTS” TAB

The comments worksheet is a simple form that provides a dedicated space to add comments and observations about the DMA or individual sampling points within the DMA. The location of the observation is noted in column A. Simply record the number of the sampling point. However, if the comment is for the entire DMA, enter “All” or “DMA” in column A. The comments and observations are recorded in column B.

The added notes and other information serve both as a record of conditions observed at the time of monitoring as well as information that can help with the interpretation and evaluation of quantitative values. For example, users might include information on:

- The recent grazing history of a pasture or DMA. Monitoring data can vary considerably if the monitoring period occurs before, during, or after a period of grazing.
- Evidence of substantial wildlife impacts, especially from large ungulates like elk or from wild horses and burros.
- The recent history of large-magnitude floods, persistent drought, wildfire in the contributing watershed, etc., which might affect the values of some indicators. It is important to understand the influence environmental and climatic conditions can have on indicator values so this influence can be factored into or differentiated from the effects of management actions.
- Current streamflow conditions, which might affect interpretation of some indicators.
- Recent gain or loss of beaver dams.
- Changes in streamflow related to flow regulation or flow diversion.
- The presence of noxious weeds or invasive species, which may require treatment.
- The presence of a rare, threatened, or endangered species that might occur off the greenline or in amounts that are not typically recorded in MIM (i.e., a minimum of 10% relative cover in the quadrat).
- Anomalies that could create misleading statistical conclusions. For example, measurement of coarse, angular particles that are not transported by the stream but as a result of rock fall and gravity could greatly skew particle-size calculations, which should really reflect only particles that are water transported.

Finally, the comments worksheet includes notes (see columns I and J) to describe any updates or recent changes and improvements in the current Data Entry Module.

7. “GRAPHS” TAB

When historical data are imported using the macro on the Instructions worksheet, “IMPORT plant list from previous data,” the data are migrated to the Graphs worksheet. This worksheet provides the source for the plant lists created in the PLANTS worksheet. In columns AG-BO there are graphic and tabular displays of the historical data, including for woody plant height, woody riparian age class, plant species composition, and greenline-to-greenline width. Finally, a biodiversity index is provided in cell G5. Generally, DMAs with higher biodiversity tend to have lower spatial autocorrelation than those that have lower biodiversity and greater homogeneity. Knowing in advance the biodiversity of the DMA can provide a leading indication

if there might be negligible or considerable spatial autocorrelation. This information might trigger an examination of existing data with the Spatial worksheet in this module and the Statistical Analysis Module to determine an appropriate sampling interval.

8. “CODES” TAB

The codes worksheet displays metric summary codes such as for "Ecological Status." Many of the MIM data entries use a code to shorten data entry and to facilitate automated analysis of data with the Data Analysis and Statistical Analysis modules. Codes must be entered exactly, or the analytical macros will not work properly. All MIM-related codes are summarized on the “codes” worksheet. Lists of codes are supplied for:

- Ecological Status,
- Slope classes (gradient),
- Substrate classes,
- Winward (2000) capability groups,
- Woody height classes,
- Streambank stability, and
- Woody riparian species age classes.

9. “CALCS” TAB

The calcs worksheet summarizes arithmetic operations for the sample-size estimator. This includes the adjustments made for non-normally distributed data using the bootstrap analysis results described in Appendix B.

10. “PLANTS” TAB

The “PLANTS” tab includes several plant lists with many of the common riparian plants of the western United States. These lists include the official NRCS PLANTS database species and genera codes. The various plant lists include:

- Alphabetic plant list by scientific name and arranged by plant functional groups (graminoids, forbs, shrubs, trees, and other) (columns A-D).
- Plants listed alphabetically by scientific name (columns E-H).
- Plants listed alphabetically by common name (columns K-L)

The MIM Data Entry and Data Analysis modules use several other codes that are not official PLANTS database codes. These codes are listed in Table 1.

Table 1. Additional MIM codes for greenline composition.

Code	Name	Wetland Indicator Status*	Successional Status	Winward greenline stability rating [#]
CAREXRH	Rhizomatous sedge	FACW	Late seral	8.5
CAREXTU	Tufted (clumped) sedge	FACW	Mid-seral	2
JUNCURH	Rhizomatous rush	FACW	Late seral	8.5
JUNCUTU	Tufted (clumped) rush	FAC	Mid-seral	2
MFE	Mesic forb early seral	FAC	Early seral	2
MFL	Mesic forb late seral	FACW	Late seral	8.5
MFM	Mesic forb mid-seral	FAC	Mid-seral	2
MG	Mesic grass	FAC	Early seral	2
MGRH	Rhizomatous mesic grass	FAC	Mid-seral	5
MGTU	Tufted mesic grass	FAC	Early seral	2
MS	Mesic shrub	FAC	Mid-seral	5
NG	No greenline	UPL	Early seral	1
RK	Embedded rock			10
UF	Upland forb	UPL	Early seral	2
UG	Upland grass	UPL	Early seral	2
US	Upland shrub	UPL	Early seral	2
WD	Anchored wood			10

* FAC = facultative; FACW = facultative wetland; UPL = obligate upland

[#] Winward greenline stability ratings vary from 1 (lowest stability) to 10 (highest stability).

In addition, the “PLANTS” tab includes an area to record a list of plants (by species codes) that occur along the greenline and to designate the key species for the DMA:

- A user-filled list of plants found on the greenline (column M). This column can be populated in two ways. Either the user can walk the DMA, making a list of plants observed on, near, or overhanging the greenline. Or in the case of an existing DMA with previous data collected, the user can run the “IMPORT plant list” macro found at cell E5 of the “Instructions” tab to automate entry of existing data into columns M, and Q through T.
- Key species are entered into columns N and O. The user can select key graminoid and key woody riparian species based on a reconnaissance of the DMA and deciding which palatable species are abundant and appropriate for measuring stubble height or woody riparian species use. Alternatively, if existing data has been imported into columns Q through T, the user can use the results of earlier monitoring to pick key species (columns Q and S) that are palatable, abundant (columns R and T), and important to management.

	M	N	O		Q	R	S	T	U
1	DMA PLANT LIST	KEY SPECIES			Key Species list from previous sample				
2		STUBBLE HEIGHT	WOODY USE		STUBBLE HEIGHT	FREQUENCY	WOODY USE	FREQUENCY	
3									
4	ACMI2	CAAQ	POTR5		ALAE	2	POTR5	1	
5	ALAE	ALPR3	SAEX		ALPR3	6	RIAU	1	
6	ALPR3	POPR					RIBES	6	
7	ARCA13						ROWO	17	
8	ARTR2								
9	BRIN2				CAAQ	1			
10	CAAQ						SAEX	17	
11	CEST8								
12	CIAR4				HOBR2	9			
13	HOBR2				MG	2			
14	JUOC								
15	JUSC2								
16	MEAR4				POPR	68			

Figure 12. Plant lists uploaded to the Data Entry Module from previously collected data at the DMA as contained in the “PLANTS” tab.

In figure 12, plants collected in the earlier sample are listed in column M. Plants used to measure stubble height and woody use are listed in columns Q and S respectively. The information from the table on the right is used to populate the Key Species list in columns N and O. The user chooses which plants will be used to assess stubble height and woody use. The plants in columns M, N, and O will then be displayed in the drop-down lists in the “DMA” tab as shown in figure 12. Once species codes are entered into columns M, N, and O, the user will be able to use drop-down lists to enter data into columns B (greenline composition), I (stubble height species), and S (woody riparian species use) on the DMA worksheet.

11. “KEYSP” TAB

This worksheet includes a master list of possible key species used for stubble height, woody riparian age class, and woody riparian species use. This is a comprehensive list, but it is not exhaustive. If the plant communities on a DMA include species that are behaving and functioning like key species, but these are not on the list, then users will have to add them to the bottom of the appropriate column (A for additional graminoid species and D for additional woody plants).

If one or more added species codes are added to the lists in column A or D, these columns can be re-sorted to insert the additional codes in proper alphabetic order. Select the top species code (row 3) in the column to reorder and then in the header select the “Data” tab and “Sort A to Z.”

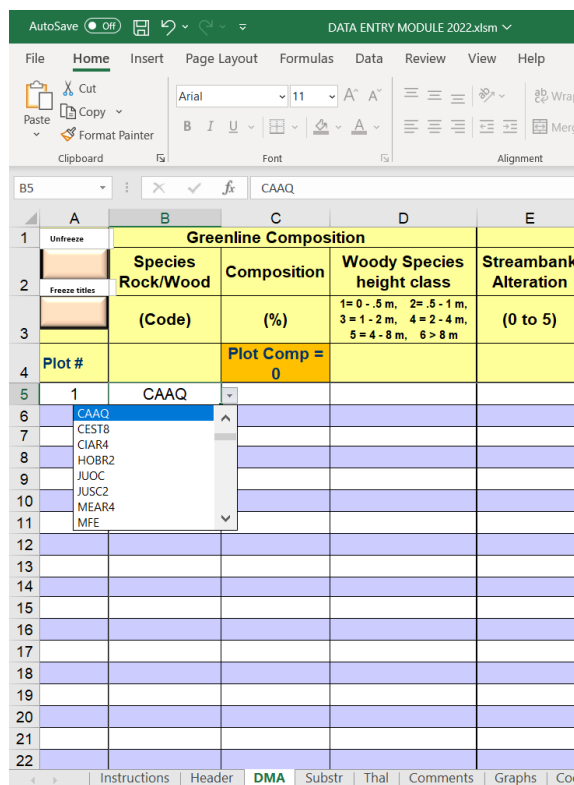


Figure 13. Species listed in the drop-down list on the “DMA” tab of the Data Entry Module.

12. “SPATIAL” TAB

Testing for spatial autocorrelation may be conducted after collecting data from about 10 to 20 sample points to see if there is any indication that the sample point spacing is not adequate for sample independence involving one or more of the indicators. The user runs the test by pressing the” Run Spatial Analysis” button on the “Instructions tab”. The output is displayed at the end of processing this test, as shown in figure 14.

	A	B	C	D	E	F
1						
2			Run spatial analysis			
3	Sample point spacing:		3.75			
4						
5	CORRELATION COEFFICIENTS (r)					
6	Indicator	Adjacent	Every other	Every third	Sample points (N)	
7	Bk Alter	-0.3049	-0.0341	-0.1827	20	
8	St Ht	0.2930	-0.3895	-0.4994	20	
9	GGW	0.5198	0.1966	-0.1130	20	
10	Wdy use	0.0293	-0.2775	0.1415	20	
11	Significance test using the t score (adjacent plots):					
12		Bk Alter	St Ht	GGW	Wdy use	
13	t score	-1.358	1.300	2.582	0.124	
14	p value	0.191	0.210	0.019	0.903	
15	significant?	N	N	Y	N	
16						
17	Reference: https://www.ecologycenter.us/vegetation-ecology/correlograms					

Figure 14. Results of the test for spatial autocorrelation indicate that none of the indicators are showing spatial autocorrelation except GGW. For GGW, the correlation coefficient of 0.5198 suggests the possibility of spatial autocorrelation as sampling continues (there were only 20 sample points so far in the survey). In these instances, the user can adjust sampling so that only every other sample point is measured for GGW, which had a much lower correlation coefficient of 0.1966, which is not significant.

A table of correlation coefficients and results of the t-score test are displayed showing whether spatial autocorrelation is present in the data. If not, sampling continues without interruption. If there is indication of spatial autocorrelation in one or more of the indicators, the user has two options: 1) Continue sampling at the current sample spacing and address autocorrelation issues after completing the survey (i.e., use every other sample point for the questioned indicator(s) or treat the entire DMA as one sample); or 2) increase the sample spacing to reduce the potential for spatial dependence for the indicator(s) in question. This latter approach may be preferable if two or more indicators are showing spatial autocorrelation.

To view an example correlation table used to evaluate spatial autocorrelation, see Table A1 in Appendix A (Spatial Autocorrelation Assessment).

B. DATA ANALYSIS MODULE

The Data Analysis Module is used to analyze and store data. Although it can be used to enter data manually in the field, **it is NOT as useful for field data collection as the Data Entry Module because it does not have all the data entry functions of the Data Entry Module.** This file then becomes a record of the data collected and the analytical results. Files are stored and named according to the user's choice. A good convention is to name files by their DMA name, and date or year. The Data Analysis Module is organized much like the Data Entry Module with worksheets as shown in figure 15.

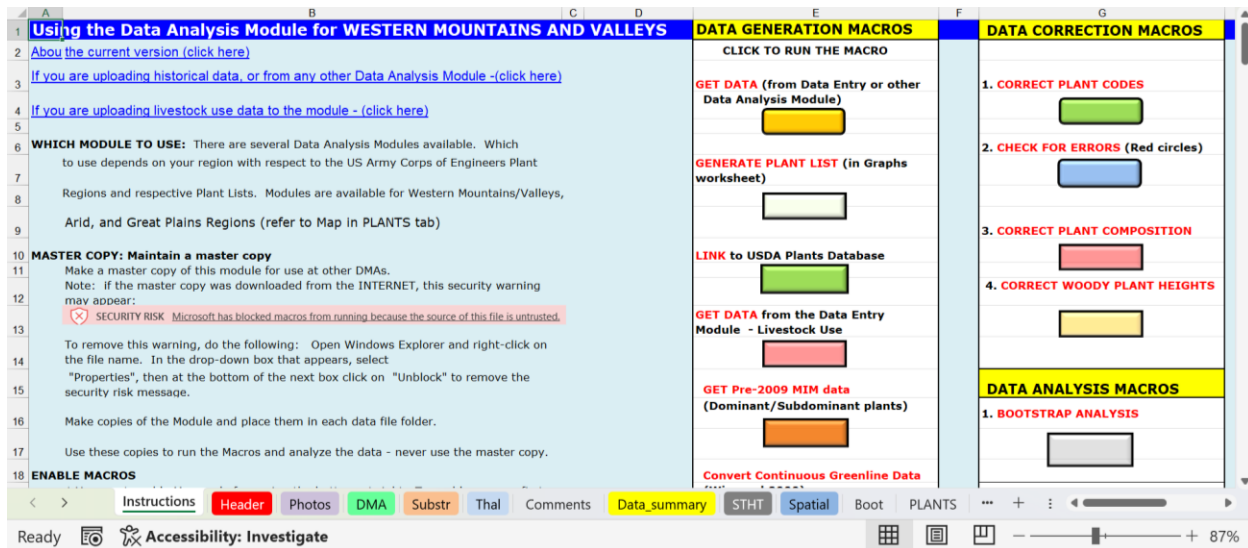


Figure 15. The “Instructions” tab of the Data Analysis Module. Note the tabs at the bottom for most of the 19 worksheets in the module. There are macros for Data Generation, Data Correction, and Data Analysis.

There are 3 Data Analysis Modules as shown in figure 16.

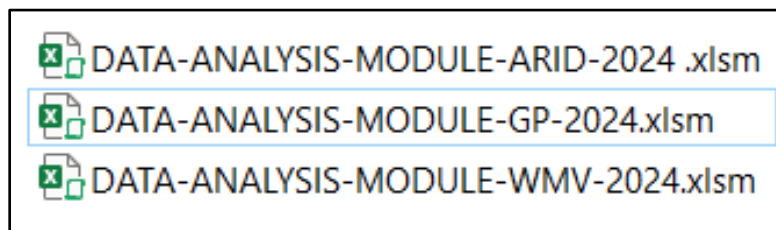


Figure 16. Data Analysis Modules for the Arid (ARID), Great Plains (GP), and Western Mountains/Valleys and Coast (WMV) plant regions.

Each module is applicable to a plant region (WMV – Western Mountains/Valleys and Coast, GP – Great Plains, and ARID – Arid West) as defined in the 2014 National Wetlands Plant List, <http://plants.usda.gov/core/wetlandSearch> . Or https://wetland-plants.usace.army.mil/nwpl_static/v34/home/home.html

The map in figure 17 depicts each of the three plant regions in the Western US: ARID WEST, WESTERN MOUNTAINS/VALLEYS and COAST, and GREAT PLAINS.

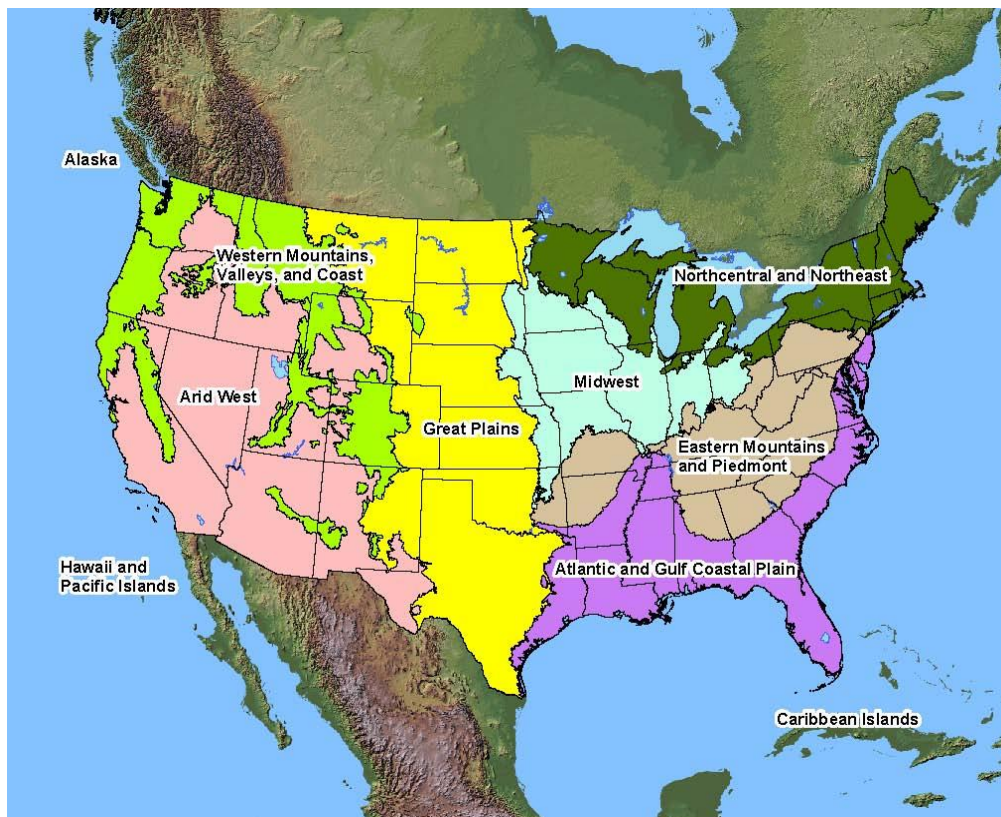


Figure 17. Map of US plant regions used to attribute riparian plant species in the Data Analysis Modules.

The layout of the “Header”, “DMA”, “Substr”, “Thal” and “Comments” worksheets is the same as that of the Data Entry Modules. This module also contains a proper functioning condition assessment “PFC” validation worksheet that displays many of the metrics relative to PFC checklist items (USDI 2015). The “Graphs” worksheet has more detailed metric summaries for plant species, including a graph of relative species composition. The “Calcs” worksheet describes how each metric is mathematically derived. The “Data Summary” worksheet has 45 metric summaries such as average stubble height and percent stable streambanks. There are several statistical analyses available in the module including confidence intervals for several of the key metrics, analysis of potential spatial autocorrelation for the relevant indicators, cumulative frequency distributions for substrate data, and data distributions for the short-term indicators allowing the user to identify the need to use re-sampling (bootstrap) to calculate means, medians, and confidence intervals for non-normally distributed data. Figure 18 shows

user reads and understands the instructions prior to using the module for the first time. The module is updated regularly, so the ***“about the current version”*** link takes the user to a location describing the latest improvements/updates.

Which module to use and ***Master copy*** are self-explanatory. More details on which module to use are described above.

Enable macros describes how to allow macros to be run in Excel. To enable the macros, first select File (top left), then select "Options" then "Trust Center, followed by "Trust Center Settings". Here select "Macro Settings", then "Enable All Macros". **The file should be saved as type: ".xlsm" to support the macros.** Once saved, close the file then reopen to allow the macro settings to take effect. At this point it is good practice to save the primary file and set it aside so that it can be available for other DMAs in the future. Then close it, and re-open it again and save according to the DMA name or other file naming convention. With this file and the master enabled, future Excel files will also allow macros to be run, however in recent WINDOWS updates, the developer of Excel has added a security warning when macros are embedded in a file downloaded from the internet. This warning can be disabled by simply right clicking on the file in Windows Explorer and then choose "Properties". At the bottom of the dialog box that appears, select "Unblock" (place a checkmark in the small box) at the bottom of the properties box. Once you have enabled macros, be sure that future Excel files downloaded from the internet are from known, safe sources.

Worksheets in the Data Analysis Module gives a brief description of each of the Excel tabs in the workbook. Much more detail on each tab is contained in the following descriptions.

Entering data into the module describes how to upload data electronically into the module using the macros provided on this worksheet. It also describes how to enter data by hand without using the macros. However, copying and pasting by hand can be quite cumbersome and time-consuming and may result in data errors where formatting conventions are lost.

Entering pre-2009 data into the module describes how to use the module to convert pre-2009 MIM data to the current module. It is always a good practice to upload historic MIM data into the current version of the module, especially before using the data, no matter the year it was collected. This supplies the most up-to-date metrics and statistics important for data interpretation and analysis. Pre-2009 MIM data were based on plant dominance rather than relative plant composition. As such, MIM modules for those years are not compatible with the current version. This tool allows for conversion of plant data collected by dominance to plant

data by relative composition. It basically assigns a higher percent cover to plants designated as “dominant” and lower composition to those designated as “sub-dominant.” While not mentioned here, the module has another macro that converts plant data collected using the line intercept method (or continuous greenline data) to relative composition. This conversion then provides for importing the continuous greenline data (i.e., the method of Winward 2000) into the Data Analysis Module.

“Macros” describes each of the macros and how to use them. The Data Generation macros are run to upload data into the module or to generate plant lists from the corrected data or to access the USDA PLANTS Database. The Data Correction macros are used to correct field data that has been uploaded to the module. It is important that all data uploaded to the module be corrected by running all four data correction macros. Failure to do so will not only leave errors in the data but will prevent execution of several crucial functions leaving some outputs absent. For example, a failure to run Macro “Correct plant composition” will leave the data table at columns AT to DF in the comments sheet blank. That table is used to upload data from the module to the Statistical Analysis Module. The Data Analysis macros allow the user to evaluate the data for potential spatial autocorrelation and to run the bootstrap analysis to resample for computation of means, medians, and confidence interval for non-normally distributed data (which is common for streambank alteration and woody riparian species use).

Adding or replacing plant codes describes how to add and/or replace plant codes in the module and is self-explanatory.

Export data to the MIM database explains how to use the “Export” tab. The MIM database is currently unsupported, however some users still use this tab to export and store data on a local server or computer.

Suggested steps for use of this module supplies a list of steps that users can follow to properly upload and correct MIM data.

Greenline plant composition: There are 498 rows in the data table on the “DMA” tab. The number of rows is limited for a particularly important reason. Generally, individual quadrats should have no more than an average of 6 plants per quadrat. This is because the MIM protocol requires that only those plants having at least 10 percent or more foliar cover by composition are recorded. Plants with less cover supply minimal contribution to calculation of the plant metrics. The protocol does allow recording important plants with less than 10 percent foliar cover, but this should be the exception not the rule. Use the “Comments” tab to

record plants with less than 10% foliar cover - basically the minority and trace plant species, to capture their presence.

2. "HEADER" TAB

This worksheet contains descriptive information about the designated monitoring area (DMA). It is identical to the Header Sheet contained in the Data Entry Module as described above ([link to Data Entry Module – Header tab](#)). This tab may or may not have been completed in the Data Entry Module prior to uploading data to the Data Analysis Module. It is good to complete as much as possible in the field (in the Data Entry Module) so that some of this information does not have to be retrieved from memory. An example is the selection criteria when starting a new DMA, most of which are answerable while viewing the DMA in the field. Some information may have to be added to the "Header" tab in the office. A satellite image of the DMA, for example, is likely best copied and pasted in the office. Any additions to the "Header" tab made in the office can be done in the Data Analysis Module after uploading data from the Data Entry Module.

3. "PHOTOS" TAB

This tab has photos of the upper/lower markers and photo comparisons provided by the user. Photos are copied (i.e., from .jpg files) and then pasted to this worksheet as shown in figure 19 below.



Figure 19. Photos of the marker locations at a MIM DMA.

4. "DMA" TAB

This worksheet contains raw data table for greenline cover and composition, woody species height class, streambank stability and cover (bank type, cover, erosion feature), greenline-to-greenline width, bankfull width, woody riparian species age class, and grazing-use indicators (streambank alteration, stubble height, woody riparian species use). It is the same as in the Data Entry Module described above ([link to Data Entry Module – DMA tab](#)). The main difference is that this worksheet does not have the drop-down lists useful for entering raw data into individual cells. Thus, the Data Analysis Module is not as useful for data entry as the Data Entry Module. To the right of column U (columns V to EI), are several tables used for various analyses. These tables do not exist in the Data Entry Module. Just ignore these tables and do not try to alter any cells in these columns. Doing so will corrupt the analytical results. These columns with their data manipulations could be hidden to protect their integrity and users would be prevented from accessing them. However, these cells are locked so if the user

refrains from unprotecting the worksheet, accidental deletion of these cells can be avoided. Also, it is the desire of the developers to be transparent and leave these cells visible to the users so that they can examine how the data operations are used to derive metrics and statistics in the module if they so desire.

5. "SUBSTR" TAB

This worksheet has raw data for substrate information collected in the field. It is the same as in the Data Entry Module described above ([link to Data Entry Module – SUBSTR tab](#)). One of the useful outputs of this table, not available in the Data Entry Module is the percent fines by sample point calculated in column N. Percent fines, and other indicators of substrate conditions can then be evaluated for (1) the entire DMA, (2) just the riffles, and (3) just the pools. The summary of substrate data for the entire DMA is displayed in the “Data Summary” tab in cells B21-E23. The habitat-specific substrate data for pools and riffles are displayed on the “Substr” tab in cells AG4-AI10.

6. "THAL" TAB

This worksheet has data for calculating residual pool depth and pool frequency. It is the same as in the Data Entry Module described above ([link to Data Entry Module – THAL tab](#)). As with other worksheets in the Data Analysis Module, this sheet does not contain the drop-down lists as in the Data Entry Module.

7. "COMMENTS" TAB

This worksheet is used to display comments provided by the field data collection and has the same information in Columns A and B as in the Data Entry Module. By contrast, this worksheet also has several tables used for a variety of purposes including:

Columns O to Q – to upload dominance plant data from a pre-2009 module.

Columns Z to AM – to convert dominance plant data to plant composition by quadrat. This section is referenced by the macro “Get pre-2009 MIM data” executed from the “Instructions” tab. Once completed, don’t forget to click on the orange button that sends the converted data to the DMA spreadsheet.

Columns AT to BF – to organize quadrat data into a table for statistical analyses including computation of correlations displayed in the “Correl” tab and for upload to the Statistical

Analysis Module.

Columns BQ to CD – to correct plant composition data and send it back to the DMA sheet.

This section is referenced by the macro “Correct plant composition” executed from the “Instructions” tab. Follow the instructions supplied and then don’t forget to click on the orange button that sends the corrected data back to the DMA spreadsheet.

Columns CO to DJ – to convert continuous greenline data to plant composition by quadrat.

This section is referenced by the macro “Convert continuous greenline data” executed from the “Instructions” tab. Follow the instructions and then don’t forget to click on the orange button (at cell DE2) to send corrected data back to the DMA spreadsheet.

8. “DATA SUMMARY” TAB

This worksheet has metric summary data for the DMA. The general format for the data summary table is shown below in figures 20 and 21. Some key features and their explanation follow:

- 1) The table is divided into two sections – short-term indicators and long-term indicators.
- 2) In the black block at the top left are the DMA identifiers and date. In the black box at the top right are links to other locations providing additional interpretations of the data. These include the PFC worksheet, Graphs, and Correl worksheets (figure 20) and the stubble height analysis, short-term data distributions, and substrate particle size analysis (figure 21). These are described in their respective Excel tabs. The stubble height analysis and short-term data distributions are described in the “Graphs” tab. Just click on the wording in white letters and the link will automatically navigate to the respective worksheet.
- 3) Blue header boxes are for the short-term indicators, green for the long-term vegetation indicators, and orange for the long-term channel indicators. Wording in these headers is depicted in blue underlining which shows that each header is also a link. Selecting and clicking on the header name automatically navigates to the “Calcs” tab where there is a description of the calculation of each metric. For example, if the user selects “Median SH....”, the following is displayed:

Median SH all Key species (inches)	The median value of all stubble height values entered into Column J of the DMA spreadsheet. As with Mean SH, all key species are integrated in this calculation.
--	--

- The two rows of data immediately below the header row contain the metric value in the first row and the number of samples used to calculate the metric in the second row. For example, average stubble height for all key species in figure 20 is 9.0 inches based on a sample size of 105.
- 4) The two rows below the number of samples contain the 95% confidence intervals (if available) for the metric. A detailed discussion of the confidence interval (CI) is contained in [Chapter III, part B, 2](#). In the first row in light gray (95% conf int¹) is the confidence interval calculated from the field data collected for this sample. This confidence interval is calculated for the mean or proportion using the standard normal coefficient from Excel's "confidence" function. It is adjusted for non-normally distributed data (streambank alteration and woody riparian species use) using the equations derived from the bootstrap analysis (see appendix B). In the second row in yellow (95% conf int²) are the confidence intervals derived from testing of the MIM indicators as described in Chapter III with results displayed in Table 9 of that chapter. These tests of repeatability (or observer variation) of each metric supply a basic estimate of the precision of the metric. Generally, the 95% confidence interval in the first row has a lower value than that in the second row. The goal is to collect enough data in the field to derive a lower confidence interval for the field data than that of the test data. However, regardless of the CI value in either row, the value for the field data in the first row is always preferable (see the discussion of 95% confidence interval provided in the "[Spatial](#)" tab).

Note that for some indicators (i.e., stubble height, streambank alteration, and greenline-to-greenline width), the 95% confidence interval is calculated from an equation rather than from Table 9 in Chapter III. An example is in cell J24, which used this equation: $0.006307 + J21 * 0.07506$. Cell J21 is the metric value for greenline-to-greenline width. As that value increases, so does the 95% confidence interval as seen in the test data, and regression equations were derived from that relationship. Basically, the CIs based on the test data increased at a greater rate with increasing mean than the CIs based on a single measurement of the sample point. So rather than using just the average confidence interval derived from the testing, it was decided that the regression value would be more correct (e.g., the higher the GGW, the higher (i.e., wider range) the 95% confidence interval from the replicate samples).

METRIC DATA SUMMARY				LINK TO PROPER FUNCTIONING CONDITION (PFC) ANALYSIS				
DMA = DMA-01				LINK TO GRAPHS WORKSHEET				
Pasture = Spears meadow				LINK TO CORRELATION MATRIX				
Date = 6/10/2022								
SHORT-TERM INDICATORS								
Stubble Height						Woody Use	Streambanks	
Median SH for all key species (in)	Average SH for all key species (in)	Average SH for all Key species GRAZED (in)	Average SH for all Key species UNGRAZED (in)	Dom key species for SH	Avg SH of dom key species (in)	Woody species use - all woody species MEDIAN (%)	Streambank alteration (%) altered)	Streambank stability(%)
8.00	9.0	7.6	10.5	JUBA	10.83	10	20%	100%
n=	105	49	48	29		25	46	45
95% conf Int ¹	0.80			*	1	*	*	*
95% CI ²	1.12					6%	6%	5%
LONG-TERM INDICATORS								
Vegetation Ratings				Miscellaneous Vegetation Metrics				
	Greenline ecological status rating	Site wetland rating	Winward greenline stability rating	Vegetation biomass index	Percent rhizomatous woody	Percent forbs	Plant Diversity Index	Hydric plants (% by constancy)
Rating	81	74	7.02	59	2%	2%	10.68	76%
n=	Late	FACW-	High					
	*	*	*	120	1	3	151	115
95% conf Int ¹	*	3.5	*	*	*	*	*	*
95% CI ²	5.75	3	0.16					6.2
Substrate:					Pools			
	Percent fines	D16 particle size (mm)	D50 particle size (mm)	D84 particle size (mm)	Total number pools	Pool frequency (#/mile)	Mean residual depth - All (m)	Mean residual depth - >.06 (m)
	44%	0.9	7.75	18	10	180	0.33	0.33
n=	70	70	70	70	20	20	20	20
95% conf Int ¹	12.06	*	*	*	*	*	0	0
95% CI ²	11.6					14	0.06	0.06
¹ 95% conf Int: 95% confidence interval based the data in this DMA								
² 95% CI: the 95% confidence interval from all test sites (see Table F7 in TR 1737-23) MORE								

Figure 20. Left half of the Data Summary table showing results for the Spears Meadow DMA on June 10, 2021.

LINK TO STUBBLE HEIGHT ANALYSIS LINK TO SHORT-TERM DATA DISTRIBUTIONS LINK TO SUBSTRATE PARTICLE SIZE ANALYSIS				
Streambank cover (%)	Covered - stable (%)	Covered - unstable (%)	Uncovered - stable (%)	Uncovered - unstable (%)
100%	96%	2%	2%	0%
45	43	1	1	0
*	*	*	*	*
5%	5%	5%	5%	5%

Woody Riparian Species Age Class					
Woody composition (%)	Woody species frequency (N)	Hydric herbaceous (%)	Percent seedlings	Percent young	Percent mature
18%	0	58.9%	26%	37%	37%
26		89	20	14	14
*	*	*	0	*	*
5.9		6.2	7%	7%	7%

Width and Shade			
Greenline-greenline width (m)	Average woody plant height (m)	Shade index	Bankfull width (m)
3.99	2.2	0.34	
42	31	120	0
0.30	1	*	*
0.31			0.30

Figure 21. A portion of the right half of the data summary table.

9. “STHT” TAB

The STHT worksheet analyzes stubble height statistics for selected key species, one at a time, or in combination for up to 4 stubble height key species. As stated in the Technical Reference “Generally, no more than four key species are used at a DMA”. **Key species** are plants that are relatively palatable to grazing animals, relatively abundant, important for stream/riparian function and habitat, and serve as indicators of environmental and management changes. This worksheet allows the user to examine the results of data collected at the DMA and to see which species are relatively abundant and palatable. The

target sample size is 50 or more of these samples to give a reasonably precise estimate of mean stubble height.

Some key species may have potential growth heights different than others in the survey. For that reason, it is often desirable to examine the stubble height statistics separately for those plants rather than combine them with all key species in the survey. A good example of this is shown in the following graphic describing the stubble height analysis in the “STHT” tab. The results are presented in a BASIC STATS table as seen in the example below. The data in this table are automatically copied to the Data Summary Tab at column G, rows 6 to 8.

1	Note: For proper execution, run the “Generate plant list” macro before using this analysis				F	G	H	I
2					Enter species code(s) below to view statistics (up to 4 species)			
3	STUBBLE HEIGHT ANALYSIS							
4	STUBBLE HEIGHT		*Proportion of “Y”s			Species*	* leave these cells blank to use all species collected in f	
5	Key Species	N	Avg Height (in)	% Grazed*		1	PHAR3	
6	AGST2	1	2	100%		2		
7						3		
8						4		
9	CYDA4	1	4	0%				
10								
11	MG	1	5	0%		BASIC STATS (before bootstrap and spatial analysis)		
12						N	STANDARD DEVIATION	95% CONFIDENCE INTERVAL
13	PHAR3	48	9	63%		48	4.31	1.22
14								COMBINED MEAN STUBBLE HEIGHT
15	POPR	18	4	72%				8.51
16						Bootstrap these data		
17						Sample Mean	8.51	
18						Bootstrap mean	7.50	

There are 5 key species listed in the sample. Three of these, AGST2 (redtop), CYDA4 (thistle cholla), and MG (mesic grass) are represented by just a single sample collected in the field.

[Note: CYDA4 is a cactus, thistle cholla, and not a key graminoid. The correct species code was likely CYDA (*Cynodon dactylon*, bermudagrass). This error would have been detected when running the “CHECK FOR ERRORS” data correction macro and emphasizes the need to run all macros in the correct sequence.]

Two key species, PHAR3 (reed canarygrass) and POPR (Kentucky bluegrass) were sampled at 48 and 18 respectively. Since these represent the bulk of the samples, the inclusion of AGST2 and MG would not be consistent with the direction to use plants that are “relatively abundant”. Notice the difference in average stubble height between POPR and PHAR3 (5 and 9 inches respectively). Using just PHAR3 as in this graphic produces a mean of 8.51

inches with a N of 48 and 95% CI of 1.22. The addition of POPR to the metric produces the following:

1	Note: For proper execution, run the "Generate plant list" macro before using this analysis				Enter species code(s) below to view statistics (up to 4 species)			
2								
3	STUBBLE HEIGHT ANALYSIS							
4	STUBBLE HEIGHT		*Proportion of "Y"s		Species*		* leave these cells blank to use all species collected in t	
5	Key Species	N	Avg Height (in)	% Grazed*				
6	AGST2	1	2	100%	1	PHAR3		
7					2	POPR		
8					3			
9	CYDA4	1	4	0%	4			
10								
11	MG	1	5	0%	BASIC STATS (before bootstrap and spatial analysis)			
12					N	STANDARD DEVIATION	95% CONFIDENCE INTERVAL	COMBINED MEAN STUBBLE HEIGHT
13	PHAR3	48	9	63%	66	4.11	0.99	7.48
14								
15	POPR	18	5	72%	Bootstrap these data			
16								
17					Sample Mean	7.48		
18					Bootstrap mean	7.50		

The effect of adding POPR is to reduce the overall mean stubble height from 8.51 to 7.48 with a N of 66 and 95% CI of .99. The observer may want to report both stubble heights, 9 inches for PHAR3 and 7.5 inches for the combined PHAR3 and POPR. Note that POPR had more grazed plants than PHAR3 (72% compared with 63%) and may be more palatable.

This tab provides the opportunity to execute bootstrap and spatial statistics for the chosen stubble height key species. As shown in the following screenshot of this portion of the "STHT" tab, there are two macro buttons, one for bootstrapping the data and the other for executing the spatial analysis. The user has only to click on each of these buttons to run the analysis.

[illegible]

Running the spatial analysis produces a table exactly like the spatial autocorrelation table in the “Spatial” tab. In fact, that table is simply copied here for convenience. Note that there was spatial autocorrelation for adjacent sample points (the column “Significant” has “Y” for that scenario). These data were collected in 2021 when the standard sample-point spacing (plot spacing) was 2.5 m, thus every other sample point is located 5 m apart. At that spacing autocorrelation is eliminated and the final statistics for stubble height are mean stubble height of 4.5 inches, a sample size (N) of 20, and a 95% CI of .71.

DMA: Bear Creek		Link to INSTRUCTIONS		 Select indicator		
INDICATOR:		Adjacent sample points	Every other sample point	Every third sample point	Every fourth sample point	Every fifth sample point
Seral status	Left bank correlation coefficient	0.0549	-0.1366	-0.2050	-0.2514	0.0562
	Right bank correlations coefficient	0.2625	0.0340	0.0460	0.1131	0.3889
	N	41	40	39	38	37
Sample point	80	Seral status	Seral status	Seral status	Seral status	Seral status
1	50.0					
2	80.0	50				
3	80.0	80	50			
4	32.0	80	80	50		
5	74.0	32	80	80	50	
6	32.0	74	32	80	80	50
7	32.0	32	74	32	80	80
8	32.0	32	32	74	32	80
9	62.0	32	32	32	74	32
10	56.0	62	32	32	32	74
11	68.0	56	62	32	32	32
12	50.0	68	56	62	32	32
13	20.0	50	68	56	62	32
14	80.0	20	50	68	56	62
15	80.0	80	20	50	68	56
16	74.0	80	80	20	50	68
17	59.0	74	80	80	20	50
18	37.5	59	74	80	80	20
19	50.0	37	59	74	80	80
20	80.0	50	37	59	74	80
21	74.0	80	50	37	59	74
22	68.0	74	80	50	37	59
23	32.0	68	74	80	50	37
24	20.0	32	68	74	80	50
25	68.0	20	32	68	74	80
26	56.0	68	20	32	68	74

Figure 22. The Spatial analysis table showing correlation coefficients for adjacent, every other, every third, every fourth, and every fifth sample point. This example is for seral status at Bear Creek.

As shown, this table produces correlation coefficients for adjacent sample points, every other sample point, every third sample point and so on to the fifth sample point (figure 22). Each sample point is separated by the sample point spacing used in the survey and as recorded in the “Header” tab (cell C6). The user clicks on the “Select indicator” button to run the analysis. This macro goes directly to a table of indicators as shown below (figure 23).

LIST OF INDICATORS - SELECT ONE	
Plot Seral status	Select
Plot Wetland Rating	Select
Plot Veg Stability	Select
Average Plot Woody Use	Select
Bank Stability Rating	Select
Bank Alteration Rating (hits per plot)	Select
Stubble Height	Select
Greenline-greenline width	Select

Figure 23. Spatial autocorrelation macro table. The buttons on this table are selected to run spatial autocorrelation for the indicated metric/indicator.

The module is limited to automated analysis of this specific list of indicators for spatial autocorrelation. However, the user can choose any other indicator not on this list and manually enter the data into column B – the blue colored cells and derive the same outputs. In the automated procedure, upon selection of the indicator, spatial analysis is processed for that indicator. This produces a correlogram like the one shown in figure 24:

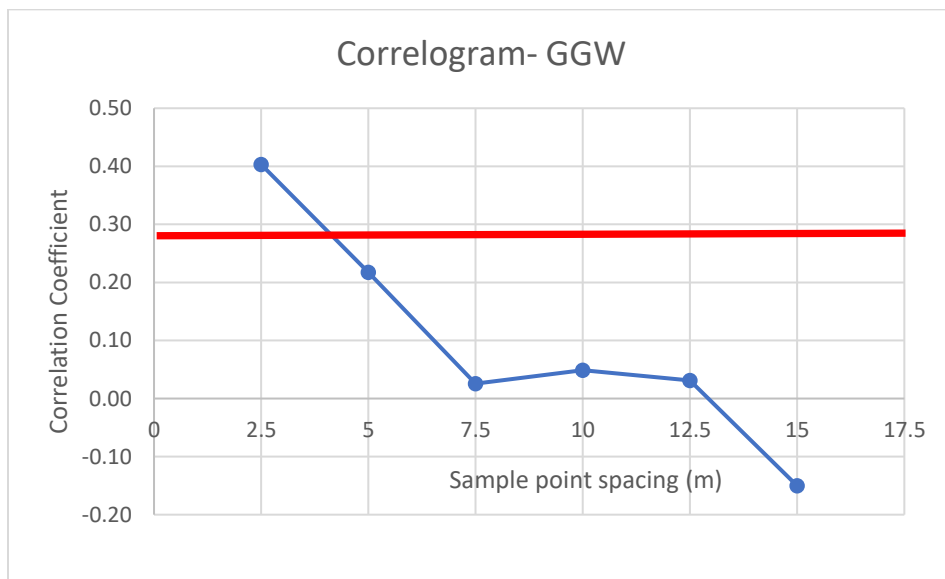


Figure 24. Correlogram of greenline-to-greenline width (GGW) data showing how the correlation coefficient changes with each interval of sample point spacing. If declining with increasing sample point spacing, then spatial autocorrelation may be suggested. The first sample point above the red line is spatially autocorrelated. The red line represents the results of the t-score test of significance which in this case has an r value of .28 ($p < .05$).

The correlogram displays the correlation coefficient by sample point (plot) spacing. This supplies an estimate of the presence of spatial autocorrelation. Typically, correlation coefficients greater than 0.3 are statistically significant (using a T-score test). Another evidence of spatial autocorrelation is the shape of the line on the correlogram. If the line on the correlogram is decreasing with increasing sample point spacing, spatial autocorrelation may be indicated. This is because sample points found spatially closer together are correlated more highly than those spaced farther apart. A scatter plot (figure 25) is supplied so that the investigator can assess spatial clustering of the data. As data appear clustered on the plot, that may suggest spatial dependence.

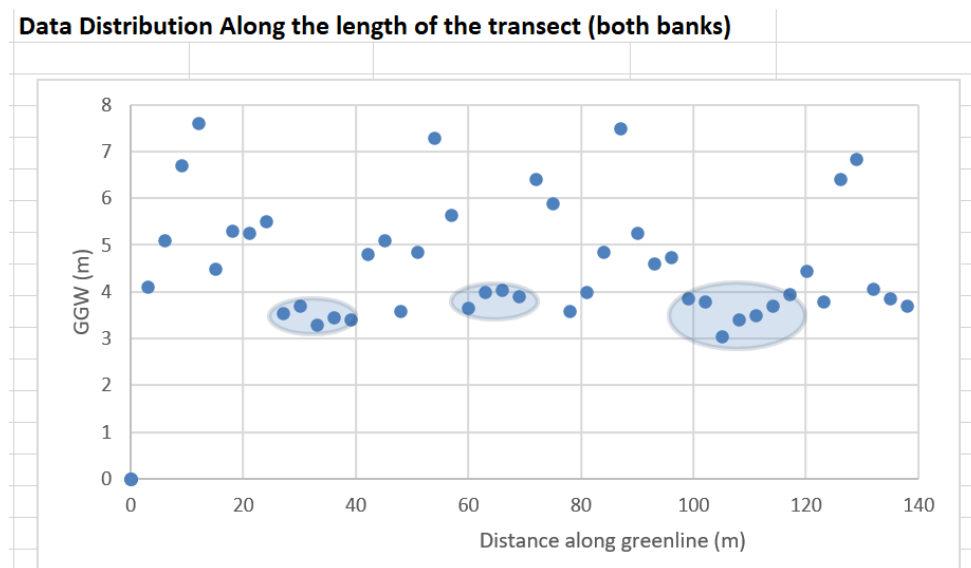


Figure 25. Clustering of samples with distance along the greenline may suggest spatial dependence. Several adjacent sample points cluster (shown in shaded light blue ovals) in this example for GGW.

Note that for these data there appears to be minor clustering of sample values, in this case for GGW, which relates to the moderate levels of correlation seen between adjacent sample points.

A summary table is supplied (figure 26) to indicate the distance at which the correlogram is indicating independence of samples (the interpolated distance) at which the correlation coefficient is below 0.3 (insignificant).

Sample point Spacing:	4.00
Interpolated distance (m)	

Figure 26. Summary table suggesting the sample point spacing at which spatial autocorrelation is negligible.

Instructions for using the “Spatial” tab are provided starting at cell Z7.

Finally, the correlation coefficients for each side(streambank) of the DMA, left and right are tested for significance using the t-score test as described in Statology at:

<https://www.statology.org/p-value-correlation-excel/>

The test results are summarized as follows (figure 27):

Significance test using the t score					
The null hypothesis (Ho): The correlation between the two variables is zero. (reject if p value<.05)					
The alternative hypothesis: (Ha): The correlation between the two variables is not zero, e.g. there is a statistically significant correlation.					
Adjacent Sample points	r	t score	p value	Significant?	Reject Ho for adjacent sample points (y/n)?
Left side	0.074	0.450	0.655	N	Ecological status
Right side	0.290	1.845	0.073	N	Wetland rating
Every other Sample point					Veg Stability
Left side	0.1326	0.813	0.421	N	Woody Rip Spec Use
Right side	0.3739	2.452	0.019	y	Bank stability
Every third Sample point					Bank alteration
Left side	0.1546	0.952	0.347	N	Stubble height
Right side	0.2030	1.261	0.215	N	GGW
EQUATIONS FOR THE SIGNIFICANCE TEST					If 'Y' then spatial autocorrelation is likely present
t	=r*sqrt(n-2)/sqrt(1-r^2)				
p	=T.DIST.2T(t,n-2)				
Reference: Statologyj					
https://www.statology.org/p-value-correlation-excel/					

Figure 27. Test for significance of the correlation coefficient using the t score.

The significance test table to the left (figure 27) describes the results for one indicator (in this case ecological status). Note that all the correlation coefficients (r) are less than 0.30, the approximate r value that is significant based on the typical MIM sample size of 40 for each bank of the stream. The nested table to the right summarizes the test results for the full list of indicators. For this sample, spatial autocorrelation was found for greenline-to-greenline width (GGW), which in this case the alternative hypothesis is accepted - that the correlation between adjacent plots is statistically significant.

If spatial autocorrelation is likely to be present, there will be a “Y” in the “Significant?” column.

Figure 28. Mean, standard deviation, and confidence intervals for all sample points (both streambanks) collected at a DMA. Similar tables are provided for subsets of the DMA including every other sample point, every third sample point and on either the left or right banks.

Figure 29 Spatial autocorrelation table for Stubble Height showing statistical results for various scenarios with those NOT having autocorrelation on the left and all values regardless of autocorrelation on the right.

samples. This produces a mean stubble height of 16.7 with a 95% CI of 1.66 inches. The table on the right shows what would have been calculated for all samples regardless of autocorrelation. It shows that there were a total of 54 samples for stubble height with a mean of 15.6 inches and 95% CI of 1.17. A smaller 95% CI would be expected for a larger sample size of 54, still the difference between stubble heights for the two scenarios is not great.

The results for the non-spatially autocorrelated statistics (row with red numbers in figure 29), are then sent to a summary table as shown in figure 30, and these values are then automatically uploaded to the “Data summary” tab.

RESULTS having no spatial autocorrelation SENT TO DATA SUMMARY TAB				
INDICATOR/METRIC	CONFIDENCE INT*	METRIC VALUE	N	
Ecological status	4.21	98.01	39	
Wetland rating	5.36	78.47	68	
Veg Stability	0.25	7.98	39	
Woody Rip Spec Use	0.12	10.06	16	
Bank stability (%)	6.9%	91.0%	67	
Bank alteration (%)	4%	10%	67	
Stubble Height	1.66	15.69	24	
GGW (m)	0.31	3.73	66	

Figure 30. Results of the spatial autocorrelation tests reflecting the scenarios having the greatest possible sample size (N) not likely to have spatial autocorrelation. Note the stubble height sample in the second to last row matches the data in figure 29. All these data are automatically sent to the “Data Summary” tab.

As shown in figure 30, statistical values associated with the highest sample size likely to not produce spatial autocorrelation are presented. These values are then automatically sent to the “Data Summary” tab. The “Data Summary” tab will not present the most accurate 95% confidence interval until this spatial analysis has been executed. The user can run all macros at the same time by clicking on the Spatial button on the “Instructions” tab or run indicators one at a time (figure 23) using the macros for those provided in the Module.

In rare circumstances there is no scenario that avoids having spatial autocorrelation. In those circumstances, the table in Figure 30 will provide the original metric value from the entire data set and the confidence interval to revert to the 95% confidence interval derived from observer variation on the indicator in question. This approach is used as a last resort when spatial autocorrelation cannot be avoided. In this case the entire DMA is the sample, a sample of one represented by the metric value. A discussion of the use of this confidence interval is provided in Appendix F.

11. "CORREL" TAB:

This worksheet has a standard correlation table for 7 important indicators. The correlation table extracts data from the compilation of plot data in columns AT to BF of the "Comments" tab. An interpretation table (figure 31) shows the strength of the relationship between indicators. The correlation coefficient is a measure of the lineal statistical relationship between two variables and has two directions – negative or inverse, and positive. An inverse relationship indicates that as one variable increases the other decreases (i.e., if bank alteration increases then streambank stability decreases). A positive relationship indicates that as one variable increases the other also increases (i.e., if stubble height increases, streambank stability also increases).

INTERPRETATION
If $r = +.70$ or higher Very strong positive relationship
+.40 to +.69 Strong positive relationship
+.30 to +.39 Moderate positive relationship
+.20 to +.29 weak positive relationship
+.01 to +.19 No or negligible relationship
-.01 to -.19 No or negligible relationship
-.20 to -.29 weak negative relationship
-.30 to -.39 Moderate negative relationship
-.40 to -.69 Strong negative relationship
-.70 or higher Very strong negative relationship
Source:
http://faculty.quinnipiac.edu/libarts/polsci/statistics.h
tml

Figure 31. Correlation coefficient interpretations showing the strength of the association. Correlation coefficients less than +/- 0.2 are considered as negligible or no correlation, and less than +/- 0.3 are often not statistically significant based on the T-score test from sample MIM data sets.

12. "GRAPHS" TAB

This worksheet has a table of plants with associated statistics and graphs. It is the same as in the Data Entry Module described above ([link to Data Entry Module – GRAPHS tab](#)). To the far right in columns CL to EJ are statistical analyses for the short-term indicators, stubble height, streambank alteration, and woody riparian species use, as referenced from the "Data Summary" tab. These are unique to the Data Analysis Module. They are used to assess whether

the data fit a normal probability distribution. Streambank alteration and woody riparian species use typically do not fit such a frequency distribution, but stubble height usually does. Figure 32 shows a typical distribution for stubble height data with just a weak positive skew, while the distribution for streambank alteration data (figure 33) shows a strong positive skew. In the first case the mean would be appropriate, in the second the median would be more appropriate. For calculation of the 95% confidence interval, the standard normal coefficient would be proper for the stubble height data. For the non-normal streambank alteration data, conversion based on resampling (using bootstrapping) would be applied. The calculated 95% confidence interval shown is based on this conversion for the streambank alteration data. In the streambank alteration example, the mean and 95 % confidence interval were derived from the bootstrap analysis which derives from a normal probability distribution and could therefore be used to describe the central tendency and associated margin of error (95% confidence interval) around the mean. The Data Analysis module has a routine for bootstrapping non-normal data in the “Boot” tab.

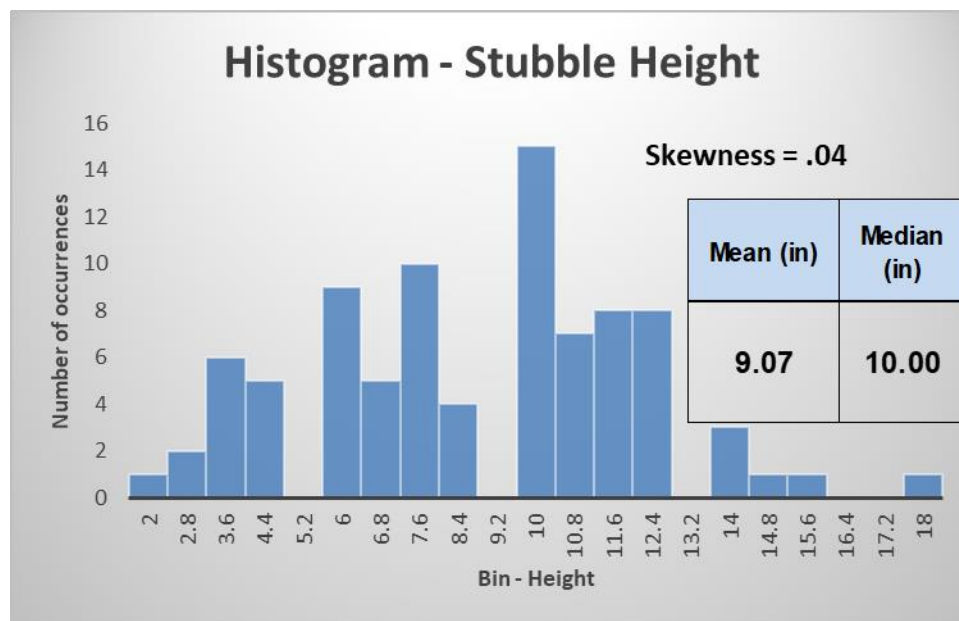


Figure32. Frequency histogram for stubble height showing a normal distribution of the data. Here the mean would be used to describe the central tendency.

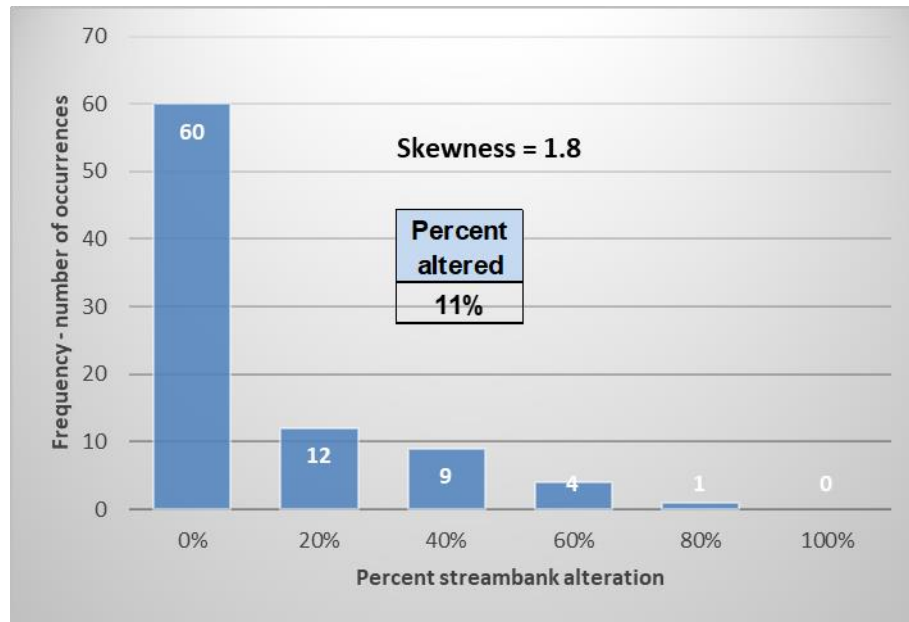


Figure 33. Frequency histogram for streambank alteration showing a strongly positive skew (non-normal distribution). Here the mean and confidence interval are derived from the bootstrap analysis which does fit a normal distribution and thus applies to the mean percent altered by plot.

Here is a general rule of thumb for assessing whether a distribution is normal or not normal using the skewness coefficient:

- If skewness is less than -1 or greater than 1, the distribution is highly skewed.
- If skewness is between -1 and -0.5 or between 0.5 and 1, the distribution is moderately skewed and may be considered normal.
- If skewness is between -0.5 and 0.5, the distribution is approximately symmetric and considered normal.

13. "CODES" TAB

Displays ecological status codes, derived from slope class, substrate class, and riparian capability groups, as summarized in table form. This tab also has tree height and streambank stability codes. Woody riparian species age class codes are also described. These are the same as in the Data Entry Module described above ([link to Data Entry Module – CODES tab](#)).

14. "EXPORT" TAB

Contains two rows of data to be copied and then pasted into an ACCESS database. Also includes a table of plants collected in the survey. The ACCESS database is no longer supported and will not be described in this document. The plant list table is provided as a reference that can be printed. There is a button to the right of the table which when selected and executed will format the plant list table for printing.

15. "CALCS" TAB

Summarizes important arithmetic operations in the workbook. As described above in the "Data summary" tab, this worksheet supplies a description of how each metric is derived. Also, to the right are several metric rating tables that can be used to communicate the general significance of a metric value. For example, if the Winward greenline stability rating is greater than 6.0, this shows a "high" rating (value in cell E8) suggesting that the plants at that DMA are contributing to streambank and channel stability.

16. "PLANTS" TAB

This worksheet has the master plant list used to analyze vegetation. Included are maps of the Plant Regions, the plant list for the DMA, and the Key Species for the DMA. This tab is different than the one with the same name in the Data Entry Module. Rather than a general listing of common riparian plants, as contained in that module, this module contains a list of common riparian plants specific to the plant region for which the module is named (e.g., "Arid West," "Western Mountains and Valleys," or "Great Plains") as described in figure 17. The general format of the table is displayed in figure 34:

WESTERN MOUNTAINS AND VALLEYS								
Species	Scientific - Common Name	Woody?	Hydrophytic?	Herb?	Forb?	Wetland Indicator Status Rating	Plant successional Status	Winward Greenline Stability Rating
ACCI	ACER CIRCINATUM - Vine maple	y	y			50	E	5
ACCO2	ACACIA CONSTRICTA - Whitethorn acacia	y	y			50	M	5
ACCO4	ACONITUM COLUMBIANUM - Columbian monkshood		y	y	y	75	L	5
ACER	ACER SPP. - Maple spp.	y				25	L	5

Figure 34. A portion of the Plants table showing the kinds of descriptor information associated with each plant.

Displayed are the species code (from the NRCS PLANTS database), name (scientific and common), and several characteristics associated with the plant (i.e., woody? hydrophytic? etc.). These characteristics apply to the appropriate plant region, in this case "Western Mountains

and Valleys”. This table is used to derive plant metrics in the analyses associated with the module. A discussion on how to derive these characteristics is contained in the Technical Reference, Appendix G. To the right of this table is a map of the plant regions and a listing of the plants in alphabetical order by common name. There is also a list of less common riparian plants that may occur in the plant region, in addition to the DMA plant list and the key species used in the collection of stubble height and woody riparian species use for that DMA.

Additional plants can be added to the plants table in this tab. Instructions for doing so are contained on the “Instructions” tab. Figure 35 provides a summary of the procedure:

Adding to or Replacing Plant Codes

The two worksheets, "PLANTS" and "KeySP" are designed to facilitate users' changes to the master plant lists. These sheets are not locked, so be careful while making changes that sections are not deleted or lost. You can add or change plant codes, change plant ratings, and re-sort the data. You can examine a plant species characteristics using the Fire Effects Information System at this link:

[FEIS](#)

To sort the data:

1. Place the cursor in the cell immediately below "Species" in the "PLANTS" worksheet (cell A3)
2. Select - "Data" - then "Sort" in the drop-down menu

Figure 35. How to add or replace plant codes in the “plants” tab.

17. "KEYSP" TAB

This worksheet has the master list of potential key graminoid species for stubble height in Column A, key woody riparian species for woody riparian species use in column D, and woody riparian species list for woody age class in Column G. The lists here are used in the macro “Check for errors” when evaluating the plants selected by the observer in the stubble height woody riparian species age class, and woody riparian species use columns of the “DMA” tab. If a species is absent from this list, it will be selected (red circles from the “Check for errors” macro) for any of the indicators. In that case, that plant code can be added to this list at the bottom of the table in either columns A (stubble height), column D (Woody use) or column G (woody age class. Once it is added, the “Check for errors” macro will no longer place a red circle around the plant code on the DMA tab.

18. "PFC" TAB

In this worksheet, data from the module are displayed for each of the PFC assessment items and can be used by the ID team in making or updating PFC ratings for the associated stream reach. Each PFC question is associated with a list of relevant MIM indicators as shown in figure 36. There is a note describing how the indicators apply to the PFC question. The question is then answered “yes” or “no” and whether that item needs attention. The “PFC” tab is provided as a means of quantifying at least part of the results of a PFC assessment, which allows for a more objective determination. In addition, quantitative data provide a more precise tool for determining statistically significant trends, which is not possible with qualitative PFC assessments.

PFC VALIDATION*				ID TEAM RATINGS			
LINK TO DATA SUMMARY				LINK TO INSTRUCTIONS WORKSHEET			
PFC QUESTION	MIM INDICATOR	MIM INDICATOR VALUE	OTHER POTENTIAL QUANTITATIVE INDICATORS	Note	PFC QUESTION	RATING (YES/NO/NA)	Needs Attention
1) Floodplain is inundated in “relatively frequent” events.			Discharge at bankful (cfs)	Survey the cross section then use a cross section program such as: xsecAnalyzerVer13.xlsm. Compute the frequency of Qbf, for example using USGS's StreamStats program.	1		
			Bankful cross section	xsecAnalyzer.xlsm online			
			Flow at bankful	StreamStats USGS - On-line			
2) Beaver dams are stable.					2		
3) Sinuosity, gradient, and width/depth ratio are in balance with the landscape setting (i.e., landform, geology, and bioclimatic region).	Width/Depth ratio GGW/Res Depth	38.76	Pool data in the “Thal” tab are required for this estimate	Width to depth ratio may indicate an imbalance if wider and shallower than the reference DMA or other comparison	3		
	Bed Material - Median Particle	12.65		Median particle size is useful in determining the Channel Type which in turn helps to define the landscape setting.			
	Gradient			Helpful in determining Channel and Valley Type			
4) Riparian area is expanding or has achieved potential extent.	GGW	8.14		Narrowed GGW may indicate widening extent of riparian vegetation. Compare this measurement to a reference DMA.	4		
5) Riparian impairment from the upstream or upland watershed is absent	Width/Depth ratio GGW/Res Depth	38.76	Pool data in the “Thal” tab are required for this estimate	Width to depth ratio may indicate an upland watershed influence.	5		
	GGW	8.1		Wider than expected GGW may indicate upland watershed effects			

Figure 36. A part of the PFC validation table.

19. “BOOT” TAB:

This worksheet contains the bootstrap analysis for the short-term indicators of stubble height, streambank alteration, and woody riparian species use. Stubble height data usually fit a normal probability distribution, but streambank alteration and woody riparian species use commonly do not. All three undergo random resampling with replacement in this tab for mean, median, and 95% confidence interval. The values presented in the “Data summary” tab for these indicators are always the bootstrapped values if the data skewness is greater than 0.5 or less than -0.5. The 1000 resamples in this tab are contained in columns J to ALU and are hidden. As directed in the worksheet, you can select columns I and ALV and then select “unhide” from the drop-down menu to view these data. The bootstrap analysis produces several tables as shown below.

Sample Mean	33.75		Sample Median	10.00
Bootstrap mean	33.75		Bootstrap median	10.00
Sample SD	31.60		Sample SD	31.60
Bootstrap SD	4.45		Bootstrap SD	8.13
Sample 95% CI	8.94		Sample 95% CI	8.94
Bootstrap CI	8.60		Bootstrap CI	10.00
Sample Skew	0.89		Sample Skew	0.89
Bootstrap Skew	0.10		Bootstrap Skew	1.83

Figure 37. Tables in the “Boot” tab displaying mean with its associated SD, CI, and skewness on the left, and median with its associated standard deviation (SD), confidence intervals (CI), and skewness (Skew) on the right. Note that EXCEL displays confidence intervals without $\pm$.

These two tables show the sample and bootstrapped mean, median, and 95% confidence intervals for a single indicator, in this case woody riparian species use. In the following table, the data are summarized for all three short-term indicators, stubble height, streambank alteration, and woody riparian species use. The 95% confidence intervals displayed on this summary table (figure 38) are for the bootstrapped mean.

Bootstrap	Stubble Height (In)	Streambank alteration %	Woody riparian species use (%)
Mean	4.23	14.17	33.54
Median	4.00	0.00	10.00
95% ci	1.09	6.38	9.02

Figure 38. The bootstrap summary table showing bootstrapped mean, median, and 95% CI for the 3 short-term indicators.

Metrics Calculated in the Data Analysis Module

Table 2 summarizes the metrics applicable to each of the monitoring indicators and how they are derived. For more details on these metrics, see the header for each metric in the “Data summary” tab, click on it to go to the “Calcs” tab to view how it is processed within the module.

Table 2. Metrics calculated in the Data Analysis Module on the “Data Summary” tab. Some metrics may also be found in the “Substr” and “Thal” tabs.

Indicator	Metrics in the “Data Summary” tab	Description
Stubble height	Median SH for all key species	Median stubble height value for all key species (in or cm)
	Average SH for all key species	Average stubble height value for all key species (in or cm)
	Average SH for all Key species GRAZED	Average stubble height value for all key species that are grazed
	Average SH for all Key species UNGRAZED	Average stubble height value for all key species that are not grazed

Indicator	Metrics in the "Data Summary" tab	Description
	Dom key species for SH	The most frequent species for which a stubble height is measured
	Avg SH of dominant key species (cm)	The average stubble height for the dominant key species
Streambank alteration	Percent streambank altered	The relative percentage of alteration hits along the greenline as measured on the MIM frame.
Woody riparian species use	Woody riparian species use MEDIAN (%)	The median value of all woody use measurements in the survey.
	Woody riparian species use AVERAGE (%)	The average value of all woody use measurements in the survey. Because woody use is not normally distributed, this value comes from bootstrap resamples of the data.
Greenline composition	Winward Ecological Status rating	Ecological status is calculated using plant successional status ratings and Winward's Riparian Capability Groups. It is further adjusted where a woody overstory component should be present but currently is not present.
	Site Wetland Rating	The average wetland ranking of plants as computed using the Wetland Indicator Status of Reed (1996) Lichvar et. al. (2012) as provided on the "Plants" tab for each species
	Winward Greenline Stability Rating	The average vegetation stability rating of plants as defined in Winward (2000) and provided on the "Plants" tab for each species.
	Woody composition (%)	The percentage of all plants that are designated as "woody." Woody designations are provided on the "Plants" tab in column C.
	Woody composition by plot (%)	The percentage of plots on the DMA spreadsheet that contain a "woody" plant using the designations on the Plants tab in column C.
	Hydrophytic plants (%)	The proportion of plots containing "hydrophytic" plants as defined for each species on the Plants tab in column D
	Hydrophytic Herbaceous (%)	The percentage of plants that are designated as both "hydrophytic" and as "Herbaceous" as defined on the "Plants" tab in columns D and E.
	Percent forbs	The percent of all plants in the sample identified as "forbs" as defined in the "Plants" tab column f.
	Plant diversity index	Calculated by multiplying the number of plant species by average species composition on the plots and dividing by the standard deviation of plant species composition.
Woody species height class	Shade index	The average tree height divided by GGW times the metric "Woody composition by plot".
Streambank stability and cover	Streambank stability (%)	The number of plots classified as "stable" in column X of the "DMA" tab, divided by the total number of plots.
	Streambank cover (%)	The number of plots classified as "covered" in column X of the DMA worksheet, divided by the total number of plots.

Indicator	Metrics in the "Data Summary" tab	Description
	Covered - stable (%)	The percentage of all plots classified as both "stable" and "covered" in column X of the DMA worksheet divided by the total number of plots
	Covered - unstable (%)	The percentage of all plots classified as both "unstable" and "covered" in column X of the DMA worksheet divided by the total number of plots
	Uncovered - stable (%)	The percentage of all plots classified as both "stable" and "uncovered" in column X of the DMA worksheet divided by the total number of plots
	Uncovered - unstable (%)	The percentage of all plots classified as both "unstable" and "uncovered" in column X of the DMA worksheet divided by the total number of plots
Absolute streambank cover	Perennial vegetation	The proportion of the streambank that is perennial foliar cover within .5 meter of the ground surface.
	Embedded rock	The proportion of the streambank that is embedded rock >15 cm in diameter, on the ground surface.
	Anchored wood	The proportion of the streambank that is anchored wood > 10 cm in diameter on the ground surface.
	Bare ground, litter, moss	The proportion of the streambank that is bare ground, litter, and/or moss on the ground surface.
Woody riparian species age class	Percent seedlings or young	The proportion of woody shrubs encountered in all plots classified as seedlings or as young.
	Percent Mature	The proportion of woody shrubs and trees encountered in all plots classified as mature.
	Percent Rhizomatous Woody	The total number of all plants identified as rhizomatous on the "DMA" tab (the woody regeneration portion of the table) divided by all plants identified in the woody riparian species age class section of the "DMA" tab.
	Age class evenness	The Shannon-Weiner index of evenness - based on the relative proportions within each age class. A value of 1 is perfectly even, .8 moderately even, .6 uneven
Greenline-to-greenline width (m) - GGW	GGW (and bankfull width)	Average of width measurements on the DMA spreadsheet for the respective method - GGW or Bankfull.
Substrate	Percent Fines	Ratio: the number of substrate particles less than 8 mm in size divided by the total number of substrate particles sampled (expressed in percent).
	D ₁₆ /D ₅₀ /D ₈₄ Particle Size	The 16th, 50th or 84th percentile of particle size distribution. The D ₁₆ particle size is approximately one standard deviation from the median particle size.
	Substrate habitat analysis	Percent fines, median/ D ₁₆ /D ₅₀ /D ₈₄ particle sizes for pools and for riffles.

Indicator	Metrics in the “Data Summary” tab	Description
Residual pool depth and frequency	Number, frequency, and mean residual depth of pools	As derived from the Thalweg Depth procedure. Includes a count of all pools encountered and their relative frequency or number of pools per mile. Mean residual depth is calculated as the average of all differences between crest depth and pool max depth in the survey.
Woody indicators	Woody Species Frequency (N)	In the "Graphs" tab woody species frequency is calculated for each woody species by summing the "N" values for Woody Species Height, Woody Species Use, and Woody Species Age class displayed in column N of that spreadsheet. In the "Data summary" spreadsheet it is the total of frequencies for all species taken from the total of column N in the "Graphs" spreadsheet.
	hydrophytic woody plant composition	Proportion of plots containing both hydrophytic and Woody plants derived from greenline composition, woody riparian species age class and use.

Processing data from multiple DMAs in the same complex: In complexes having abundant shrub cover, livestock use may vary across the complex and throughout the grazing unit. As such, one DMA may not be completely “representative” of conditions across the unit. In this case multiple DMAs may be more appropriate.

The Data Analysis Module is set up to summarize data from one DMA, thus each DMA would be entered into separate modules for analysis. If multiple DMAs are placed in the same complex and within the same grazing or management unit (i.e., within the same stratum), and because of variability within the complex in question, each DMA can be treated as a sample, and therefore metric means and confidence intervals can be computed for the multiple DMAs. The approach necessitates using the Statistical Analysis Module to accommodate analysis of up to 6 DMAs. A table like the following (Table 3) is included in the “Comp” tab allowing the computation of statistical summaries using each DMA as a sample.

Table 3. Mean and 95% confidence intervals for metric summary data at 6 DMAs in the Spears Meadow riparian complex on Marks Creek, Oregon.

Mark Creek - Spears Meadow complex								
DMA	DATE	Streambank stability(%)	Streambank cover (%)	Greenline Ecological Status Rating	Site Wetland Rating	Winward greenline stability rating	Hydric plants (% by composition)	Greenline-greenline width (m)
DMA 1	8/3/2005	90%	96%	76.79	90.78	7.76	93%	3.55
DMA 2	8/3/2005	86%	87%	61.20	85.32	6.78	94%	3.53
DMA3	8/3/2005	96%	98%	72.49	85.42	7.60	94%	3.09
DMA4	8/3/2005	84%	87%	65.84	85.83	7.53	88%	3.02
DMA5	8/3/2005	83%	95%	69.08	92.05	7.57	92%	3.77
DMA6	8/3/2005	81%	86%	71.07	85.44	7.57	89%	3.05
Average all DMAs		87%	92%	69.41	87.47	7.47	92%	3.34
95% CI (±)		4%	4%	4.0	2.3	0.3	2%	0.2

This approach facilitates investigation of any significant differences or patterns among the DMAs, such as differences in species diversity, abundance, or composition. For example, the confidence interval for hydric plant composition is $\pm 2\%$. The average for all DMAs is 92%. This means that DMAs 4 and 6 would be outside the range of the confidence interval and therefore there is a 95% chance that they are significantly different than the other DMAs with respect to hydric plant composition.

Uploading and Correcting Historic Data in the Data Analysis Module

To ensure that historical data are uploaded correctly into the Data Analysis Module, the following steps should be followed.

Step 1. Use the GET DATA macro to upload data from a historical data set. The data should come from a Data Analysis Module. If the historical data are contained within an older Data Entry Module (5+ years), the data may have to be manually copied, pasted, and corrected in the “DMA” tab.

Step 2. Run the CHECK FOR ERRORS macro and look for data with red circles in the data tabs, especially the “DMA” tab. If plant codes were used in the original data set that have since changed, or if new plant codes unique to the DMA were added, those plant codes along with their associated characteristics must be added to the “PLANTS” tab. Just add them to the bottom of the plant list. Also, bank stability feature codes have changed over the years, and the old codes will have to be replaced. Don’t forget to check the “HEADER” tab as some important conventions have also been updated there.

Step 3. Run the CORRECT PLANT COMPOSITION macro. While the historical data were likely corrected for composition, this macro populates an important table used for data analysis elsewhere in the module.

Step 4. Run the GENERATE PLANT LIST macro. This list must be generated before many of the metrics are calculated.

Step 5. Run the SPATIAL ANALYSIS macro. This will populate the tables in the “SPATIAL” tab.

Step 6. Run the BOOTSTRAP ANALYSIS macro. This will check for normal distribution of bank alteration and woody riparian species use data and if skewed, will bootstrap the data to produce a

corrected confidence interval. Metric values for these indicators will not be displayed in the “Data_summary” tab until this analysis has been completed.

Appendix D contains an example from real MIM data showing application of these steps plus more detail on how to run and interpret the bootstrap and spatial analyses.

The Statistical Analysis Module supplies a platform for comparing multiple samples, either at the same DMA through time for trend analysis, or separate DMAs (for example, a representative and a reference DMA in the same complex) for condition analysis. Statistical tests are used in this module to assess the significance of trends, compliance with grazing-use criteria, progress toward achieving a management objective, and evaluation of linkages or correlations between monitoring indicators. Figure 39 shows the “GetData” worksheet (or first tab) where raw data from the DMA(s) of interest are uploaded and filtered so that data reside in columns without row separations. Data from the DMAs uploaded to the module can then be processed for statistical analyses, as shown in figure 40. Note that up to 6 samples can be uploaded to the module.

Figure 39. The “GetData” tab in the Statistical Analysis Module showing how point data are compiled into columns for analysis.

63

Figure 40. Comparison samples collected at the Big Elk Creek DMA showing metric summary data (green cells), 95% confidence intervals (yellow cells), and maximum difference between the DMAs (grey cells).

C. STATISTICAL ANALYSIS MODULE - EXCEL TABS:

1. “INSTRUCTIONS” TAB:

This worksheet introduces the module with explanations and general instructions. It is important to read and follow those instructions to effectively use the module. The same conventions for saving a master file, unblocking security warnings, and enabling macros apply to this module as with the [Data Analysis](#) and [Data Entry](#) modules.

2. “GETDATA” TAB:

This worksheet is used for importing data from up to 6 different DMAs for analysis. Figure 41 displays the general form for the workbook macros on this tab:

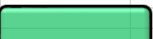
IMPORT DATA			
Click to run the Macro			
Get Data - 1st DMA			
			
Filter data			
			
DMA ID	PASTURE	STREAM	DATE
Dominant Key Species:			
Validate data			
			

Figure 41. The get data macro buttons are used to upload and filter the data.

These macros are for importing and filtering data. The “validate data” macro is only used if after import, the system leaves a blank at the top of the data column for any indicator.

3. “COMP” TAB:

This worksheet has a comparison table for up to six DMAs describing basic summary metrics for each, as described above in figure 38. The 95% confidence intervals (CI) displayed in this table (see figure 40), are the highest of the CIs produced for each sample. This worksheet is especially useful for comparing two or more DMAs, either the same DMA at multiple points in time, or for comparing

separate DMAs to assess condition or variability among DMAs. Chapter III contains a good discussion on the subject of detecting change and testing the significance of that change.

Of interest in this tab is the Table of Confidence Intervals at the bottom of the tab, showing the metric values for all samples and their associated confidence intervals in graphical form. Figure 42 is an example of the output showing change through time at the Big Elk Creek DMA.

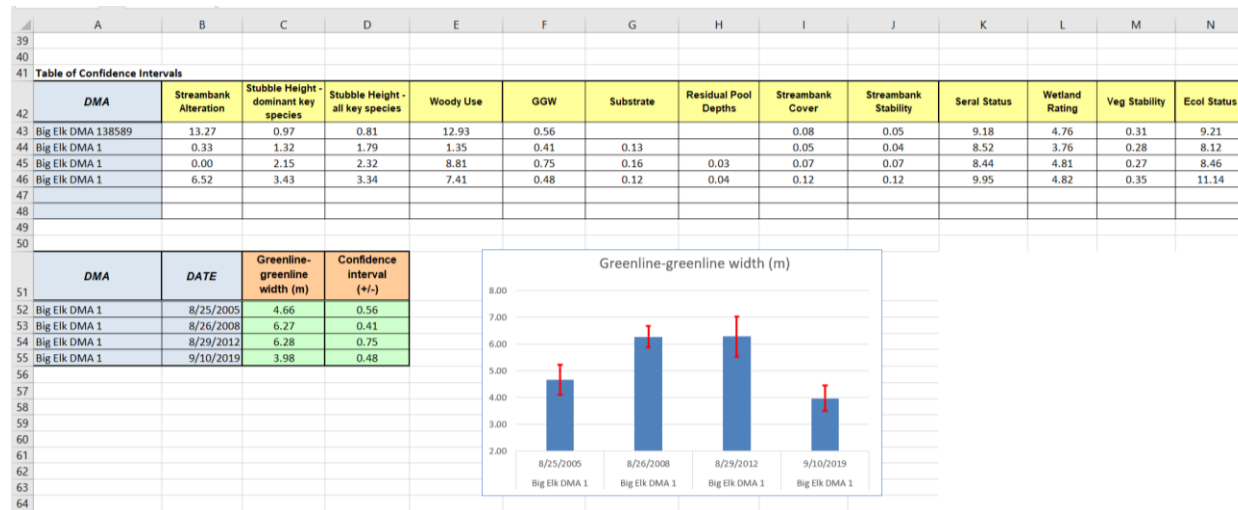


Figure 42. Table of Confidence Intervals showing the margins of error for each of the indicators/metrics. Below is a table used to create the graphic for greenline-to-greenline width. This graph shows trend through time with error bars representing the confidence interval for each year. Where confidence intervals do not overlap, a significant change is indicated.

The remaining tabs in the module are for various statistical tests and analyses. Each tab has an explanation and instructions for using the tab. The following is a brief description of each.

4. "NORMAL PLOTS" TAB:

The user inputs data following the instructions and a plot is produced (figure 43):

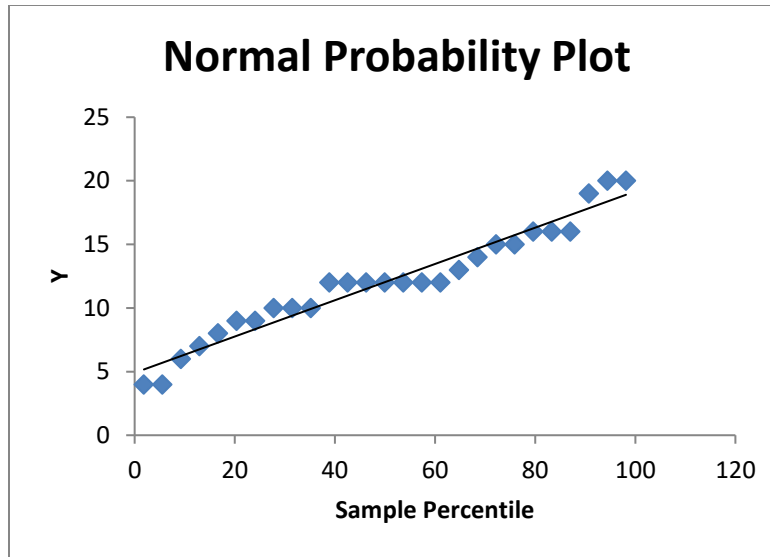


Figure 43. A sample normal probability plot.

According to Elzinga et al. (1998) "If the data come from a normal distribution, the plotted values fall along a straight line extending from the lower left corner towards the upper right corner." These stubble height data are near the straight line, suggesting a normal distribution.

5. "HISTOGRAMS" TAB:

Histograms are another way to assess normality. Following the instructions in this tab, and using the same data as that for the normal probability plot above, produces the following output (figure 44):

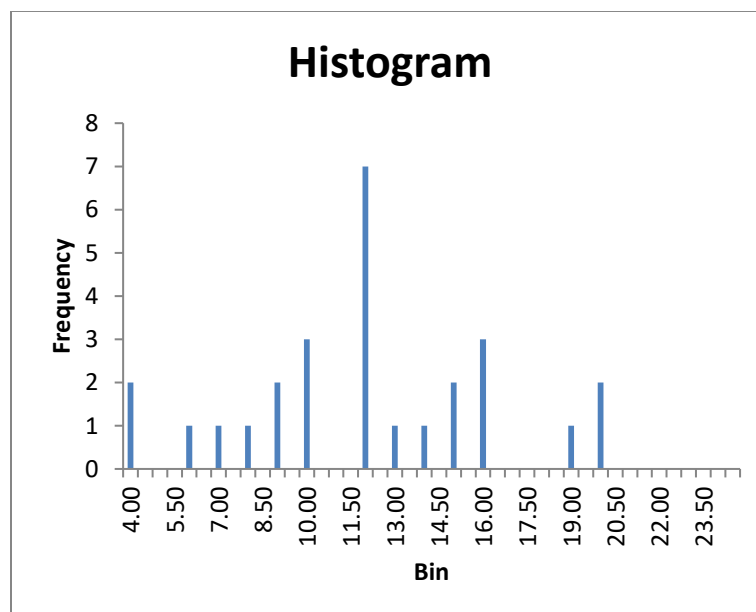


Figure 44. Sample histogram.

Notice that the bars fall generally into a bell shape with the lowest number of values in both the low and high range, and the most frequent values in the mid-range. This suggests a normal probability curve. This tab also displays the skewness coefficient for the data, in this case it is -0.03. If skewness is less than 0.5 and more than -0.5 then the distribution is likely normal ([Link to discussion on skewness/normal distribution](#)).

6. “SPATIAL” TAB:

This worksheet is somewhat like the same worksheet by this name in the Data Analysis Module described above ([link to Spatial in the Data Analysis Module](#)). The difference is that the data to be analyzed must be pasted into the data cells in Column B. Instructions are supplied starting at cell Z7. To the right of the instructions are two tables having the sample point data for each DMA uploaded to the module. These tables will NOT populate unless the data uploaded to the module is derived from a MIM Data Analysis Module 2023 or later. Thus, MIM data collected prior to 2023 must first be uploaded to the Data Analysis Module 2023 (or later). Once done, use the “GetData” (or GetData2) macro to upload the data into this module. Doing so will now populate the table to the right of column AM. ***Spatial analysis, as shown in the instructions, should be executed on data from each side of the stream separately. For this reason, the sample number at which the left side ends is provided in the “Header” tab of the Data Analysis Module (also in the Data Entry Module).*** The user simply copies and pastes those data from sample points 1 through that number (the top or end of the left bank) into the cells in Column B to run the analysis. The analysis is then run for the remaining sample points on the right bank.

7. “SHTERM” TAB:

This worksheet calculates statistics for stubble height, streambank alteration, and woody riparian species browse – the short-term indicators. As shown below (figure 45), the user supplies a stubble height criterion and the worksheet tests for compliance with the criterion using the 95% confidence interval comparison.

STUBBLE HEIGHT (I)									
DMA ID	PASTURE	STREAM	DATE						
BC-01	Bear Meadow	Bear Creek	7/14/2022						
						USER SUPPLIES THIS VALUE			
Stubble Height - dominant key species	Stubble Height - all key species	DOMINANT KEY SPECIES=	JUBA		Grazing use criterion	6			
12	11	Units = Inches							
			STATISTICS	Stubble Height - dominant key species	Stubble Height - all key species				
23	4								
7	3								
18	6		MEAN	11.32	8.94				
23	3		MEDIAN	12.00	7.00				
6	5		CONF INTERVAL	1.12	1.12				
13	12		N	34	117				
13	10								
13	11		CI RANGE						
14	23		MAX	12.45	10.06				
8	25		MIN	10.20	7.82				
24	7								
12	21	EXCEEDS CRITERION?(Y/N)		NO	NO				
16	18								
5	23								

Figure 45. The statistical summary table for stubble height. Similar tables are available for the other short-term indicators.

Note that for this example, the mean stubble height minus the confidence interval of 1.12 (10.2 inches) is well above the criterion of 6.0 inches and therefore the criterion is not exceeded. The upper limit of the confidence interval would have to be less than 6.0 inches to exceed the criterion.

8. "CHANNEL" TAB:

This worksheet calculates statistics for GGW, substrate, pools, streambank cover, and streambank stability. Two DMAs are compared statistically to assess whether they are significantly different using the 95% confidence interval (CI) test with the null hypothesis that the two samples are not significantly different.

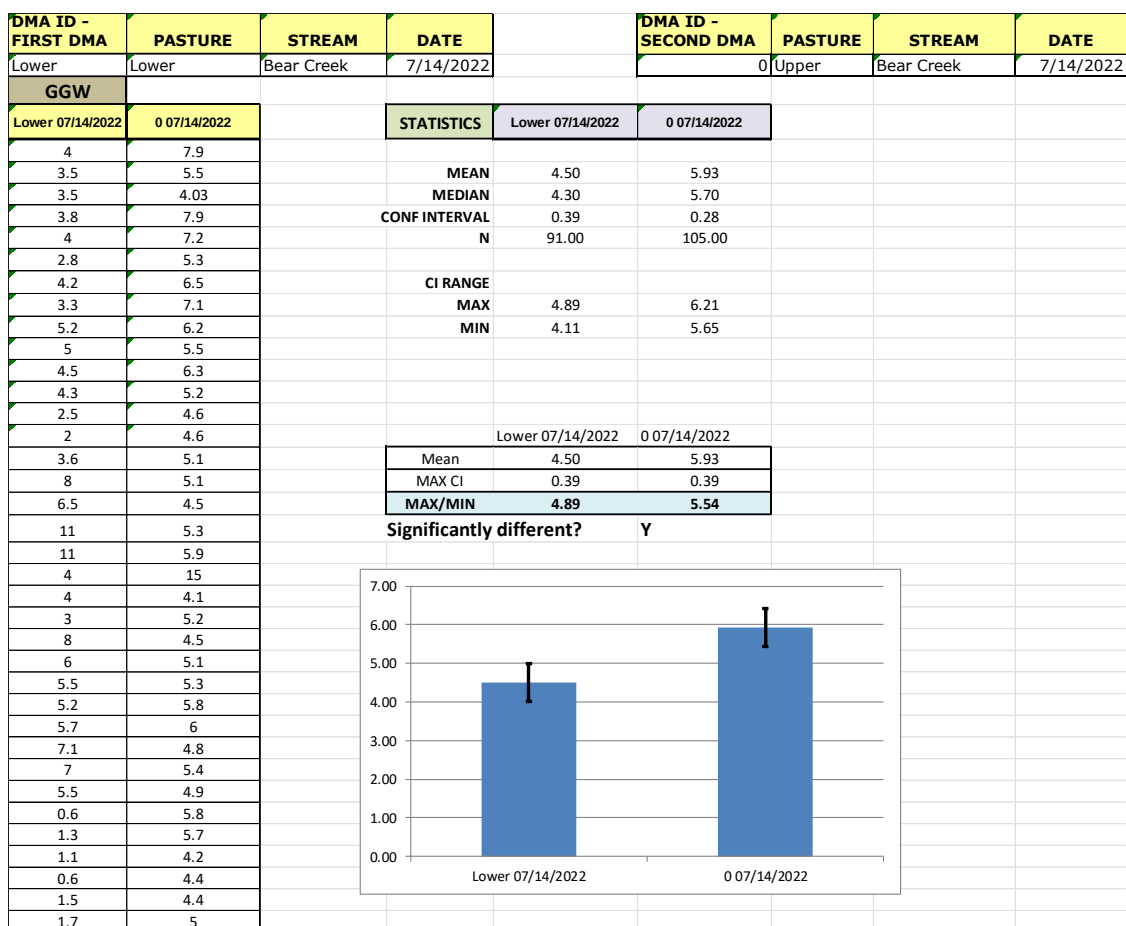


Figure 46. Statistics and graphics for GGW in the “channel” tab.

The worksheet produces a comparison of both mean and median and graphs the two samples with their corresponding error bars (figure 46). The error bars represent the “maximum” CI or highest CI value between the two samples, in this case 0.39 (figure 46).

9. “VEG”TAB:

This worksheet functions identically to the “Channel” tab but calculates statistics for woody frequency, seral status, wetland rating, and vegetation stability. Figure 47 provides an example for seral status:

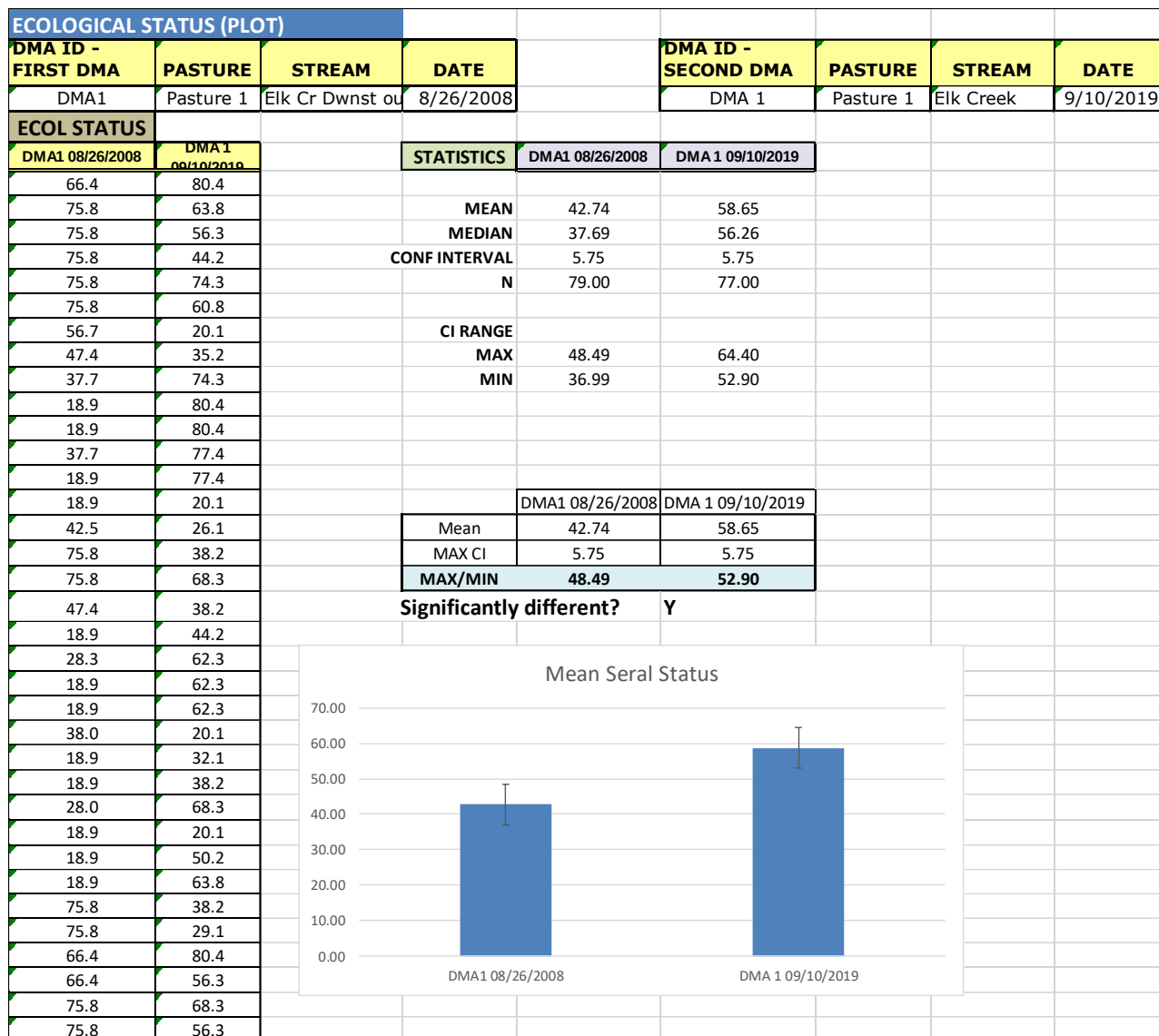


Figure 47. Statistics and graphics for Ecological Status in the “Veg” tab.

10. “TTEST” TAB:

This worksheet runs the student’s t-test for assessing the difference between two data sets using the built-in statistical function available in EXCEL. Follow the instructions and the following output is produced for the same set of data used above in the “Veg” tab.

t-Test: Paired Two Sample for Means		
	Elk 08	Elk 19
Mean	39.39	57.81
Variance	719.93	624.20
Observations	75	75
Pearson Correlation	0.11	
Hypothesized Mean Difference	0.00	
df	74.00	
t Stat	-4.60	
P(T<=t) one-tail	0.00001	
t Critical one-tail	1.67	
P(T<=t) two-tail	0.00002	
t Critical two-tail	1.99	

Figure 48. Results of the t-test in the “Ttest” tab.

This output table (figure 48) shows the mean for each variable, in this case the ecological status for two separate DMAs, and the *t* statistic (t Stat). To reject the null hypothesis that these two are not significantly different (at probability level of 95%), the P value should be less than .05, which in this example is far less than that. The t-test is a powerful statistic and may detect a difference between DMAs when the confidence interval test does not. The interpretation for this test is provided in the “Ttest” tab (figure 49):

INTERPRETATION:		
"Pearson Correlation": Tests correlation between the two samples, this value should be low		
df - degrees of freedom - the number of categories or variables minus 1		
tStat: the Students <i>t</i> statistic (if negative ignore the sign). If t Stat > t Critical, we reject the null hypothesis.		
P(T<=t) one-tail: The probability that Ho (the null hypothesis) is true - should be less than .05 to reject the null hypothesis		
t(Critical) one-tail: The minimum value of the t-Stat needed to reject the null hypothesis at alpha = .05		

Figure 49. Interpretation of the t test.

11. “MANNW” TAB:

This worksheet runs the Mann-Whitney U test for assessing the difference between two data sets for data that does not have a normal probability distribution (figure 50). While these tests are less powerful than the t-test, they may be more powerful than the CI test for data in which the median is more appropriate. Woody riparian species use, for example, is a type of data for which the median is sometimes more appropriate for describing the central tendency.

MANN-WHITNEY U TEST						
Median A=	10.00	Mean A=	16.53			
Median B=	10.00	Mean B=	13.64			
R1=	2397					
R2=	1974					
N1=	49					
N2=	44					
U1=	984					
U2=	1172					
U=	984					
Z=	-0.723319497					
P=	0.46948362					
Accept/Reject Null Hypothesis?	Accept					
The Null Hypothesis: The central tendencies of the two populations are not significantly different						

Figure 50. The Mann-Whitney U test results.

In this example, the null hypothesis is accepted, and the two populations of woody riparian species use data are not considered to be significantly different.

12. “CHISQ” TAB:

This worksheet runs the Chi Square Test for assessing the difference between two sets of categorical type data. The test uses a contingency table that shows the frequencies or counts of the two categorical variables (for example stable or not stable, covered or not covered). This tests the null hypothesis that there is no significant difference between the two samples (e.g. DMA 1 and DMA2). The Chi Square is calculated using the formula shown in the tab. The degrees of freedom are also shown on the tab. The chi-square distribution table (next tab in the module gives the p-value associated with the calculated chi-square statistic and degrees of freedom. If the p-value is less than the significance level (usually <0.05), then reject the null hypothesis and conclude that there is a significant difference between the two monitoring samples. If the p-value is greater than

the significance level, then fail to reject the null hypothesis and conclude that there is no significant difference between the two samples.

The Module uses a Chi-square test for independence. In this type of test, it is appropriate to use both "streambanks stable" and "streambanks unstable" in a Chi-Square analysis, because they represent distinct categories, and all observations fall into one of them. Since their combined frequency equals 100%, this suggests a **dichotomous (binary) variable**, meaning each observation falls into one of two mutually exclusive categories. Two events are mutually exclusive if they cannot happen at the same time. In the MIM, a point on the streambank cannot be both stable and unstable simultaneously, so the categories are mutually exclusive. The data here is working with raw counts instead of percentages, so a **2xN contingency table**, including both stable and unstable categories is appropriate. In this case, the module performs a **chi-square test of independence** or a proportion test to see if stability distributions differ across DMAs or across time.

The following (figure 51) is an example for streambank stability:

CHI SQUARE TEST

Null Hypothesis: Frequency of stability at A is not significantly different than B

Contingency Tables

Stability	A	B	Percentage: Streambank Stability	
	DMA 1	DMA 2	DMA 1	DMA 2
Stable	59	39	65%	51%
Unstable	32	38	35%	49%
Total	91	77	100%	100%

Observed	DMA 1	DMA 2	Totals
Stable	59	39	98
Unstable	32	38	70
Total	91	77	168

Expected	DMA 1	DMA 2	Totals
Stable	53	45	98
Unstable	38	32	70
Total	91	77	168

$$\chi^2 = \sum \frac{(O - E)^2}{E}$$

Chi Square Stability	3.45	Lookup 1 df and alpha of .10	2.71
Reject the null hypothesis	p = .10		
Accept the null hypothesis	p = .05	Lookup 1 df and alpha of .05	3.84

P VALUE =	0.063
-----------	-------

If the p-value is less than 0.05, there is a significant difference in bank stability between sites.

If p > 0.05, there is no significant difference.

Figure 51. Results of the Chi Square test show a case in which the null hypothesis is rejected at p=.10 and accepted at p = .05.

Note that both the 95% and 90% levels of probability are available for this statistical test, and that in this example the null hypothesis is rejected only at the 90% level, or $p=.10$. The calculated p is .063, and since this is greater than .05, then you would accept the null hypothesis that there is no significant difference between the two samples at the 95% level of probability, the preferred level.

13. "REPORTS" TAB:

This worksheet has examples of documentation and graphics that can be used to report results. There are two kinds of reports provided on this tab. The first describes how to develop criteria, such as grazing-use criteria for the short-term indicators. The second includes several examples for assessing condition and trend for long-term indicators.

REFERENCES

- James E. Bartlett, II, Joe W. Kotrlik, and Chadwick C. Higgins. 2001. Organizational Research: Determining Appropriate Sample Size in Survey Research. *Information Technology, Learning, and Performance Journal*, Vol. 19, No. 1, S.
- Lichvar R.W. N.C. Melvin M.L. Butterwick and W.N. Kirchner. National Wetland Plant List Indicator Rating Definitions. ERDC/CRREL TR-12-1. U.S. Army Engineer Research and Development Center Cold Regions Research and Engineering Laboratory Hanover NH
- Matthew G. Johnston, and Christine Faulkner. 2020. A bootstrap approach is a superior statistical method for the comparison of non-normal data with differing variances – *New Phytologist* Volume 230, Issue1 April 2021.
- United States Department of the Interior (USDI). 2015. Riparian area management: Proper functioning condition assessment for lotic areas. Technical Reference 1737-15. Bureau of Land Management, National Operations Center, Denver, CO.

III. TESTING PRECISION AND DETECTING CHANGE

A. INTRODUCTION

Effective monitoring programs address management objectives. These objectives are driven by key questions related to determining whether progress is being made towards meeting the specific, predetermined management objectives. The best way to evaluate whether progress is being made towards achieving the management objectives is to have a monitoring program that detects change through time. To do that, a reasonable level of accuracy and precision in the method of measurement needs to be applied.

B. TESTING THE PROTOCOL

This protocol was tested in the field using a variety of approaches to evaluate its precision in detecting an effect. According to Elzinga et al. (1998, p. 66):

“There are some simple statistical tools that provide a convenient shortcut for evaluating the precision of your sampling effort from a single sample. These tools involve calculating standard errors and confidence intervals to estimate sampling precision levels.”

Standard error is simply the standard deviation divided by the square root of the sample size. The confidence interval is calculated by multiplying the standard error by a critical z value from the table of standard normal coefficients, producing a “margin of error” (ME) which is the confidence interval half-width extending on both sides of the mean or proportion. For example, the Hardtrigger Creek DMA had an average stubble height of 6.6 inches, from 136 samples collected at 80 sample points (quadrat plots). These data produced a standard error (SE) of 0.34 inch and a margin of error (ME) of 0.53 inch. The confidence interval is then 6.6 inches plus and minus 0.53 inch or 6.07 to 7.13. There is a 95% chance that the true mean stubble height occurs within this interval. Thus, precision is related to sample size, natural variability in the parameter (standard deviation), level of statistical significance (or probability of finding the true metric value), and the level of observer bias in the method itself (MacDonald et al. 1991). Because estimates of trend are made at each individual DMA, the principal sources of variability in sampling indicators are associated with measurement error (differences between repeat observations) and spatial variability within the DMA itself. If resampling is not always done at approximately the same time of year, additional sources of temporal variability may be introduced. The MIM protocol has been designed to limit dependency upon streamflow. Still, some short-term climatic influences, such as a sudden cloudburst during the low-flow season, may introduce variation that must be accounted for. A reference DMA, subject to the same kinds of temporal variations, is desirable in the latter case. The following describes how precision was evaluated in testing this protocol.

1. Precision

Precision denotes agreement between repeat observations taken at the same time (e.g., within the low-flow season) and at the same place (e.g., a DMA) to estimate observer variation, and may also be influenced by the tools used to make measurements (e.g., the calibration of the laser rangefinder to accurately measure true distances). There may be broad-scale yearly variations, such as those associated with climatic variances (Larsen et al. 2004), but within-season variations are expected to be minor in the same year because the MIM indicators are not influenced by streamflow at low flow. If an unusual streamflow event should occur within the low-flow season, the data should be interpreted according to such influences. In this case, a reference reach (reaches) would be useful for calibrating the influence of the unusual event. Potential errors associated with relocating the DMA reach (Larsen et al. 2004) could be minimized if the DMA is properly monumented.

Observations may be repeated by the same or different individuals. Differences between samples arise from the bias of individual observers. If bias in sampling occurs, results may be inaccurate (Elzinga et al. 1998). Precision is important for interpreting compliance and trend. If, for example, the stubble height grazing-use criterion is 4 inches and the precision of the measurement is ± 0.96 inches, an observation of 3.6 inches would not imply that the criterion has been exceeded. Similarly, an observation of 4.4 inches would not imply that the criterion has been met, either. With respect to condition, if the objective for streambank stability is to achieve 80 percent stability and the precision is ± 8 percent, an 85 percent observation does not mean that the objective has been met.

Another factor influencing precision is the sample size. Larger sample sizes come closer to the true mean value for the indicator. A good statistic for estimating precision is the confidence interval (or margin of error), calculated from the sample mean and standard deviation, or the sample proportion (i.e., % stable streambanks). More details on the confidence interval are discussed in part 2 below. The sample size needed to achieve desired levels of precision can be predicted by using the standard normal distribution and from field data where the mean and standard deviation are known. The equation is described in Chapter II, Data Entry Module, “Header” tab.

Electronic data entry may be used to assess the margin of error and therefore the confidence interval of the sample size that is being (or was) collected in the field. The user has the option of accepting a lower level of confidence with respect to the sample size with fewer samples and weaker statistical certainty or expanding the length of the DMA with respect to the inadequate sample in question to achieve a higher level of confidence. The former practice described in Burton et al. (2011) of collecting more samples within the DMA is no longer recommended because of the

potential of introducing spatial autocorrelation. As stated by Elzinga et al. (1998) “Data from a pilot study are the most reliable means to estimate the number of sampling units required to meet the targets of precision and power...” (p. 19). The pilot studies conducted for MIM established the desired number of sampling units based on minimizing the confidence interval (or maximizing precision) as shown in figure 52.

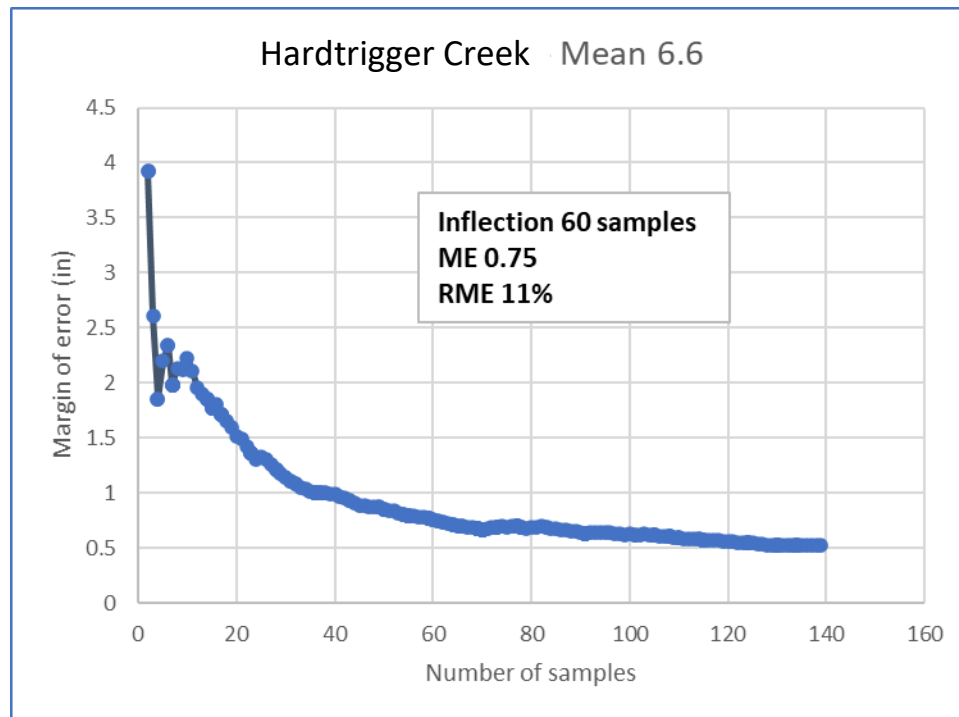


Figure 52. Relationship of precision to sample size using the margin of error (ME). As sample sizes increase, ME decreases to an inflection point or point where ME is no longer decreasing. At this point the confidence interval is minimized. The relative margin of error (RME) is the ratio of ME to mean stubble height.

The MIM protocol uses a Microsoft Excel workbook, the Data Entry Module, designed for use with electronic tablets, laptop computers or similar devices, which allow computation of the sample sizes needed to achieve the desired levels of precision.

2. The confidence interval

The confidence interval can be used to express precision and to test the differences between observations through time to assess trend and condition.

It can be calculated as follows:

$$CI = \hat{Y} \pm Z_{\infty}(\sigma/\sqrt{N})$$

Where:

CI = confidence interval

$\hat{Y}$ = the sample mean

Z_{∞} = the upper critical value of the standard normal distribution, which is found in the table of the standard normal distribution

σ = the standard deviation, and

N = the sample size

Note that as sample size increases, the confidence interval decreases. In other words, the interval is closer to the metric value. With a 95% confidence level there is a 5 percent chance that the confidence interval does not contain the true metric value. With a 90 percent confidence level, there is a 10 percent chance of the same.

A confidence interval supplies an estimate of precision around a sample mean (i.e., mean stubble height) or a sample proportion (i.e., percent stable streambanks) that specifies the likelihood that the interval includes the true value. Confidence intervals have two components: the width of the interval and the confidence level or probability that the interval contains the true value. Because confidence intervals are calculated from the standard error (the standard deviation divided by the square root of the sample size), they are directly influenced by the sample size. As sample size increases, confidence intervals decrease. Because confidence intervals reflect an amount of precision (the probability that the true value falls within its range), we can use them to compare two samples for differences. If one sample, stubble height for example, has a mean of 5.6 inches with a confidence interval of ± 0.5 inch (at a probability or confidence level of 95%), and the stubble height criterion was 6.0 inches, we could not conclude that the criterion was not met since the confidence interval around the mean (5.1 inches to 6.1 inches) contains the stubble height criterion of 6.0 inches. However, if the mean was 5.4 inches with a confidence interval of ± 0.5 inch, the interval: 4.9 inches to 5.9 inches does not encompass the criterion so that we could conclude with 95% confidence that the criterion was not met.

In the same way, confidence intervals can be used to assess trend and/or condition to see if two samples are statistically different – are the means or proportions the same? Basically, are the confidence intervals overlapping? This is one way to see if conditions have statistically and significantly improved or declined or stayed the same over time. While this visual method of assessing the overlap is easy to perform, unfortunately it comes at the cost of reducing the ability to detect differences. CIs may not overlap yet the differences could still be significant. Hypothesis tests such as the 2-sample *t* test are more powerful and may detect such differences.

There are several assumptions associated with the confidence interval statistic. First, the assumption that the sample was randomly selected (independence assumption), second that the standard deviation of the sample is known, and third that there are few or no outliers such that the sample mean fits a normal probability distribution. The assumption of independence includes the requirement that samples are not spatially dependent, basically, that no two observations in a dataset are related to each other or affect each other in any way. Systematic sampling obtains good interspersions of samples along the greenline (Elzinga et al. 1998). The regular placement of quadrats along a transect is an example of systematic sampling. The starting point for the regular placement must be selected randomly. Many natural populations of both plants and animals exhibit a clumped spatial distribution pattern; stream channel morphological units such as meander bends and straight channels also follow this clumped spatial pattern, as do the repeated patterns of pools and riffles or pools and steps, etc. This means that nearby units tend to be similar to (correlated with) each other, a concept referred to as spatial autocorrelation. As stated in Elzinga et al. (1998), "This spacing of sampling units (e.g., quadrats) is needed if one is to treat a systematic sample as if it were random. Indeed, the contiguous placement of quadrats along a transect or the separation of such quadrats by small distances (e.g., one "pace"), practically ensures that adjacent sampling units will be correlated. This will result in an underestimation of the standard error", (and therefore the confidence interval). Spatial autocorrelation is addressed in the MIM protocol and details of the approach are summarized in Appendix A.

3. The coefficient of variation

The coefficient of variation (CV), a dimensionless index of variability between and among observers' repeated samples, has been used to estimate observer agreement (Kauffman et al. 1999, Coles-Ritchie et al. 2004, Heitke et al. 2008, and Roper et al. 2002). The CV is calculated as follows:

$$CV = (\text{standard deviation} / \text{mean}) \times 100\%$$

Where:

CV = the coefficient of variation

Mean = the mean value of the repeat samples

The CV may be expressed as a percentage and represents a proportion of the mean. If the standard deviation is less than 20 percent of the mean ($CV < 20$), then by comparison, a CV of 30 would be less precise. For purposes of these tests, CV values greater than 20 and less than 33 are considered moderately precise, and values less than 20 are considered precise (Kaufmann et al. 1999).

Kaufmann et al. (1999) were interested in detecting change by pooling data across many streams in a region. With the MIM protocol, the observer is more concerned with detecting change at a single

site. For this reason, CV was examined site-by-site and not pooled regionally. ***It is important to note that comparison of CVs should only be done when the samples from which the data are derived use the same unit of measurement.***

Values of CV less than 10 may be required to detect change for variables that may change slowly through time (vegetation erosion resistance, for example). Conversely, values of CV less than 25 may be adequate for detecting change in variables more responsive to management, such as greenline-to-greenline width or percent streambank stability (Archer et al. 2004).

When comparing the precision of data, the choice between coefficient of variation (CV) and confidence interval (margin of error - ME) depends on the context and the specific requirements of the analysis. Both measures have their own strengths and are used in different scenarios.

Coefficient of Variation (CV): The coefficient of variation is a relative measure of variability that compares the standard deviation to the mean of a dataset. It is expressed as a percentage and provides a standardized measure of dispersion relative to the mean. The CV is useful when evaluating the precision of a particular indicator and/or metric.

A lower CV indicates higher precision, as it implies less relative variability around the mean. Conversely, a higher CV suggests lower precision. For more details on the coefficient of variation see: Natrella, M. G. (1996). NIST/SEMATECH e-Handbook of Statistical Methods. Section 4.3.3: Coefficient of Variation.

<https://www.itl.nist.gov/div898/handbook/apr/section3/apr433.htm>

4. Margin of error (ME):

The margin of error is a measure used in statistical inference. It represents the amount of random sampling error expected in the results since only a subset of the population is surveyed, as is the case in DMA sampling. The ME provides an interval estimate within which the true population value is likely to fall with a certain level of confidence (e.g., 95% confidence interval - CI). It considers both sample size and variability.

The confidence interval is typically expressed in a range, such as $\pm 3\%$, where the ME is one side of the range, either the plus 3% or the minus 3%. It indicates the maximum expected deviation of the estimated value from the true value in the population. The ME (CI) is useful in comparing two or more samples for trend and/or condition. Overlapping CIs indicate that the two samples are within the range that likely includes the true population value and are therefore not significantly different.

5. Testing Observer variation

In testing the measurement error associated with the MIM protocol, the confidence interval was used to assess differences between repeat samples (repeatability). The confidence interval quantifies the uncertainty or variability around the estimated parameter, such as the mean, or proportion. It provides a range within which the true value of the parameter is likely to fall. A narrower confidence interval indicates higher precision or less variability, while a wider confidence interval indicates lower precision or more variability.

In the context of repeat samples from multiple observers, the confidence interval can help assess the precision of the estimated mean or agreement measure. One method for comparing the measurements of the same riparian indicator obtained by different teams is to calculate the mean, standard deviation, and confidence interval for each team's measurements. The width of the confidence intervals will indicate the degree of variability among the teams and provide an estimate of the precision of the mean measurements. This approach assumes that each measurement is independent as described above. Such independence requires that samples not be spatially autocorrelated, which was not determined in the original pilot tests. Consequently, a better method for comparing observer differences is to evaluate the variation between observers using DMA summary metrics and computing MEs, CVs, and RMEs from that variation. Figure 53 describes this approach at one of the test DMAs, Big Elk Creek, Idaho.

Big Elk Creek, Idaho												
Site	Mean SH (inches)	Mean Alteration (%)	Mean Woody Use (%)	% Stable Bank	Covered Bank (%)	Percent saplings + young	Percent Mature	Percent hydric	Greenline stability rating	Ecological Status	Site Wetland Rating	Greenline- greenline width (m)
BIGELK3-1	5.1	40.9%	39.0%	36%	60%	40%	60%	54%	6.62	72	84	4.74
BIGELK3-2	5.2	39.5%	40.0%	39%	66%	3%	97%	55%	6.51	68	86	4.59
BIGELK-TR	6.4	37.2%	12.0%	30%	67%	50%	50%	80%	6.60	78	88	4.65
Mean/percent	5.56	0.39	0.30	0.35	0.64	0.31	0.69	0.63	6.58	72.79	86.05	4.66
SD	0.75	0.02	0.16	0.04	0.04	0.25	0.25	0.15	0.06	5.19	2.19	0.07
N	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00
ME	0.84	0.02	0.18	0.05	0.04	0.28	0.28	0.16	0.07	5.87	2.48	0.08
CV	0.13	0.05	0.52	0.12	0.06	0.80	0.36	0.23	0.01	0.07	0.03	0.02
RME	15%	5%	59%	14%	7%	90%	41%	26%	1%	8%	3%	2%

Figure 53. Observer (team) variation of 12 MIM metrics at Big Elk Creek, Idaho, showing CI, CV, and RME for 3 monitoring teams.

Note in figure 50 that RMEs are particularly high (less precise) for woody age class metric - percent saplings/young and percent mature. This is common when woody plants are minimal in numbers at the DMA. Because of this lack of woody plants, each team will typically encounter a different number and condition of plants.

6. Field testing the MIM protocol

Three different field tests were used to assess the precision of MIM metrics. In the first approach, several test sites were established in Idaho and monitored over the years as part of the development of the MIM protocol. Repeat tests between and among the observers were used to evaluate precision at a limited number of sites (5). In the second approach, a much larger sample (30 sites) was obtained involving participants in a number of regional training sessions at many locations in the Western United States over several years. Observers were instructed to repeat observations immediately after obtaining their first sample set to obtain a replicate using the same observers. The instructors would also sample the same reach to evaluate replication by different observers. It has been suggested that this approach may be biased by the fact that observers made the estimates at the same time they were learning the protocol. Such bias may result in better agreement, due to the immediacy of the training, or worse agreement, due to lack of experience with the rule sets. As shown in Table 4, differences between observers were often greater at these training sites than at other test sites. Also, this approach does not account for “revisit variance” across the sample season, which accounts for differences that may result from natural changes during that time period. In particular, streambank alteration and stability may change between the time grazing ends and the end of fall or onset of winter. Thus, the MIM protocol emphasizes the importance of revisits occurring at the same time of year to minimize this environmental noise. The long-term vegetation variables would not be expected to change dramatically during the sample season; however, the ability to identify plants could vary depending upon the presence of reproductive structures.

To further address revisit variance, in the third approach, a more controlled experiment was established to evaluate the variability among trained observers, consisting of three separate, 3-man teams, well trained in the protocols and that visited eight monumented sites (Pacfish Infish Biological Opinion (PIBO) Monitoring Program sites) at different times in the same sampling season (late summer). Teams visited the sites at varying times within a 2-month period during the low-flow season and independently relocated the sites. The advantage of these tests is that opportunities for assessing variation due to site revisits and within-season variations are better. Results of three types of tests are presented in Table 4.

Table 4. Mean Difference Among Observers at Authors' Test Sites, at MIM Training Sites, and at the PIBO Repeat Sites

Metric	Test Sites (5)	MIM Training Sites (30)	PIBO Repeat Sites (8)	All Sites
Average Stubble Height for All Species (in)	0.75	0.88	na *	0.86
Streambank Alteration (%)	10.12	6.21	na	6.76
Woody Riparian Species Use All Species (%)	24.49	5.05	na	8.00
Streambank Stability (%)	8.82	8.16	5.61	8.23
Streambank Cover (%)	10.24	8.29	5.43	8.51
Percent Young	14.27	14.47	9.96	14.44
Percent Mature	15.00	14.18	9.44	14.30
Hydrophytic Plants (%)	10.26	6.22	8.84	6.66
Winward Stability Rating (1-10)	0.97	0.42	0.50	0.48
Greenline Ecological Status (1-100)	14.20	10.51	6.93	10.93
Site Wetland Rating (1-100)	5.46	3.76	4.09	3.94
Greenline-to-Greenline Width (m)	0.26	0.46	0.62	0.43
Woody Composition (%)	4.37	8.95	10.51	6.09
Hydrophytic Herbaceous (%)	11.55	8.00	4.58	9.97
Average Height of Dominant Key Species (in)	1.20	1.47	na	1.44
Percent Fines	2.52	5.09	na	4.72
Median Particle Size (phi)	0.11	na	na	0.11
Pool Frequency (pools per mile)	na	22	na	22
Mean Residual Pool Depth (m)	na	0.01	na	0.01

* na: "not available"—data not collected.

Differences among observers were often greatest at the authors' test sites where indicators were being tested early in the development of the protocols and prior to refinements that included more detailed methods. Such tests were retained in the analysis because they represented only a small proportion of the total number of tests, and observers in these tests were well trained. Training sites also tended to have higher differences as compared to the more controlled repeat sites likely because observers were not as proficient in the methods and other details of the protocols. Precision was also evaluated within and among observers to determine if same observer replicates would be less biased than those of different observers.

Table 5 summarizes the coefficients of variation for same and different observers. Variation was greatest among different observers, as might be expected.

Table 5. Percent Agreement and Coefficients of Variation (ratio of standard deviation to the mean) for Repeat Sampling.

Metric	Agreement for categorical variables	CV - same observers (30)	CV - different observers (33)
Average Stubble Height for All Species (in)		10%	15%
Streambank Alteration (%)		18%	27%
Woody Riparian Species Use All Species (%)		23%	52%
Streambank Stability (%)	81%		
Streambank Cover (%)	85%		
Percent Seedlings + Young		20%	32%
Percent Mature			32%
Hydrophytic Plants (%)	82%		
Winward Stability Rating (1-10)	88%		
Greenline Ecological Status (1-100)	74%		
Site Wetland Rating (1-100)	77%		
Greenline-to-Greenline Width (m)		6%	8%
Woody Composition (%)	79%		
Hydrophytic Herbaceous (%)		17%	38%
Average Height of Dominant Key Species (in)		17%	23%
Percent Fines		7%	6%
Median Particle Size (phi)		5%	23%
Pool Frequency (pools per mile)			13%
Mean Residual Pool Depth (m)			13%

For categorical variables, agreement matrices were also used to estimate observer agreement and differences by comparing rating results among repeat tests of the same DMAs. Table 6 describes the results of that analysis for ecological status. Agreement was good among all categorical variables, 74% to 85%.

Table 6. Agreement Matrix for Ecological Status (units within each cell represent the rating comparison from the test samples, e.g., 16 repeat tests agreed with an early ecological status rating)

Agreement Matrix - Ecological Status					
	Very Early	Early	Mid	Late	PNC*
Very Early	7	1	0	0	0
Early		16	10	2	0
Mid			16	5	0
Late				14	1
PNC*					1

#Tests: 73

Agreement: 74%

*PNC is potential natural community.

All variables are measured or estimated quantitatively at each sample point, including stubble height, streambank alteration, greenline-to-greenline width, woody riparian species use, and substrate particle sizes. At the replicate test sites, monitoring indicators were recorded for the reach and revisits or resamples again recorded the same indicators. Because individual sample locations depend on the randomly selected starting point, revisit samples were likely not located at the same sample location.

7. Testing sample size versus margin of error

At several of the author's test DMAs, more than 100 sample points were selected for sample size analysis. Figure 54 describes graphically the sample sizes at these DMAs in relation to their respective margins of error for stubble height.

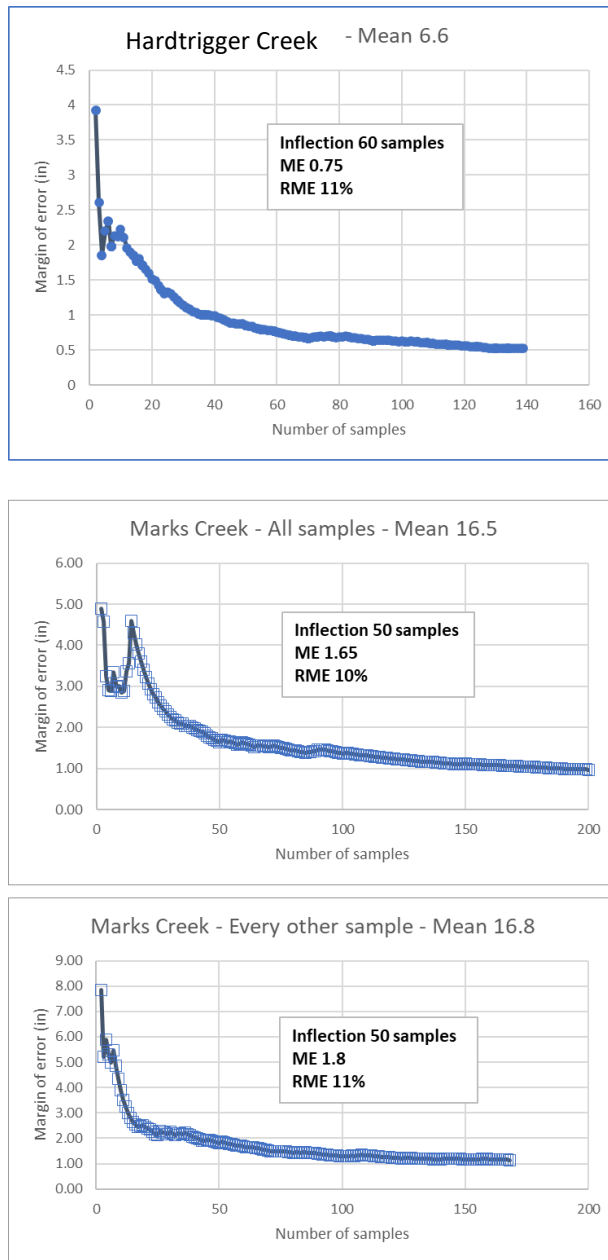


Figure 54. Margin of error versus increasing sample size for stubble height (inches) showing the number of samples at the inflection point as well as the ME and RME at that same point.

Two DMAs, one with low stubble heights and another with taller plants were selected. Analysis of spatial autocorrelation indicated that all samples are independent at the Hardtrigger DMA at sample point spacing of 2.75 m. Adjacent sample points at the Marks Creek DMA were spatially autocorrelated suggesting that the ME was underestimated at 1.65. Every other sample point at a spacing of 5.5 m, however, was not spatially autocorrelated. Here the ME is 1.8 and the RME is 11% and the number of samples at the inflection point was reduced from 60 to 50. Note

that although the sample size was smaller, the mean stubble height did not change substantially.

Figure 55 displays the same graphic for streambank alteration (percent). Two DMAs with differing levels of streambank alteration were chosen for the analysis. There was no spatial autocorrelation at either DMA.

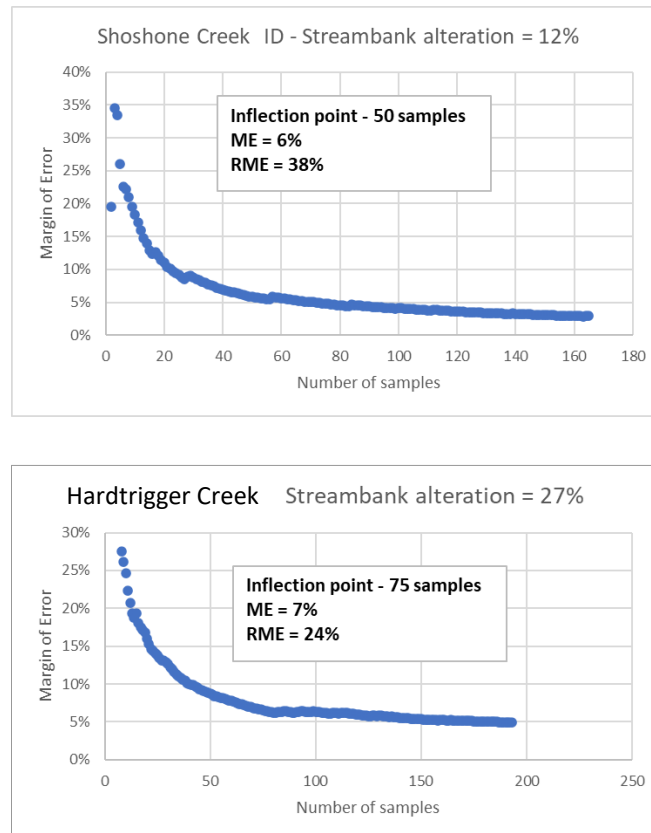


Figure 55. Margin of error versus increasing sample size for streambank alteration (%) showing the number of samples at the inflection point as well as the ME and RME at that same point.

Figure 56 describes the results for woody riparian species use. No spatial autocorrelation was detected for adjacent sample points at these DMAs. Although the Long Tom DMA had much higher levels of woody use, the RME was the same (9%) at both DMAs.

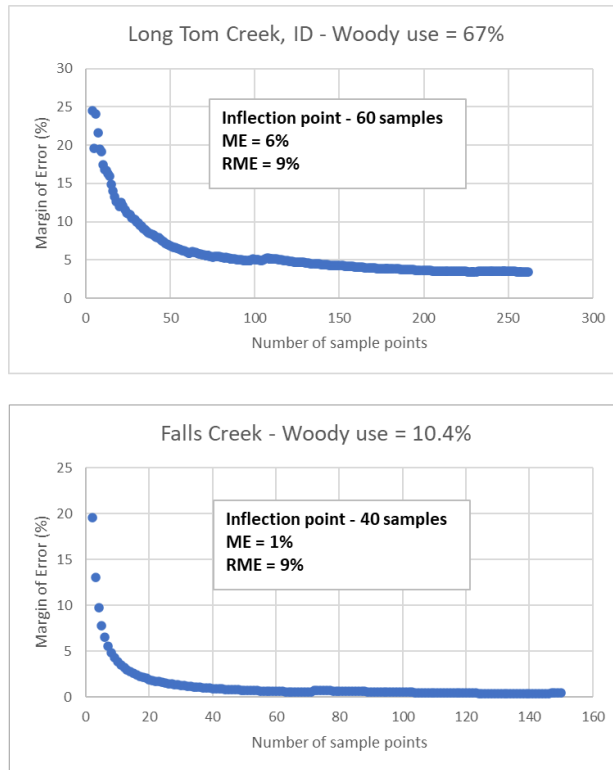


Figure 56. Margin of error versus increasing sample size for woody riparian species use (%) showing the number of samples at the inflection point as well as the ME and RME at that same point.

The sample size analysis for greenline-to-greenline width is summarized in figure 57. The Deadwood River was over 12 m in width, while Hardtrigger Creek is narrow at less than 2 m width. Given that the protocol for GGW was reduced to just 40 samples minimum, this would increase the CI to 0.15 m at Hardtrigger Creek and 1.5 m at the Deadwood River.

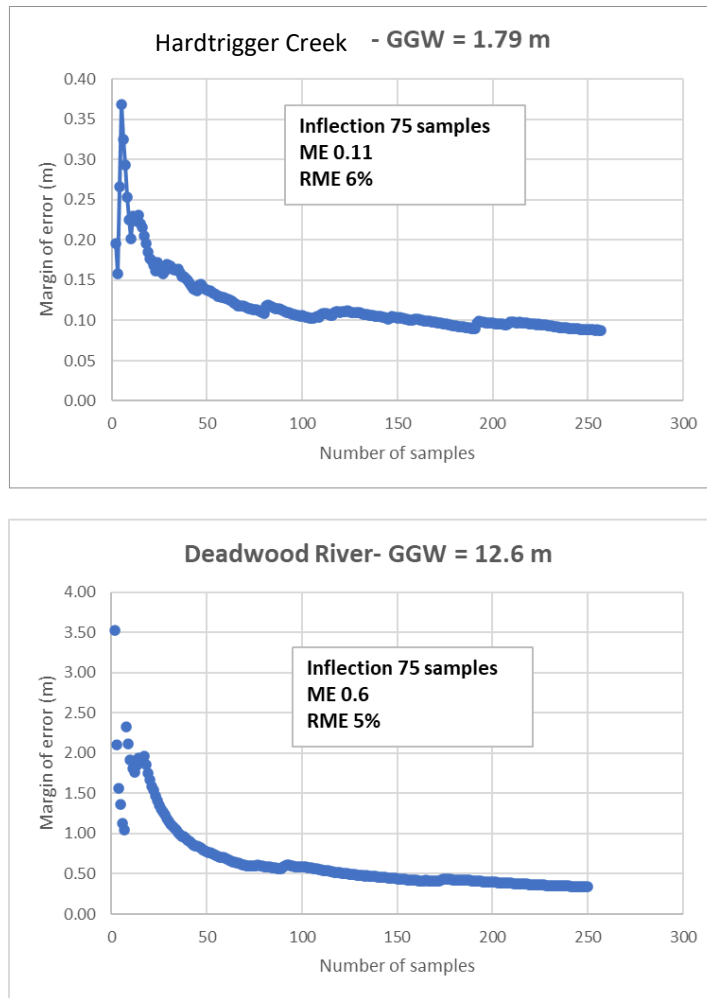


Figure 57. Margin of error versus increasing sample size for greenline-to-greenline width (GGW) showing the number of samples at the inflection point as well as the ME and RME at that same point.

8. Estimating sample size

Bias in statistics is a tendency to underestimate or overestimate the value of a parameter. Bias includes the difference between the population mean of a measured indicator and the mean of a subsample of the population. This source of error may be influenced by the size of the sample. The more samples in the subsample, the closer the mean would be to the whole population mean. Thus, larger samples come closer to the true mean value for the indicator. They produce a lower standard error or standard deviation from the mean. The larger a sample, the closer the resulting mean or proportion is to the true or population mean or proportion and therefore facilitates a better

comparison of samples drawn from separate populations. One statistic for comparing means or proportions of separate populations is the confidence interval, as described above.

A target sample size can be estimated by solving for N in the confidence interval. Thus, at a preselected confidence interval or desired precision (MIM uses the confidence interval for the observer variation in Table 9) and while collecting data in the field at the DMA, the standard deviation can be used to predict the number of samples needed to achieve the desired precision. This comes from the following equation:

$$N = (Z)^2(\sigma)^2/(\beta)^2$$

Where:

Z = the standard normal coefficient

σ = the sample standard deviation

β = the desired precision level expressed as half of the maximum acceptable confidence interval width.

This equation is provided in the Data Entry Module and gives the user some idea of the sample size at which the data gives a good estimate of the metric value. The default precision levels provided in the module are based on confidence interval widths from field tests of sample size and observer variability for the indicators at a 95 percent confidence level.

To incorporate site variability into the assessment of precision, replicate samples were combined to calculate the mean of all observations, and the confidence interval of those combined samples was used to describe range of variability around samples drawn from tests of the same DMA. For example, 4 teams collected data at the Fawn Creek DMA. All combined these teams produced a mean greenline-to-greenline width of 4.2 m from 270 samples with a standard deviation of 1.5 m to produce a combined 95% confidence interval of +/- 0.18 m. This value represents the margin of error associated with all samples collected by multiple observers at the same DMA. Each team visited the DMA at the same time period to minimize temporal variation but on different days to avoid communication between teams and biasing the data.

Table 8 provides an estimate of the number of samples that would be needed to meet the precision level (β) using the standard normal coefficient for the indicators presented. Using the data from this table, it was determined that 80% of the time, 80 samples would be adequate for all but substrate which would require about 200 samples (average of all sites). Thus 80 sample

points and 20 transects with 10 particles each for substrate were established as the default sample size in the MIM protocol.

Table 8. Estimates of Sample Size Needed at the 95 Percent Confidence Level and Precision (β) using the standard normal coefficient (Z) for proportions (Streambank alteration, streambank stability, and woody riparian species use) and for means (stubble height, GGW, and substrate).

Stream	Streambank Alteration	Streambank Stability	Stubble Height	GGW	Woody Species Use	Substrate
Beaver Creek	67	72	23	69	47	311
Big Creek	79	83	56	23	na	206
Darling Creek	125	78	53	102	19	179
Ditch Creek	135	81	51	61	15	242
Trout Creek	69	45	62	68	47	245
Hardtrigger Cr	47	53	29	99	11	384
Lawson Creek	50	55	56	109	12	269
Long Tom Cr	55	13	95	47	5	162
Blanchard Cr	15	22	60	91	19	74
Smart Creek	116	102	9	60	30	184
Telephone Cr	31	39	22	59	81	346
Mill Creek	38	5	28	21	14	149
Burr Creek	50	64	35	47	50	327
Crooked Creek	36	43	21.5	10	60	415
Indian Jack Creek	62	67	40	26	54	131
Little Lost Creek	42	67	37	23		506
Marks Creek	73	57	89	76	81	
Rio Bonito	67	84	56	41.5	50	304
Shoshone	72	85	43	93	50	
Taylor	69	64	30	90	61	432
WF Blacktail Deer	44	69	27	20	56	229
Average	69	54	45	67	27	229
MAX	135	102	95	109	81	384
MIN	15	5	9	21	5	74

Note that sample size adequacy varies considerably by stream. These differences reflect the unique characteristics of diversity associated with each site. Adding additional samples at a site could be done by increasing the DMA length and maintaining the sample interval of at least 3.75 m, not by inserting new samples into the existing sample set at less than the sample interval so as to avoid creating spatial autocorrelation. The Data Entry Module provides an

indication of spatial autocorrelation in the data, allowing the user to adjust the sample spacing if desired. This may cause a decrease in sample size and result in reduction of the precision (i.e., confidence interval) of the sample population.

Some metrics do not fit a normal probability distribution as needed to apply the standard normal coefficient in the confidence interval equations above. Streambank alteration, for example, often has a heavily skewed, non-normal distribution. Samples usually contain no or few alterations, with two or more alterations less common. Such a sample distribution tends to be positively skewed or right skewed. For this reason, it is best to evaluate streambank alteration as a proportion (or percent streambank altered) and then use the confidence interval for a proportion rather than for a mean.

9. Testing observer variation

In the context of repeat samples from multiple observers, the confidence interval can help assess the precision of the estimated mean or agreement between observers. This represents one method to assess the precision of the method or metric. For example, in comparing the measurements of the same riparian indicator obtained by different teams, the mean, standard deviation, and confidence interval can be calculated for each team's measurements. The width of the confidence intervals will indicate the degree of variability among the teams and provide an estimate of the precision of the mean measurements. The relative differences between observers were described in Table 5 above. Because some of the metric indicators were found to have spatial autocorrelation of adjacent sample points, analysis of the margin of error (and confidence interval) was not conducted on the sample point (quadrat plot) data. Analysis of observer variation was conducted at the DMA scale where differences in the metric summary results were assessed for variation between observers' results for each DMA. One such analysis is described above for Big Elk Creek in figure 50. This produced a margin of error (ME), coefficient of variation (CV) and relative margin of error (RME) for each DMA. In this figure, the mean stubble height is 5.56 inches with a ME of 0.84 inch or a confidence interval of: lower bound = 4.7 inches (5.56 minus 0.84), upper bound = 6.4 inches (5.56 plus 0.84). The ratio of the ME to mean stubble height, or RME is 15%. These analyses were conducted and summarized for all test sites. Results are presented in Table 9.

Table 9. Average margin of error (ME), coefficient of variation (CV) and relative margin of error (RME) for tests of observer variation on MIM indicators and metrics.

<i>Metric</i>	ME	CV	RME	Number of tests
Mean SH (inches)	0.83	12%	16%	33
Bank Alteration (%)	6%	25%	33%	32
Woody Use (%)	15%	38%	27%	25
% Stable Bank	8%	15%	19%	40
Covered Bank (%)	7%	11%	13%	39
Percent saplings + young	15%	22%	30%	24
Percent hydric	7%	12%	14%	40
Greenline stability rating	0.34	5%	6%	40
Ecological Status	6.11	15%	19%	40
Site Wetland Rating	3.18	4%	6%	40
Greenline-to-greenline width (m)	0.45	6%	12%	40
Residual pool depth (m)	0.03	13%	14%	8

Of interest is the comparison of these observer variations with that of spatial variation or variation at the inflection point of the sample sizes. For stubble height, ME on spatial variation averaged 1.2 inches, and observer variation .86 inches. For streambank alteration, spatial variation was 6% and observer variation 7%. For woody riparian species use, spatial variation was 3% while observer variation was much higher at 28%. For greenline-to-greenline width, spatial variation was 0.6 m while observer variation was 0.45 m.

Observer variation was analyzed for a much larger set of DMAs across a broad geographical range within all three ecoregions: Arid West, Western Mountains, Valleys and Coast, and Great Plains. For this reason, the MEs for observer variation are used to define the desired precision levels in the Data Entry Module. In addition, by considering both site and revisit variability, these MEs, or ranges of variation, were used to assess expected levels of variability at MIM monitoring sites and help set targets for DMA sample sizes.

10. Displaying results.

The MIM method for detecting changes in long-term indicators or to detect a failure to meet short-term grazing-use criteria is to use a confidence interval around the mean that combines

sources of variability or variations due to site complexity and revisit variance (observer plus within-season variation).

One way to display the results of the statistical analysis is to use bar charts with error bars. The following example is derived from the MIM test site on Big Elk Creek in North Central Idaho. Sampling of greenline-to-greenline width (GGW) occurred in the grazed pasture and in each of two exclosures, one that had been fenced for 10 years and the other for 20 years. Confidence intervals around the mean, are ± 0.3 m. The following graphic (figure 58) describes these results:

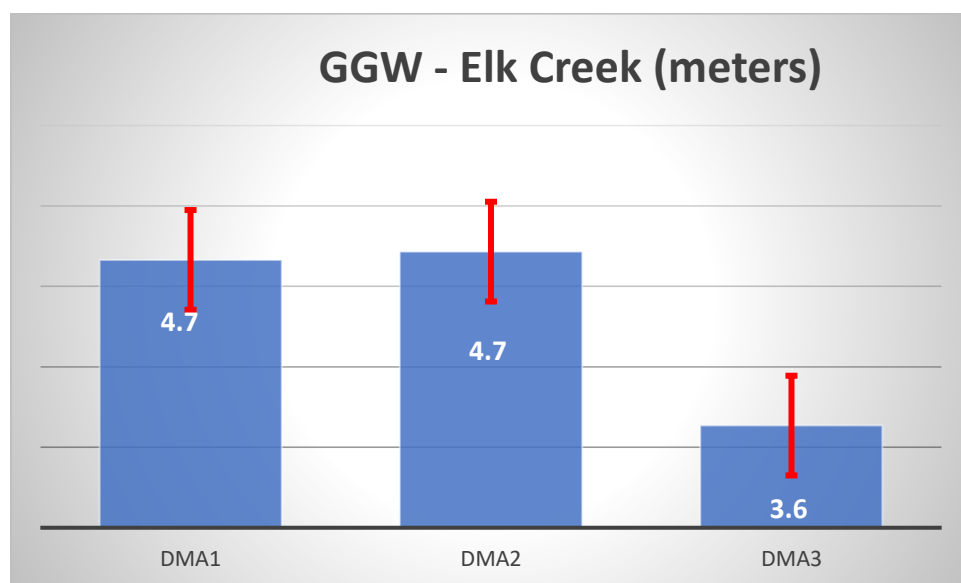


Figure 58. Bar graph of greenline-to-greenline width within the grazed pasture (DMA1) as compared to the 10-year (DMA2) and 20-year exclosures (DMA3)

It is apparent from this graph that the confidence intervals for the grazed DMA (DMA 1) and the 10-year exclosure (DMA 2) overlap, which suggests there is no statistical difference between these two monitoring sites. In contrast, the confidence interval of the 20-year exclosure DMA does not overlap with the other two monitoring sites, which indicates there is a statistically significant difference in GGW between this site and the others.

APPENDIX A – SPATIAL AUTOCORRELATION ASSESSMENT OF THE MULTIPLE INDICATOR MONITORING (MIM) PROTOCOL

The Multiple Indicator Monitoring (MIM) protocol (Burton et al. 2011) uses the confidence interval statistic to estimate precision (margin of error) in tests of significance. Because the confidence interval statistic assumes that each individual sample is both random and independent, the collection of samples should have a sample point spacing large enough to assure spatial independence. Spatial autocorrelation occurs when the values of variables sampled at nearby locations are not independent from each other. When present on a transect, sample points or quadrats in proximity tend to be more similar than distant sample points or quadrats. As stated by Elzinga et al. (1998):

“...adequate spacing of sampling units (e.g., quadrats) is needed if one is to treat a systematic sample as if it were random. Indeed, the contiguous placement of quadrats along a transect or the separation of such quadrats by small distances (e.g., one “pace”), practically ensures that adjacent sampling units will be correlated. This will result in an underestimation of the standard error (and therefore the confidence interval).”

Weixelman and Riegel (2012) assessed spatial autocorrelation of plant species occurrence among quadrats on mountain meadow transects. They found spatial autocorrelation was present, especially in plant communities with low species diversity. More recently Adam Green (personal communication, 2022) found spatial autocorrelation associated with several MIM indicators from 58 designated monitoring areas (DMAs). At these sites sample point spacing was 2.5 to 2.75 m apart. Using Moran’s I (Moran 1950), he found positive correlation for: wetland rating (0.3); greenline-to-greenline width (0.45); streambank alteration (0.4); stubble height (0.35); and slightly less positive correlation for streambank stability.

Methods

Spatial autocorrelation was analyzed at 80 DMAs for nine MIM indicators: streambank stability, streambank cover, streambank alteration (%), stubble height of all key species, woody riparian species use on all key species (%), greenline-to-greenline width, wetland rating, vegetation stability (Winward greenline stability rating), and ecological status. These are indicators that tended to

experience spatial autocorrelation based on the analysis. We did not observe spatial autocorrelation for others such as residual pool depth and percent fines. We analyzed spatial autocorrelation using a technique specifically designed for samples taken in one dimension – along a transect line, as per the MIM protocol. This approach (Ecology Center: <https://www.ecologycenter.us/vegetation-ecology/correlograms-morans-i.html>) employs simple correlation analysis using a correlation matrix (see Table A1), and a correlogram as shown in figure A1. Somewhat analogous to a time series autocorrelation analysis using correlograms (see: Real Statistics using EXCEL), the correlation matrix compares adjacent sample points, every other sample point, every third sample point, and so on out to every sixth sample point. Because sample points are normally 2.5 m apart, each iteration of distance is multiplied by 2.5 to derive the distance assessed. Thus, adjacent sample points are 2.5 m apart, every other sample point, 5 m, every third sample point, 7.5 m, and so forth out to every 6th sample point at 15 m. Using this kind of analysis, a correlogram is produced as shown in figure A1.

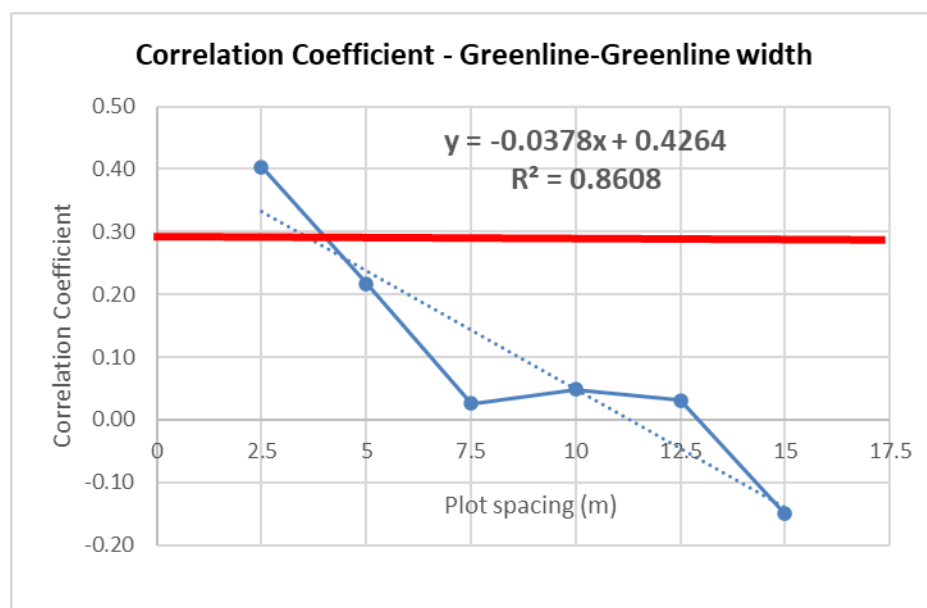


Figure A1. Correlogram for Pacific Creek (2022) showing the correlation coefficient versus sample spacing. A linear trendline is included to help assess sample independence. When the correlation is less than 0.2, the spatial relationship between samples was considered negligible in the analysis. The trend line (with R^2 of 0.86) suggests a spatial correlation between sample points that are closer in proximity. The significance of the correlation was determined using a t -score test at $p < .05$ was at an r value of .28.

As shown in figure A1, there is a declining correlation as spacing increases between samples. Although correlation coefficients are small (less than 0.4), this trend suggests more spatial relatedness among nearby measurements than those further apart. Many plants, animals, and stream features exhibit a clumped spatial distribution in nature. Streams typically have a downstream-ward pattern of alternating curved and straight channel segments that tend to be similar in morphology, such as width, and therefore exhibit this kind of clumped distribution. When the correlation coefficient declines to below 0.2, the spatial autocorrelation between samples is considered negligible (see: <http://faculty.quinnipiac.edu/libarts/polsci/statistics.html>). To estimate the sample point spacing from the graph above, we used both a linear line fitted to the data points and used the regression equation to calculate distance. In this case (where $r < 0.2$), $x = 6.0$ m. We also estimated the distance by interpolating between line points connecting the correlation coefficients from each sample point. The interpolated distance is 5.23 m. Thus, a spacing of 5 to 6 m would result in a negligible spatial autocorrelation between measurements of greenline-to-greenline width.

A somewhat similar approach was used by Myers and Swanson (1997) to assess spatial autocorrelation in estimating the precision of stream widths. Transect cross-sections perpendicular to the streams' centerline were used to make the measurements, a method very similar to the transect method of the MIM protocol. Autocorrelation was assessed using the covariance of transect lags (adjacent, every other, every third, etc.) divided by the variance of the width. This is very similar to the regression coefficient, which is basically the correlation coefficient squared. Myers and Swanson (1997) found that a spacing of 3 channel widths appears to be the minimum spacing for measuring width without spatial autocorrelation. Thus, they recommend 10 transects at a spacing of 3 channel widths.

Upon examination of the 80 DMAs selected for spatial analysis, we found an average of 2.3 channel width (GGW) for spacing of transects at a correlation coefficient of 0.2 or less. The average GGW for these 80 DMAs was 3.0 m, so the indicated spacing would be about 7 m. For 80 percent of the DMAs, a spacing of 3 channel widths would be adequate, which seems consistent with the findings of Myers and Swanson (1997).

In a different approach, Weixelman and Riegel (2012) used semivariograms to assess spatial autocorrelation of plant species occurrence among quadrats on mountain meadow transects. The MIM protocol employs a similar quadrat and transect approach to estimate plant composition, but along stream margins rather than in meadows.

As stated by Weixelman and Riegel (2012):

“Normally, points in close proximity are more similar than points farther apart, so that semivariance among points increases with distance until a maximum semivariance, called the sill, is reached”.

In spatial statistics, the semivariance is described by a datum (z) at a known location. The distance (h) is the distance between ordered data (such as quadrats on a transect line). Then the number of paired data at a distance of (h) is the sample size (n). The semivariance is half the variance of the increments squared, at a given separation distance (h). A graphic showing semivariance versus distance between sample points in a graph is known as a semivariogram. Figure A2 contains a semivariogram for greenline-to-greenline width, showing an increasing trend in semivariance up to approximately 8 m, at which point the trend flattens or declines indicating that beyond that distance samples are spatially independent. The best-fit line, based on the r -squared value, is used to describe the relationship. The best-fit line was determined from linear, logarithmic, or exponential models. The model with the highest regression coefficient (r -squared value) was chosen as that with the best fit.

Weixelman and Riegel (2012) examined semivariograms for transect data and found that three models could be fit to their data (figure A3). They describe these as model “types”. Type A communities had a flat model in which there was basically little or no covariance between sample points (no increasing trend in semivariance; figure A3). Also, with type A models, the best-fit lines had an r^2 value less than 0.2. These sites exhibited no spatial autocorrelation at a distance of 1 meter. Type B communities were positively autocorrelated with a curve like the one in figure A2 that clearly displays a flattening at some distance above 1 meter (the sill; figure A3). Type C communities had a continual rise with no sill visible. These sample points were positively autocorrelated to distances greater than 20 m (figure A3).

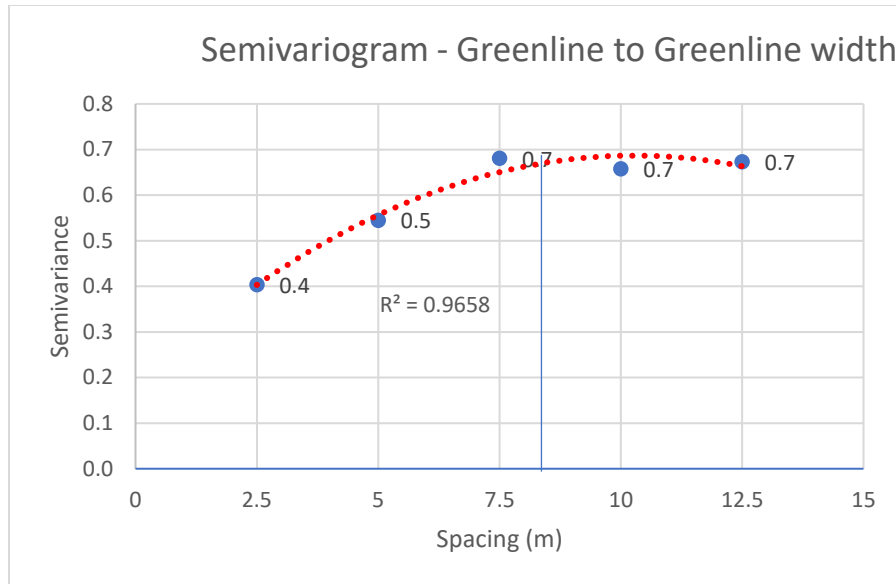


Figure A2. Semivariogram fit to average semivariance at the indicated spacing. Note that the sill (or top of the curve) is reached at a spacing distance of approximately 8 m. This is a Type B model based on Weixelman and Riegel (2012).

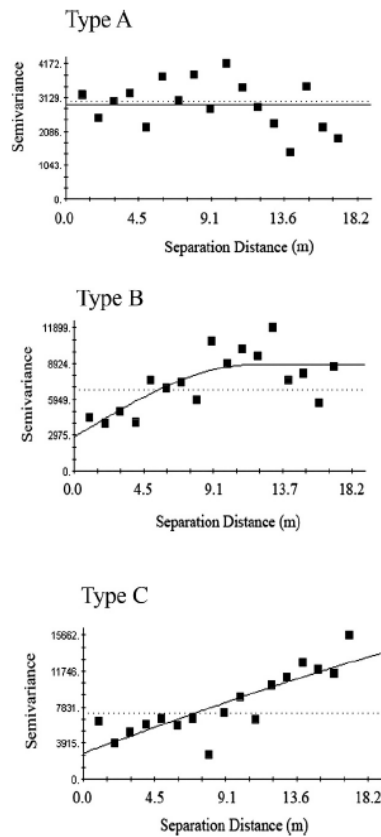


Figure A3. Three semivariogram models from Weixelman and Riegel (2012).

Semivariogram models fit to the data, related well to the correlogram results. This is shown in the figure A4 box-and-whisker plots for all MIM indicators combined, showing that A type models are associated with a median sample point spacing of 2.5 m, B type with 3.7 m, and C type with 4.3 m. Thus, as expected, A types have no spatial autocorrelation at 2.5 m. B and C type medians have higher distances, meaning a larger sampling interval is suggested to avoid spatial autocorrelation of adjacent samples.

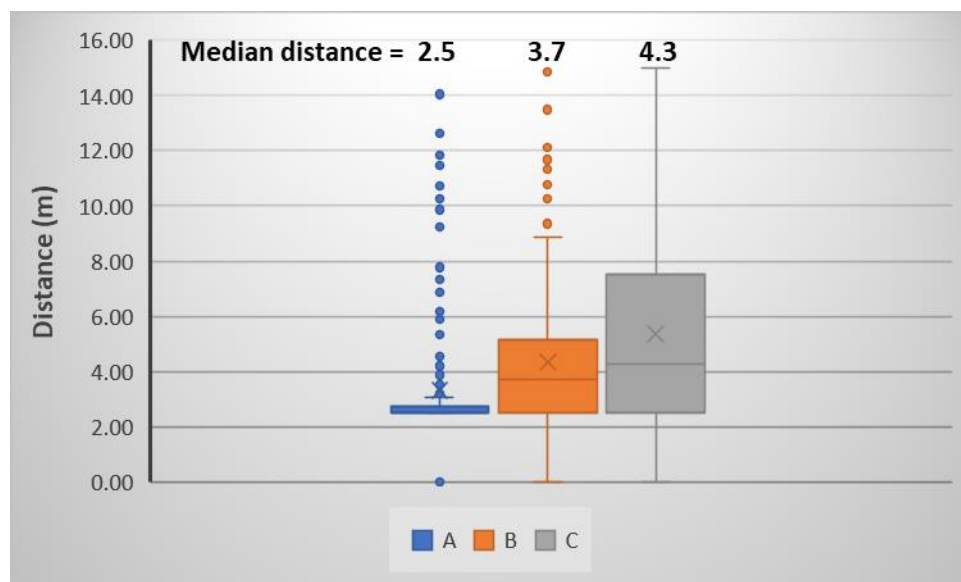


Figure A4. Box and whisker plots for correlogram quadrat spacings for each of the 3 semivariogram models showing the average distance for each model type (derived from the correlation data for all DMAs and MIM indicators combined).

We also analyzed end-of-season, short-term indicators from data collected at 12 DMAs in Nevada in 2015 and 2016 separately (Table A2). These samples were isolated to assess the presence of spatial autocorrelation in short-term indicators at the end of the grazing season. This is because MIM is often used specifically for this kind of monitoring apart from the long-term indicators of riparian and stream channel condition. Also, sometimes the data sets for short-term indicators included very few measurements (i.e., very few alterations or very few browsed woody plants to evaluate woody riparian species use) when evaluated early in the grazing season when long-term indicators are best collected. Therefore, timing of data collection matters when interpreting short-term indicators and for analyzing spatial independence.

Results

For assessment of the MIM protocol, we examined spatial autocorrelation using the correlogram method described for 80 designated monitoring areas (DMAs) located throughout the western United States, from New Mexico to California in the south, Washington to South Dakota in the north, and states in between. Although we also examined spatial autocorrelation using semivariograms, we did not find this approach especially useful for estimating auto-correlation distances. The large spacing intervals (multiples of 2.5 m) generated low-resolution sample points and commonly produced best-fit lines with low r-square values (poor line fit) that often were influenced by a single outlier value. We believe this approach would have been more useful if a higher resolution semivariogram could have been produced using a sampling interval of 1 m for testing purposes, as applied by Weixelman and Riegel (2012). This shorter sampling interval could have generated a semivariogram with more data points, tighter spacing between data points, and a better model fit. We did find the generalized semivariogram models, A, B, and C were useful in evaluating the MIM data.

The 80 DMAs selected represent a cross section of samples including multiple stream sizes and types, streams in meadows, canyons, high elevations, deserts, plains, and prairies. Also, these samples represented a diversity of hydrologic sources, including snowmelt-fed perennial streams, monsoon-driven streams, intermittent streams, and groundwater-fed streams, including small spring brooks and vegetated drainageways. And these samples represented a range of conditions resulting from intensive and chronic overgrazing to short-duration rotational grazing and long-term rest in livestock and wildlife exclosures.

We used a correlation coefficient of 0.2 as the estimator of distance at which spatial autocorrelation is negligible or absent (red line in figure A1, for example). Table A3 contains a listing of all DMAs with their corresponding sample point spacing estimates. Not all indicators were measured at all DMAs, so some DMAs do not have a sample point-spacing estimate for each indicator. Using the data in figure A2, we calculated median sample point spacing and calculated a frequency distribution for each metric indicator. We chose the median, rather than the mean, as the spacing-estimate data are heavily skewed. Also, for each indicator, we examined

semivariograms and calculated the frequency of occurrence in each of the semivariance models of Weixelman and Riegel (2012).

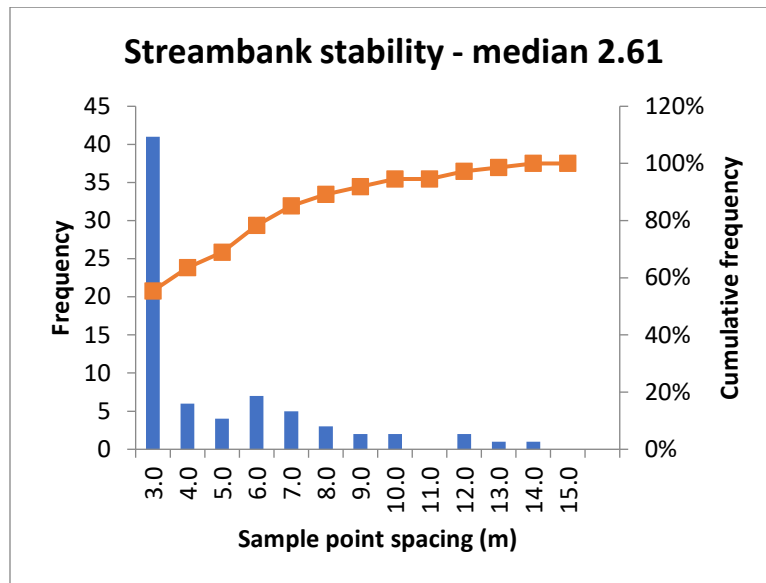
Summary

For each indicator, we analyzed the correlogram results at a correlation coefficient of 0.2 and fit the data to a cumulative frequency graph to evaluate the frequency at which various sample point spacing intervals were required to avoid spatial autocorrelation ($r < 0.2$). The median value represents the mid-point distance at which spatial autocorrelation was negligible. The point at which the frequency graph flattens (at the top of the curve) suggests no spatial correlation beyond that sample interval.

We assessed the percentage of DMAs that exhibited spatial independence at a sample point spacing of 3.75 m. This distance would provide for collecting 80 samples in a 150-meter-long DMA.

The following summarizes the results for the MIM indicators evaluated for spatial autocorrelation. The graph displays the cumulative frequency distribution of sample point spacing with no spatial autocorrelation. The inflection shows the distance at which almost all DMAs showed no spatial autocorrelation.

A. Streambank stability:



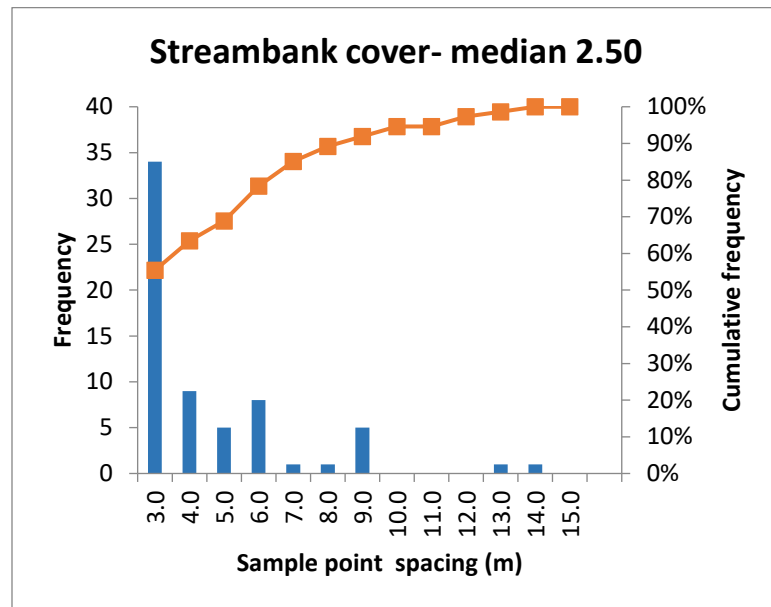
Semivariance model type: With respect to the model categories, 44% were type A, 29% the B type, and 26% C.

The highest frequency for sample spacing was less than 3 m, however almost one-third had a semivariance model type C, indicating longer sample spacing. A sample point spacing of 3.75 m would accommodate negligible correlation in about 61% of the DMAs.

For end-of-season samples, the analysis shows much less spatial autocorrelation: 62% A, 23% B, and 8% C.

90 percent of samples had a sample spacing of less than 3.75 m.

B. Streambank cover



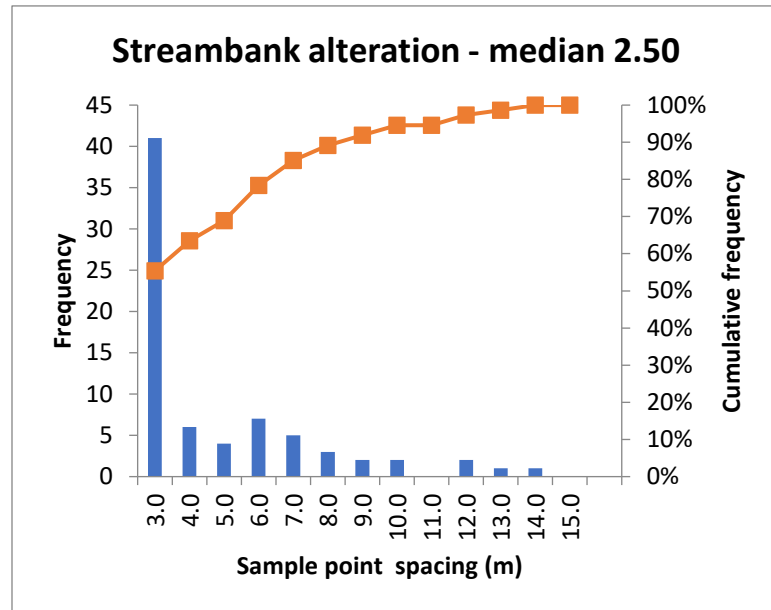
Semivariance model type: With respect to the model categories, 39% were type A, 46% the B type, and 14% C.

The highest frequency for sample spacing was less than 3 m, however almost half had a semivariance model type B, indicating longer sample point spacing. A sample spacing of 3.75 m would accommodate negligible correlation in about 63% of the DMAs.

For end-of-season samples, the analysis shows much less spatial autocorrelation: 62% A, 15% B, and 15% C.

72 percent of samples had a sample spacing of less than 3.75 m.

C. Streambank alteration:



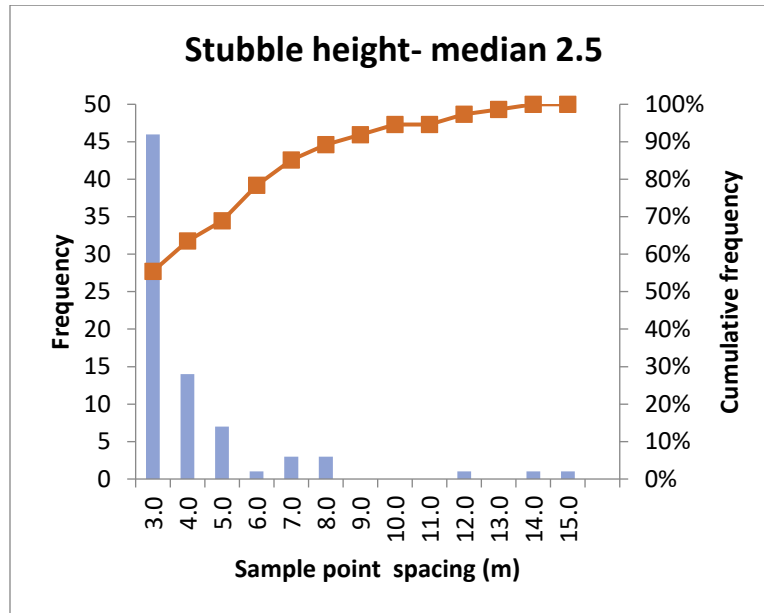
Semivariance model type: With respect to the model categories, 58% were type A, 14% the B type, and 26% C.

The highest frequency for sample spacing had a distance of less than 3 m. A sample point spacing of 3.75 m would accommodate negligible correlation in about 80% of the DMAs.

For end-of-season samples, the analysis shows much less spatial autocorrelation: 73% A, 18% B, and 5% C.

90 percent of samples had an end-of-season sample point spacing of less than 3.75 m.

D. Stubble height:



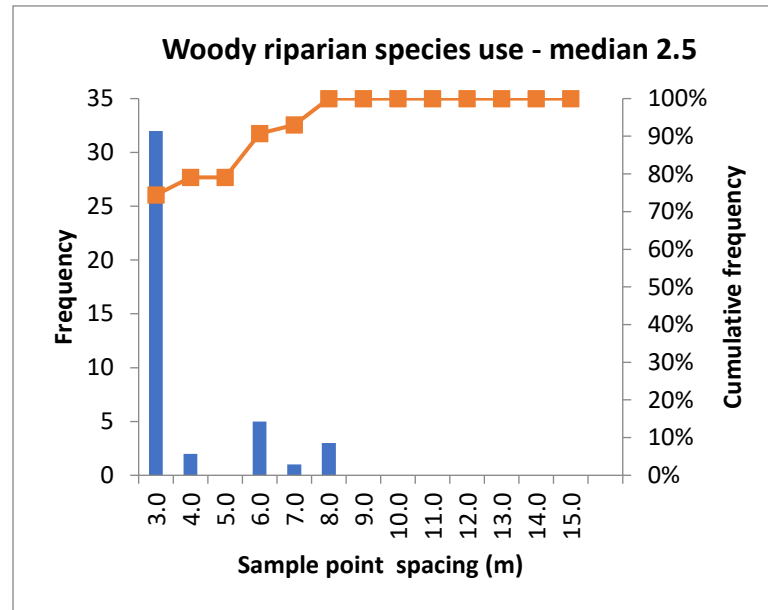
Semivariance model type: With respect to the model categories, 63% were type A, 23% B type, and 11% C.

The highest frequency for sample spacing had a distance of less than 3 m; however, some had a semivariance model type C, indicating a longer sample interval. The few sites that fit a type C model tended to be either lightly grazed (for example, it might not have been the appropriate time to measure grazing-use criteria, or use was focused on a few accessible quadrats that created a clustered use pattern), or alternatively heavily grazed (in which case a few inaccessible quadrats could greatly influence results for the entire reach). A sample spacing of 3.75 m would accommodate negligible correlation in about 73% of the DMAs.

For end-of-season samples, the analysis shows less spatial autocorrelation: 67% A, 14% B, and 10% C.

100 percent of end-of-season samples had an end-of-season sample point spacing of less than 3.75 m.

E. Woody riparian species use:



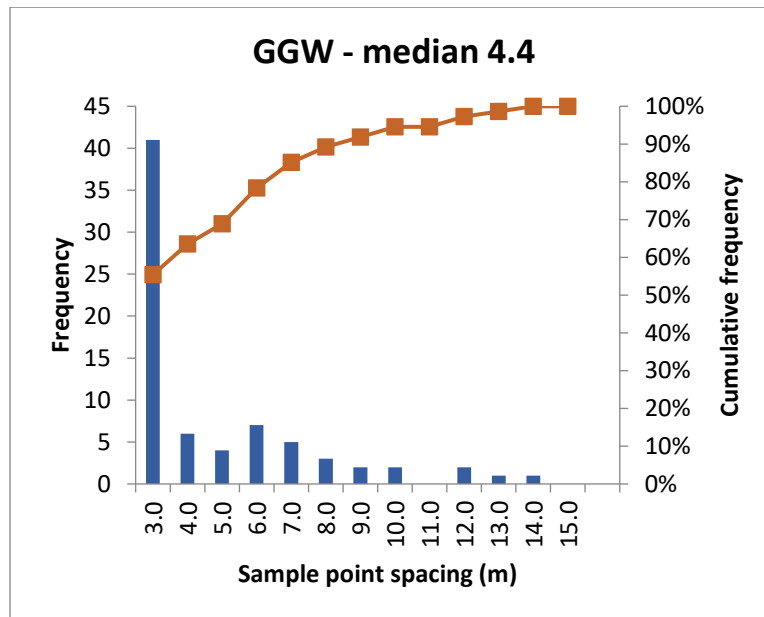
Semivariance model type: With respect to the model categories, 78% were type A, 7% the B type, and 9% C.

The highest frequency for sample spacing had a distance of less than 3 m, and the majority of DMAs had a semivariance model type A, suggesting that most of the time there is no spatial autocorrelation. However, the semivariogram models for this indicator were not robust and tended to produce low r-square values. A sample spacing of 3.75 m would accommodate negligible correlation in about 78% of the DMAs.

For end-of-season samples, the analysis shows less spatial autocorrelation: 71% A, 18% B, and 6% C.

88 percent of samples had an end-of-season sample spacing of less than 3.75 m.

F. Greenline-to-Greenline width:

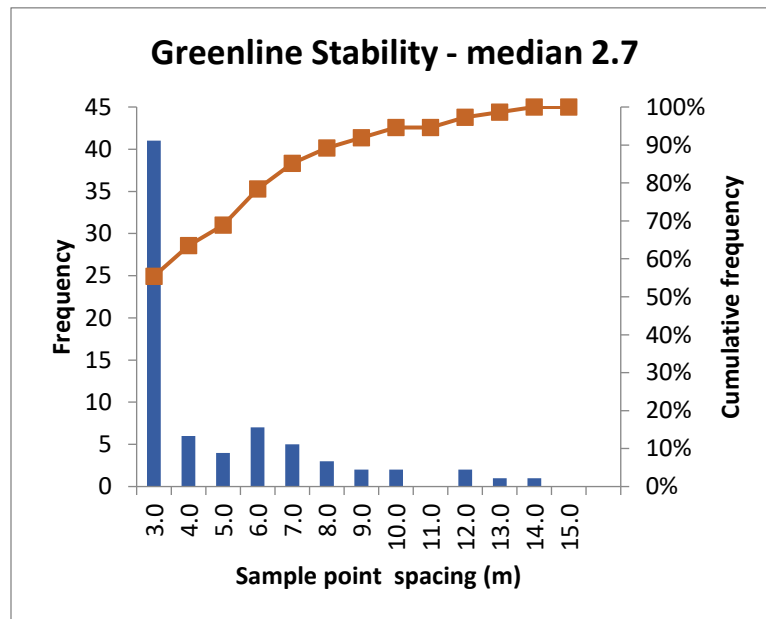
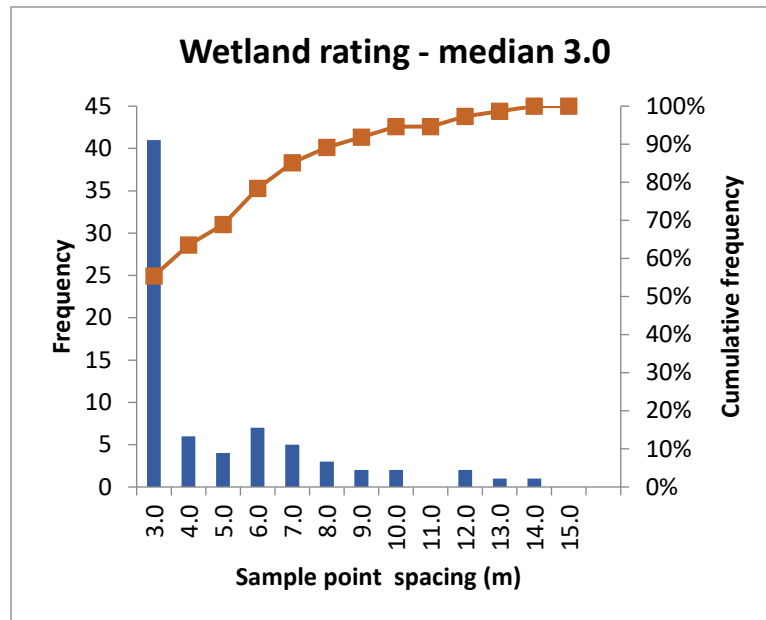


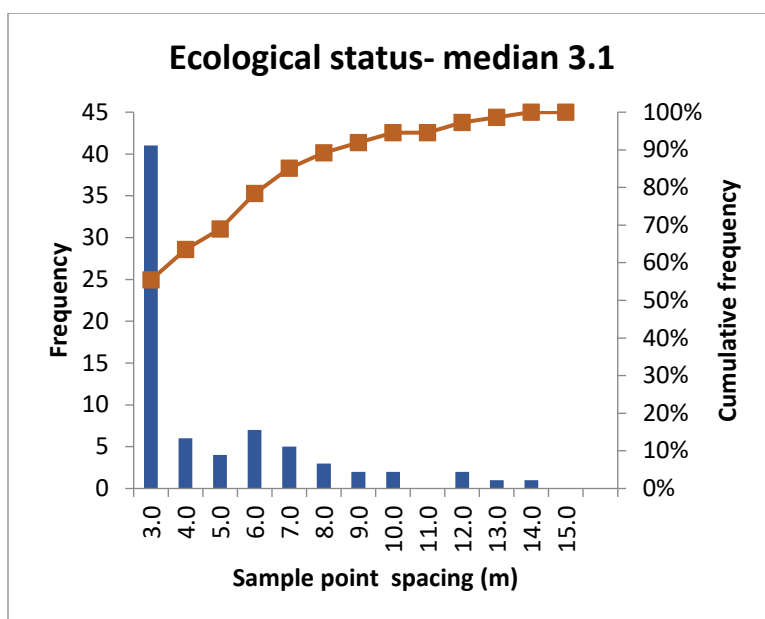
Semivariance model type: With respect to the model categories, 24% were type A, 60% the B type, and 16% C.

The highest frequency for sample spacing had a distance less than 3 m, however GGW had the largest median distance of 4.4 m and only 24% of test sample points fit the A type model (no spatial autocorrelation). A spacing of 3.75 m would accommodate negligible correlation in 41% of the DMAs. Based on Myers and Swanson (1997), this spacing of 4.4 m may be too small for streams whose width is greater than 1.5 m ($4.4/3$). They found 3 channel widths and only 10 cross-channel transects would be more appropriate for estimating stream and/or channel width. Our own data suggests that 3 channel widths would be adequate for 80% of the DMAs. Thus, in 3.75 m transect spacing design, perhaps every other sample point would be the appropriate spacing for GGW.

G. Wetland rating, greenline stability rating, and ecological status:

Wetland rating, greenline stability rating, and ecological status were quite similar in response to spatial autocorrelation.





As with other indicators, the highest frequency for sample spacing was 3 m, and the median values ranged from 2.7 m (for vegetation status) to 3.1 m (for ecological status).

Semivariance models were as follows.

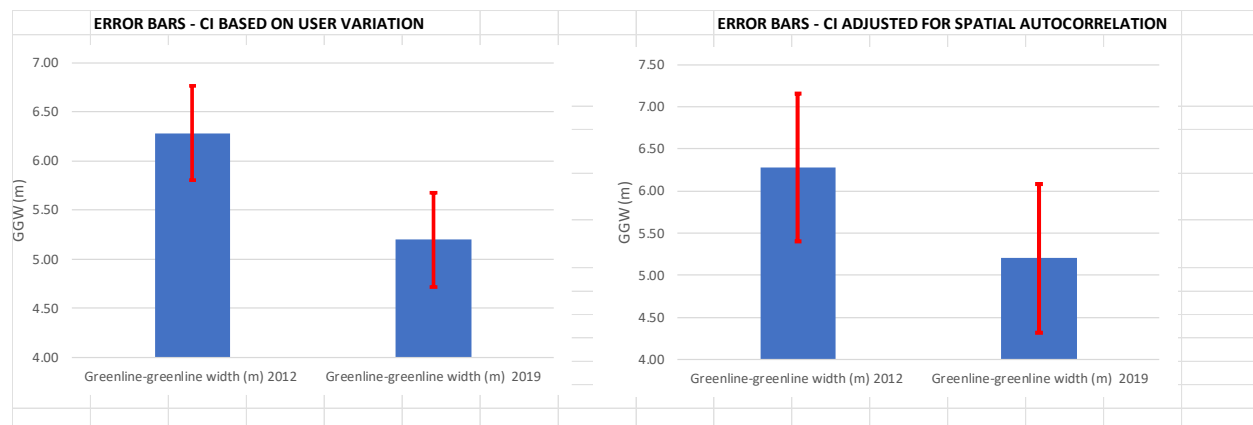
Type	Wetland Rating	Vegetation Stability	Ecological Status
A	38%	45%	42%
B	48%	39%	39%
C	13%	14%	16%

Most occurred in the A and B types, with a few in the C type, suggesting spatial autocorrelation at some DMAs. A sample spacing of 3.75 m would accommodate negligible correlation in 64% (Wet Rating), 64% (Veg Stability), and 65% (Ecological Status) of the DMAs.

Recommendations

For any sample spacing using the MIM method, the likelihood of spatial autocorrelation, for one or more of the indicators, is a real possibility. The correlation table, like the one in Table A1, is now included in the Data Analysis Module. This analysis is evaluating sample point correlations for one side of the stream (sample points 1 to 40, unless the user specifies a different sample point number for the cross-over in the “Header” tab of either the Data Analysis or Data Entry module), to avoid mixing sample points that are not on the same transect. Using that table, the user has the ability to

assess the presence of spatial autocorrelation in their data. If spatial autocorrelation is present, the user has two options. The first is to assess the distance at which the correlation coefficient is less than 0.2, then apply that spacing to the data. If it occurs at every other sample point, then the data collected from every other sample point would likely be far enough apart such that each sample point is considered independent and statistical tests could be applied to the sample point data. In this case the local variation in the quadrat or point samples could be used to calculate the confidence interval for that indicator. However, the loss of every other sample point could be a significant loss in the precision of the estimate due to the reduced sample size. Another option, if the sample size reduction is excessive, is to treat all samples as part of one plot. In this case the DMA itself is the sample. Means, medians, or proportions would simply represent the site condition without confidence intervals. The confidence intervals developed for user variation on the indicator would be applied to significance testing. The following is an example that uses these two options in the case of GGW at a DMA where spatial autocorrelation in this indicator required increasing the quadrat spacing so that every third plot separated by 8 m was chosen. This reduced the sample size from 80 to 27 and substantially increased the confidence interval. Analyzing trend from the 2012 sample as compared with the 2019 sample produced the following result:



Note that the confidence interval represented by the error bars on user variation (left graphic) do not overlap, while those from the remaining 27 sample points (every third sample point) in the DMA data (right graphic) do overlap. The user might conclude that there was a significant decrease in GGW from 2012 to 2019 based on the precision of the user variation but could not do the same for the DMA data where a limited number of samples were available to use. Given this information the user has two options going forward: 1. increase the sample size by adding length to the DMA and collecting more GGW samples, or 2. analyze the data for statistical significance using a more powerful test such as the *t*-test. Concluding that there was a significant difference based on the user variation alone would potentially result in a type I error, basically rejecting the null hypothesis

that the two samples (2012 and 2019) are not significantly different when in reality they are not different, and it should have been accepted.

A standard sample point spacing of 3.75 m would be a more appropriate design for the MIM protocol, rather than the existing 2.75 m. This would result in fewer samples per 110-meter DMAs, (58 instead of the current 80). If spatial autocorrelation is indicated, future re-sampling of DMAs that used the formerly standard 110-meter stream reach, would preferably use the same exact DMA, but would have fewer samples and therefore a wider confidence interval, based on an examination of the data from many DMAs. The option of adding additional samples to narrow the confidence interval would not be recommended to avoid having samples less than 3.75 m apart. However, it might be desirable to extend the length of the DMA to 150 m allowing a larger sample size at 3.75 m spacing. Doing so would require calibrating the new DMA to the old for future trend analyses (see Appendix E for discussion on calibration). The percentage of test sites that exhibited negligible autocorrelation at 3.75 m is summarized in Table 4; whereas the median and 80th percentiles sample point spacing are summarized in Table 5 (all indicators) and Table 6 (only the short-term (or end-of-season) indicators).

Original testing of the MIM protocol indicated the need to have at least 80 samples per DMA to adequately account for site variability. This meant that for a standard 110-meter DMA, the sample point spacing is equal to or less than 2.75 m. To reduce the potential for spatial autocorrelation, the updated protocol standardized the DMA length to 150 m with a 3.75-m spacing to attain the 80-sample target. For wide channels (i.e., those with GGW > 7.5 m), the length of the DMA is 20 times the average GGW.

Since the short-term indicators, streambank stability and cover, stubble height, streambank alteration, and woody riparian species use have relatively small median sample point spacing (from 2.5 to 3.0 m), a 3.75-m spacing is more than adequate to address these management indicators. These indicators appear to have greater spatial autocorrelation when (1) the short-term indicators were measured at the wrong time, (i.e., when there were few or no quadrats with annual use), so the level of use was too low to properly evaluate a particular use indicator, (2) patches or clusters of poor condition exist at the site, generally at livestock access points along the greenline; or (3) the monitoring site (i.e., DMA) was severely degraded. Interestingly, sites with severely degraded conditions that improved over time exhibited reduced spatial autocorrelation as riparian conditions improve. A good example is the Elk Creek DMA where these indicators had higher spatial autocorrelation in the earlier period when the stream was in poor condition. After recovery some years later, spatial autocorrelation was reduced significantly supporting a reduced sample spacing.

Greenline-to-greenline width (GGW) had the highest spatial autocorrelation distance (median 4.4 m). In this case, a spacing of more than 4 m would greatly reduce the sample size used to calculate mean GGW. Legendre (1993) discussed the problem of spatial autocorrelation and proposed several solutions.

“First, one can attempt to remove the spatial dependency among observations so that the usual statistical tests can be used, either by removing samples until spatial independence has been attained (a solution that is *not* recommended because it entails a net loss of expensive information) or by filtering out the spatial structure using trend surface analysis....”

“The alternative is to modify the statistical method in order to take spatial autocorrelation into account; this approach is to be preferred when such a method is available, especially in cases where spatial structuring is seen not as a nuisance but as a part of the ecological process under study (previous section). Cliff and Ord (1973) have proposed a method for correcting the standard error of the parameter estimates of the simple linear regression in the presence of autocorrelation. This method is extended to linear correlations, multiple regressions, and *t*-tests by Cliff and Ord (1981: Chapter 7) and to the one-way analysis of variance by Griffith (1978, 1987).”

Spatial autocorrelation in GGW is likely a result of the fluvial geomorphic patterns or spatial structure that is inherent in stream systems. Those spatial patterns appear to often produce autocorrelation and therefore spatial dependency for GGW samples that occur on the same channel type. Samples that are immediately adjacent on a meander bend, for example, would have similar widths. But so too would a sample taken on one meander bend and another downstream at great distance from the first – both separated by long distance and occurring on the same channel shape – meander bend. Also, geomorphic patterns are exhibited by pool-riffle and step-pool sequences along the longitudinal axis of a stream. GGW measurements will pick up on this spatial structure, which as Legendre (1993) points out is not a nuisance, but a part of the ecological processes under study. So, for GGW, perhaps this is one example of reducing sample size by increasing the sample interval might not always remove spatial autocorrelation. Thus, applying the 3.75 meter spacing to GGW is recommended and then analyzing trend or condition by using the confidence intervals developed for user variation rather than standard error of GGW at the DMA, or alternatively by applying some method for correcting the standard error.

Conclusions

Following MIM data collection, users will want to upload their field data into the current version of the Data Analysis Module. This module now includes a calculation of adjacent sample point

correlation coefficients for more indicators than is available in the Data Entry Module for existing DMAs, historical data should be uploaded into the current Data Analysis Module to assess the presence of spatial autocorrelation. If it is absent in the data, the sample point confidence interval(s) will be provided as contained in the Data Summary tab. Future samples at the existing DMA could be taken on the original 110-meter reach (or same length previously applied) and can be collected at the original spacing if spatial autocorrelation was absent. However, if it was found to be present, the 110-meter reach should be re-sampled at the 3.75-meter sample-point spacing.

For greenline-to-greenline width, samples should be collected only from one side of stream to avoid the potential for sampling in exactly or close to a previous measurement taken from the other side of the stream. This will reduce the sample size and therefore increase the width of the confidence interval. However, it is important to note that it may be better to be conservative and realistic about the uncertainty in the data.

Table A1. Correlation matrix for Stubble Height at Hawley Creek, Idaho showing moderate positive correlation among adjacent plots (2.5 m spacing), declining to no or negligible correlation at a distance of every other plot (5 m). These data were collected prior to grazing and likely have a higher degree of spatial autocorrelation than is typical when observed after grazing.

INDICATOR:	Stubble Height	Adjacent sample points	Every other sample point	Every third sample point	Every fourth sample point	Every fifth sample point	Every Sixth sample point
	CORRELATION	0.5185	0.0402	-0.1441	-0.1135	-0.0581	-0.0154
	N =	95	94	93	92	91	91
Sample		0	0	0	0	0	
1	6						
6	8	6					
8	7	8	6				
7	6	7	8	6			
6	7	6	7	8	6		
7	10	7	6	7	8	6	
10	8	10	7	6	7	8	6
8	10	8	10	7	6	7	8
10	12	10	8	10	7	6	7
12	6	12	10	8	10	7	6
6	7	6	12	10	8	10	7
7	9	7	6	12	10	8	10
9	11	9	7	6	12	10	8
11	6	11	9	7	6	12	10
6	6	6	11	9	7	6	12
6	9	6	6	11	9	7	6
9	20	9	6	6	11	9	7
20	14	20	9	6	6	11	9
14	12	14	20	9	6	6	11
12	14	12	14	20	9	6	6
14	12	14	12	14	20	9	6
12	12	12	14	12	14	20	9
12	15	12	12	14	12	14	20
15	11	15	12	12	14	12	14
11	4	11	15	12	12	14	12
4	4	4	11	15	12	12	14
4	6	4	4	11	15	12	12
6	14	6	4	4	11	15	12
14	18	14	6	4	4	11	15
18	12	18	14	6	4	4	11
12	8	12	18	14	6	4	4
8	8	8	12	18	14	6	4
8	5	8	8	12	18	14	6
5	8	5	8	8	12	18	14
8	5	8	5	8	8	12	18

INDICATOR:	Stubble Height	Adjacent sample points	Every other sample point	Every third sample point	Every fourth sample point	Every fifth sample point	Every Sixth sample point
5	5	5	8	5	8	8	12
5	4	5	5	8	5	8	8
4	10	4	5	5	8	5	8
10	12	10	4	5	5	8	5
12	12	12	10	4	5	5	8
12	12	12	12	10	4	5	5
12	14	12	12	12	10	4	5
14	22	14	12	12	12	10	4
22	20	22	14	12	12	12	10
20	16	20	22	14	12	12	12
16	8	16	20	22	14	12	12
8	18	8	16	20	22	14	12
18	21	18	8	16	20	22	14
21	16	21	18	8	16	20	22
16	8	16	21	18	8	16	20
8	8	8	16	21	18	8	16
8	12	8	8	16	21	18	8
12	6	12	8	8	16	21	18
6	28	6	12	8	8	16	21
28	10	28	6	12	8	8	16
10	14	10	28	6	12	8	8
14	20	14	10	28	6	12	8
20	10	20	14	10	28	6	12
10	10	10	20	14	10	28	6
10	11	10	10	20	14	10	28
11	13	11	10	10	20	14	10
13	10	13	11	10	10	20	14
10	8	10	13	11	10	10	20
8	10	8	10	13	11	10	10
10	10	10	8	10	13	11	10
10	14	10	10	8	10	13	11
14	14	14	10	10	8	10	13
14	8	14	14	10	10	8	10
8	10	8	14	14	10	10	8
10	12	10	8	14	14	10	10
12	18	12	10	8	14	14	10
18	7	18	12	10	8	14	14
7	6	7	18	12	10	8	14
6	20	6	7	18	12	10	8
20	20	20	6	7	18	12	10

INDICATOR:	Stubble Height	Adjacent sample points	Every other sample point	Every third sample point	Every fourth sample point	Every fifth sample point	Every Sixth sample point
20	12	20	20	6	7	18	12
12	27	12	20	20	6	7	18
27	14	27	12	20	20	6	7
14	28	14	27	12	20	20	6
28	28	28	14	27	12	20	20
28	26	28	28	14	27	12	20
26	30	26	28	28	14	27	12
30	19	30	26	28	28	14	27
19	28	19	30	26	28	28	14
28	8	28	19	30	26	28	28
8	17	8	28	19	30	26	28
17	11	17	8	28	19	30	26
11	11	11	17	8	28	19	30
11	10	11	11	17	8	28	19
10	14	10	11	11	17	8	28
14	12	14	10	11	11	17	8
12	13	12	14	10	11	11	17
13	11	13	12	14	10	11	11
11	10	11	13	12	14	10	11
10	12	10	11	13	12	14	10
12	6	12	10	11	13	12	14

Table A2. Sample spacing distance (meters) at which the correlation coefficient is less than 0.2 at end-of-season DMA samples in one allotment in Nevada for 2015 and 2016.

DMA	Parameter spacing- Correlogram at r = .2				Wdy Use
	Stab	Cover	Alteration	Stub Ht	
Trout Creek End of season data	2.50	2.50	2.50	2.50	
The Park end of season data	2.50	2.50	2.50	2.50	
Slaven end of season data	2.50	2.50	2.62	2.50	
Rock Creek end of season data	2.50	2.50	2.50		2.50
Ratfink end of season data	3.79	3.79		2.50	2.50
N Fork Mill Creek end of season data	2.50	2.50	2.50	2.70	
Indian Creek end of season	2.50	3.92	2.50	2.50	2.50
Harry Canyon end of season data	2.50	2.50	2.50		2.50
Fire Creek end of season	2.50	4.05	2.50	2.50	2.50
Ferris Creek end of season data	7.48	2.50	2.74	2.68	2.50
Crippen end of season data	2.50	2.50	6.33	2.50	5.03
Corral Crk end of season data	2.50	2.50	2.50	2.50	2.77
Corral Creek Nevada 2015			3.43	2.50	2.50
Crippen Creek Nevada 2015			2.91	2.50	2.50
Ferris Creek Nevada 2015			14.11	2.50	2.50
Fire Creek Nevada 2015			2.50	2.95	
Harry Canyon Nevada 2015			2.50	2.50	5.92
Indian Creek Nevada 2015			2.50	2.50	2.50
N Fork Mill Creek Nevada 2015			2.50	2.50	
Slaven Creek Nevada 2015			2.50	2.50	
The Park Nevada 2015			2.50	2.50	
Trout Creek Tributary Nevada 2015			2.50	3.09	

Table A3. Sample spacing distance (meters) at which the correlation coefficient is less than 0.2.

DMA	Stab	Cover	Alteration	Stub Ht	Wdy Use	GGW	Wet Rat
Alkali Creek, Wyoming 2012	2.50	2.50	2.50	2.50	2.50	2.50	
Alta creek, lower Nevada 2016	3.08		4.14	2.50	2.50	5.63	2.50
Alta Creek, Nevada			6.35	2.50		7.09	9.35
Argenta Corral Creek, Nevada 2016	5.15	2.50	2.50	2.50		5.69	2.50
Argenta Crippen Creek, Nevada 2016	2.50	8.28	2.50	4.20		2.50	2.84
Argenta Ferris Creek, Nevada	4.69	4.97	2.50	2.50		4.46	5.65
Argenta Fire Creek, Nevada	13.02	2.50	2.50	2.50	2.50		2.50
Argenta Indian Creek, Nevada	7.39	3.05	2.50	2.50	2.50	2.50	2.50
Argenta N Fk Mill Creek, Nevada	9.13	2.50	3.58	4.30	2.50	4.07	4.94
Argenta Slaven, Nevada 2016 baseline	5.31	5.53					3.48

DMA	Stab	Cover	Alteration	Stub Ht	Wdy Use	GGW	Wet Rat
Argenta the Park, Nevada 2016 baseline	2.50		2.50	2.50		4.44	5.78
Argenta Trout Creek, Nevada, 2016				2.50		2.50	2.50
Bear Creek Lower, Oregon 2022	5.74	5.21	2.50	2.50	2.50	6.36	4.66
Bear Creek Meadow	3.36	4.50	2.50	2.50		11.58	4.60
Bear Creek Upper, Oregon 2022	2.58	2.50	2.50	3.59	3.54	2.50	2.50
Bear Creek, Dude, Oregon	5.93	5.93	2.50	4.19		2.50	8.02
Bear Creek, Sheep Gulch Oregon	2.86	2.80	2.50	2.50		2.50	3.98
Big Creek, Nevada	2.50	2.50	2.94	3.03	5.88	4.34	3.91
Big Elk Creek DMA 3 Idaho 2012	8.67	8.67		14.85	5.13	4.44	3.73
Bison Creek, Montana 2012	7.75	13.73	2.63	2.86	7.32	10.42	2.50
Bluebucket Creek, Oregon 2010	2.50	7.65	2.50	7.95		4.35	2.50
Burr Creek, Utah 2019	4.59	4.02	11.46	13.48	2.50	4.06	2.50
Campaign Creek, Utah	7.00	6.02	2.50			2.50	3.24
Castle Creek Exclosure, Wyoming 2013	2.50	2.50		2.50	2.50	12.12	2.50
Castle creek Watergap Wyoming 2013	6.94	2.50	3.35	3.24		2.91	2.50
Corral Cr The Park 2016 baseline	2.50		2.50	2.50		4.44	5.78
Cottonwood Creek, Oregon 2010	2.75	2.50	3.50	3.98		6.91	2.50
Cougar Creek, Oregon 2010			5.54	2.50		2.50	
Crane Creek, Oregon 2010	2.50	2.50		7.47	2.50	13.63	3.58
Crooked Creek Upper, Oregon 2022	2.50			2.50	2.50	2.50	2.50
Crooked Creek, Lower, Oregon 2022	3.05			2.50	2.50	2.50	2.50
Deer Run Exclosure 2014	9.92	12.10		4.06		2.50	3.09
Dixie Cr Meadow, Nevada 2006			2.50	2.50	7.92	3.41	5.34
Dixie Cr Upper Fence			2.50	2.54		3.19	2.50
E.F. Little Morgan Idaho 2009			4.03	2.50		6.63	2.50
East Pearl, Nevada 2021	5.85	2.50	2.50	2.55		2.50	2.50
EF L MORGAN, Idaho 2014	11.32	8.85	2.50	2.50		15.00	2.50
Elk Creek DMA 3 Idaho 2005			2.50	6.11	2.50	4.67	2.50
Elk Creek DMA 3 Idaho 2019	8.02	8.02		7.77	2.50	3.85	2.50
Elk Creek DMA1, Idaho 2005	4.97	2.50	8.33	3.37	6.00	5.39	6.44
Elk Creek DMA1, Idaho 2008	12.26	8.77	2.50	2.50		8.27	8.34
Elk Creek DMA1, Idaho 2012	3.58	3.58		4.48	3.92		4.87
Elk Creek DMA1, Idaho 2019	2.50	2.50		4.99	2.50	4.89	3.43
French Creek DMA, Wyoming 2013			2.50	2.50		3.37	2.50
French Creek Watergap, Wyoming 2013	2.50		2.50	0.00	2.50	2.50	2.50
Havens Pinto Creek, California 2015	2.50	2.50	2.50	2.50	2.50		6.49

DMA	Stab	Cover	Alteration	Stub Ht	Wdy Use	GGW	Wet Rat
Hawley Eighteenmile Creek, Idaho 2018	2.50	2.50	4.63	4.16		2.50	2.50
Indian Jack Cr Lower Nevada 2011	3.05		2.50	3.50	2.50	2.72	
Indian Jack Meadow, Nevada 2011	3.05	2.50	2.50	3.81	2.50	2.50	2.50
Little Lost Creek, Idaho 2019	5.79	6.00	2.50	6.96		5.72	2.50
Little Malheur R, Oregon 2010	2.50	2.50	2.50	2.50	2.50	10.77	4.03
Lower Big Creek, Nevada 2016	2.50	2.50	2.50	2.50	2.50	3.74	2.50
Lower BLACK HORSE BUTTE, South Dakota 2014	11.67	2.50		6.91		11.83	9.22
Lower Dog Tooth Creek, South Dakota 2014	2.50		3.38	3.21		7.26	
Lower WF Blacktail Deer Creek Montana 2020	3.83	3.80		2.88		6.44	2.92
North Fork Malheur River, Oregon 2010	2.50	2.50	4.12	2.50	11.30	12.96	7.26
Pacific Creek lower, Wyoming 2010	2.50	2.50	2.50	3.87		12.27	2.50
Pacific Creek lower, Wyoming 2022	2.65	3.49	4.03	5.42		5.23	3.19
PACIFIC CREEK, Wyoming 2010	2.50	2.50	2.50	3.87		12.27	2.50
Pacific Creek, Wyoming, 2022	2.65	3.49	4.03	5.42		5.23	3.19
Pearl Creek Nevada 2021	3.42	3.07	2.50	3.84	2.50	4.58	6.76
Rattlesnake Creek Arizona 2014	9.22	2.50	2.50	5.79		6.39	2.50
Reservoir Creek Idaho 2014	4.94	3.73	2.62	4.74		5.35	14.72
Rio Bonito, New Mexico 2021	2.50	2.50	2.50	2.50	5.32	2.50	4.24
SF Beaver, Washington	2.50	2.50	2.50	2.50	2.50	12.62	2.50
SF Flatwillow Montana 2013	2.88	5.11	2.50	3.35	2.50	6.13	2.50
Shoshone Creek, Idaho 2005			4.11	2.50			2.50
Slaven Creek baseline Nevada 2016	5.31	5.53					3.48
Summit Creek Oregon 2010	2.50	2.50	2.50	2.50	5.60	2.73	6.12
Tangle Creek Arizona 2008	2.50			2.50		5.75	2.50
Taylor Creek Lower Montana 2020	6.26	5.60	4.04	3.62			7.66
Upper Big Creek, Nevada	2.50	3.05	2.50	2.50	2.50	6.87	3.05
Upper Black Horse Butte, South Dakota 2014	2.50	2.50	9.86	11.23	2.50	9.23	3.77
Upper Dogtooth, South Dakota, 2014	4.04	3.07	2.50	3.91		10.26	5.47
Upper Rio Bonito, New Mexico 2021	2.50		10.25	3.37	2.50	5.97	10.25
Upper WF Blacktail Deer Creek Montana 2020			2.50	2.50	7.07	2.50	14.04
Water gap Castle Creek, Wyoming 2013	6.94	2.50	3.35	3.24		2.91	2.50
Water Gap French Creek, Wyoming 2014	2.50		2.50		2.50	2.50	2.50

Table A4. The percentage of test DMAs that exhibit negligible autocorrelation with a sample spacing of 3.75 m.

	Stability	Cover	Alteration	Stubble Height	Woody Use	GGW	Wetland Rating	Vegetation Stability	Ecological Status
All test DMAs	61%	63%	80%	73%	78%	41%	64%	64%	65%
End-of-season DMAs	90%	72%	90%	100%	88%	---	---	---	---

Table A5. Calculated sample spacing (in meters) at which test sites exhibit negligible autocorrelation (median and 80th percentiles).

	Streambank Stability	Streambank Cover	Streambank Alteration	Stubble Height	Woody Use	GGW	Wetland Rating	Vegetation Stability	Ecological Status
Median:	3.05	3.05	2.50	2.88	2.50	4.63	3.05	2.66	3.20
Percentile - 80th	6.53	5.60	4.03	4.33	5.24	7.23	5.78	5.43	5.27

Table A6. Calculated sample spacing (in meters) at which test sites exhibit negligible autocorrelation (median, 67th, and 80th percentiles) for end-of-season indicators collected at the proper time.

		Stab	Cover	Alteration	Stub Ht	Wdy Use
	Median:	2.50	2.50	2.50	2.50	2.50
	Percentile - 67th	2.50	3.40	2.50	2.50	2.50
	Percentile - 80th	2.50	3.89	2.74	2.54	2.61

REFERENCES

- Burton, Timothy A. Steven J. Smith, and Ervin R. Cowley. 2011 Multiple Indicator Monitoring (MIM) of Stream Channels and Streamside Vegetation. USDI Bureau of Land Management. Technical Reference 1737-23.
- Elzinga, C., R. D.W. Salzer, and J.W. Willoughby. 1998. Measuring & Monitoring Plant Populations. USDI Bureau of Land Management. Technical Reference 1730-1: 477 pages.
- Legendre, Pierre. 1993. Spatial Autocorrelation: Trouble or Paradigm? *Ecology* 74(6):1659-1673.
- Lichvar R.W. N.C. Melvin M.L. Butterwick and W.N. Kirchner. 2012. National Wetland Plant List Indicator Rating Definitions. ERDC/CRREL TR-12-1. U.S. Army Engineer Research and Development Center Cold Regions Research and Engineering Laboratory. Hanover, NH.
- Myers, Thomas J. and Sherman Swanson. 1997. Precision of Channel Width and Pool Area Measurements. *Journal of the American Water Resources Association* 33(3):647-659.
- Moran, P. A. P. 1950. "Notes on Continuous Stochastic Phenomena". *Biometrika* **37** (1): 17–23. doi:10.2307/2332142. JSTOR 2332142. PMID 15420245.
- Weixelman, D.A. and G.M. Riegel. 2012. Measurement of Spatial Autocorrelation of Vegetation in Mountain Meadows of the Sierra Nevada, California and Western Nevada. *Madroño: a West American Journal of Botany* 59:143-149, Berkeley, California Botanical Society.

APPENDIX B - OBTAINING THE CONFIDENCE INTERVAL FOR NON-NORMALLY DISTRIBUTED DATA USING BOOTSTRAPPING

As described in Chapter III, section B, part 3 of this Guide, the confidence interval is an important statistic in the assessment of monitoring trends and conditions. It is used in MIM for significance testing of the differences between observations. There are three basic assumptions for proper application of the confidence interval (Chapter III, B, 3). First, the assumption that the sample was randomly selected (independence assumption), second that the standard deviation of the sample is known, and third that there are few or no outliers such that the sample mean fits a normal probability distribution. Some MIM indicators tend to fit a non-normal data distribution. More specifically streambank alteration and woody riparian species use tend to have a positive skew in the distribution of samples, with outliers typically occurring in the higher values. Skewness reflects the asymmetry of the frequency distribution about its mean. The skewness value can be positive, zero, negative or undefined. Examples of data distributions with their skewness coefficients are shown in figure B1.

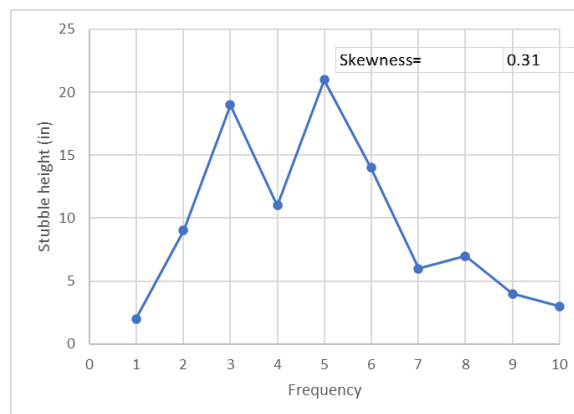
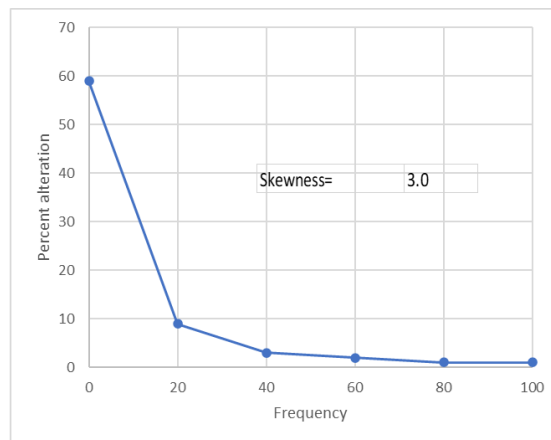


Figure B1. Examples of strong positive skew (streambank alteration - skewness 3.0) and weak positive skew (stubble height - skewness 0.31).

Experience has shown that stubble height often approximates a normal distribution while streambank alteration and woody riparian species use better describe a non-normal distribution with more values in the low range. Woody use samples, for example are more frequently in the “none” to “slight” range with fewer samples in the “moderate” to “severe” range.

METHODS

Skewness was analyzed at 53 DMAs for streambank alteration, and 35 DMAs for woody riparian species use. The average skewness was +2.99 for streambank alteration and +2.65 for woody riparian species use.

For non-normally distributed (strongly skewed) data the rules of statistics do not allow calculation of confidence intervals using the standard normal coefficient. As stated by Elzinga et al. (1998) “...other than the use of resampling techniques, ... there is no nonparametric method available to calculate confidence intervals around means”. Also:

“... resampling methods (also called computer-intensive methods) can be used to calculate confidence intervals and to conduct significance testing. Two of the most commonly used methods are bootstrapping (which involves sampling the original data set with replacement) and randomization (also called permutation) testing (which involves sampling the original data set without replacement).” (Elzinga et al. 1998, Chapter 11, Part I.)

In the present analysis, bootstrapping was the resampling method used to calculate confidence intervals for the two indicators, streambank alteration and woody riparian species use. Data from the DMAs was entered into a bootstrapping routine in EXCEL. Results were compared with confidence intervals calculated using the standard normal coefficient. In this way it is possible to evaluate the general effect of data distribution on the calculation of confidence intervals and to determine what adjustment, if any can be made within the Data Entry Module to estimate the desired sample size more accurately.

The bootstrapping technique used in this analysis has been incorporated into the MIM Statistical Analysis Module. The module uses a standard method for bootstrapping as described in:

<https://www.statology.org/bootstrapping-in-excel/>

This method was validated using an EXCEL add-in: “XrealStats.xlma” by comparing computational results using both EXCEL routines on a set of MIM data. Results of that test indicated that the module calculates the bootstrapped confidence interval with a reasonable level of precision. The MIM module uses 1000 re-samples in the bootstrapping routine. This number of samples was compared to resampling 2000 and 10,000 times to assess differences in resampling results. Computation of the confidence interval from these resampling tests resulted in differences of less than 3% on average. Therefore, the final module was developed using 1000 resampling iterations.

FINDINGS

Streambank alteration:

Results of the bootstrap analysis for streambank alteration are summarized in Table 1. The average 95% confidence interval for the bootstrap analysis was essentially the same as that calculated using the standard normal coefficient. The 95% confidence interval comparison between the two techniques was analyzed using simple linear regression to derive an adjustment based on bootstrapping. These results are shown in figure B2. With a regression coefficient (R^2) of 0.99 the relationship is excellent, and the equation indicates that there is very little adjustment of the 95% confidence interval calculated from the DMA data using the standard normal coefficient.

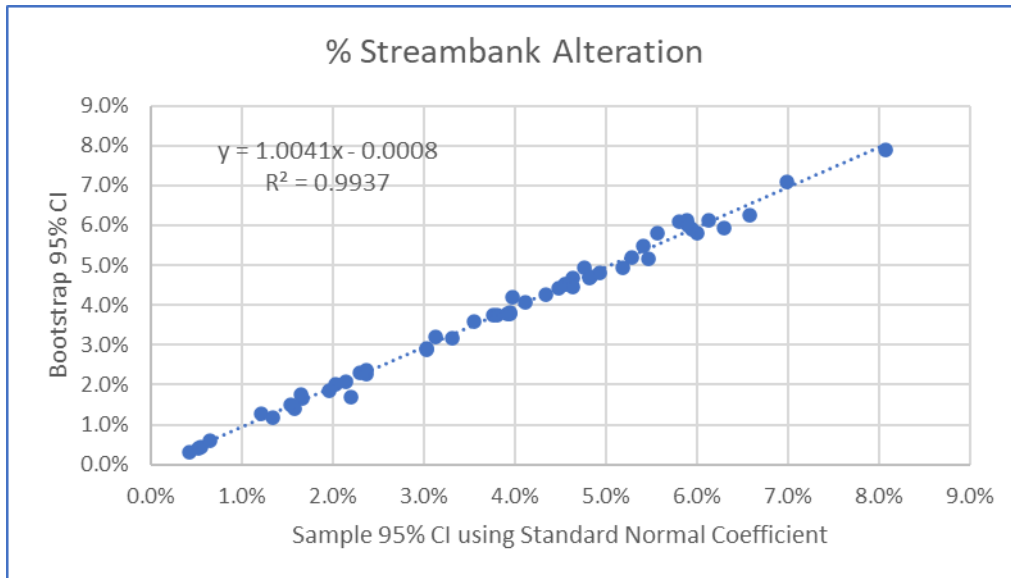


Figure B2. Regression equation for streambank alteration showing an excellent fit with R^2 of 0.99.

Woody riparian species use:

The results of the bootstrap analysis for woody riparian species use are summarized in Table 2. The average 95% confidence interval for the bootstrap analysis was essentially the same as that calculated using the standard normal coefficient. The 95% confidence interval comparison between the two techniques was analyzed using simple linear regression to derive an adjustment based on bootstrapping. These results are shown in figure B3. With an R^2 of 0.98 the relationship is excellent, and the equation indicates that there is very little adjustment of the 95% confidence interval calculated from the DMA data using the standard normal coefficient.

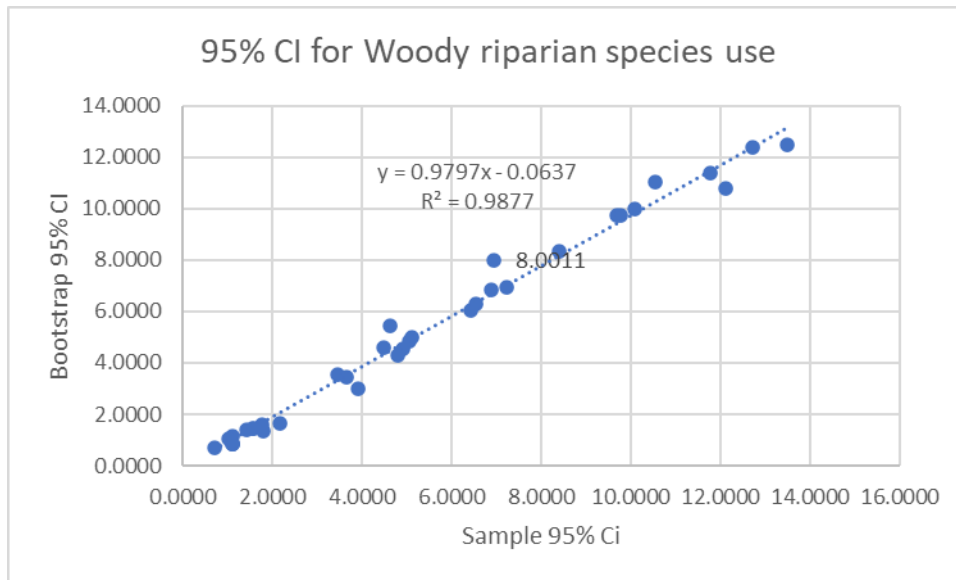


Figure B3. Regression equation for woody riparian species use showing an excellent fit with R^2 of 0.988.

CONCLUSIONS

There was surprisingly little difference between 95% confidence intervals calculated using the standard normal coefficient and by bootstrapping. Perhaps this is because there are so few categories for observation using these two monitoring indicators (streambank alteration hits per plot of 0, 1, 2, 3, 4, and 5), (woody use 0%, 10%, 30%, 50%, 70%, and 90% per plant). Thus, the overall averages per DMA are basically identical to each average in the resamples. Regardless of the reason, we can confidently compute 95% confidence intervals using the developed regression equations as an adjustment to the 95% confidence interval calculated using the standard normal coefficient (EXCEL's "confidence" function), and this adjustment will be minor. Note that the confidence intervals (CIs) in Table B1 represent the single value that is more than (+) or less than (-) the mean.

Table B1. Calculation of mean (hits per plot), standard deviation, skew, sample 95% confidence intervals and bootstrap 95% confidence intervals for streambank alteration data at 53 MIM DMAs. The average for all DMAs is shown at the bottom of the table.

Number	DMA	N	Mean	Stddev	Skew	Sample 95% CI	Bootstrap 95% CI
1	Alta Creek	83	4.01	1.25	-1.53	0.2710	0.2744
2	Berar Creek lower	92	0.70	1.16	2.00	0.2382	0.2473
3	Berar Creek ISheep Gl	83	0.54	0.91	1.58	0.1960	0.1890
4	Bear Dude	81	0.29	0.86	3.82	0.1884	0.1875
5	Black Horse Butte	90	0.80	1.32	1.76	0.2734	0.2584
6	Burr Creek	61	0.33	0.66	2.15	0.1658	0.1585
7	campaign creek	56	0.05	0.23	4.03	0.0606	0.0636
8	Corral Cr EOS	87	0.78	1.08	1.38	0.2279	0.2267
9	Cottonwood Creek	82	0.57	0.91	1.59	0.1976	0.1914
10	Cougar Creek	86	1.76	1.90	0.63	0.4040	0.3943
11	Crane Creek	88	0.14	0.51	4.00	0.1071	0.1036
12	Crippen Creek	79	0.49	0.80	1.67	0.1779	0.1795
13	Dixie Cr Meadow1	71	0.01	0.12	8.37	0.0280	0.0214
14	Dixie Cr Upper Fence	40	0.15	0.37	2.00	0.1147	0.1154
15	Dogtooth Creek	92	0.11	0.41	3.87	0.0836	0.0824
16	Dogtooth Creek upper	76	0.40	0.96	2.99	0.2169	0.2133
17	E.F. Little Morgan	76	0.31	1.07	3.83	0.2410	0.2333
18	Elk Cr Dwnst 2008	79	0.05	0.36	7.78	0.0792	0.0705
19	Eighteenmile	81	0.10	0.38	4.03	0.0825	0.0875
20	Elk Cr Dwnst 2005	90	2.01	1.68	0.41	0.3495	0.3541
21	Ferris Creek EOS	91	1.39	1.43	0.94	0.2949	0.3056
22	Fire Creek	87	1.29	1.42	1.05	0.3004	0.2907
23	Indian Creek EOS	99	1.92	1.59	0.36	0.3149	0.2959
24	Little Lost Creek	87	0.16	0.72	5.23	0.1516	0.1453
25	Bear Meadow	87	0.53	0.94	1.76	0.1990	0.2093
26	Little Malheur R	81	0.71	1.50	2.11	0.3293	0.3125
27	Lower Big Creek	70	0.35	0.87	3.37	0.2056	0.2029
28	Trout Creek EOS	83	1.73	1.34	0.32	0.2907	0.3049
29	Slaven EOS	77	0.07	0.30	4.99	0.0671	0.0592
30	Lower Indian Jack Cr	76	0.61	1.03	1.93	0.2320	0.2333
31	MARKS	86	1.15	1.44	0.97	0.3069	0.3060
32	Rock Creek	96	0.15	0.50	5.34	0.1015	0.1000
33	N Fk Mill Creek	89	2.00	1.42	0.12	0.2972	0.2955
34	North Fork Malheur River	83	0.60	1.05	1.66	0.2278	0.2256
35	PACIFIC CREEK 2010	82	0.73	1.21	1.70	0.2645	0.2593
36	Pacific Creek 2022	95	1.22	1.15	0.82	0.2318	0.2234
37	Pearl Creek	92	0.08	0.48	7.15	0.0980	0.0934
38	Pearl Creek Exclosure	78	0.51	0.88	1.98	0.1972	0.1883
39	Pinto Creek	72	0.06	0.47	8.43	0.1104	0.0845
40	Rattlesnake Creek	93	0.02	0.15	9.59	0.0213	0.0163
41	Reservoir Creek	86	0.68	1.14	1.80	0.2415	0.2353
42	SF Flatwillow	67	0.08	0.32	4.62	0.0770	0.0758
43	Shoshone Creek	76	0.95	1.30	1.38	0.2951	0.3002
44	Summit Creek	78	0.68	1.11	2.06	0.2469	0.2403
45	Taylor lower	109	0.10	0.34	1.85	0.1514	0.1435
46	Taylor upper	110	0.51	0.83	1.81	0.1567	0.1606
47	Upper Big Creek	87	0.67	1.07	1.88	0.2244	0.2209
48	WF Mill Creek	76	0.49	0.85	8.66	0.0261	0.0200
49	The Park	82	1.79	1.19	0.37	0.2593	0.2469
50	WF Blacktail Deer	94	0.25	0.58	2.59	0.1186	0.1129
51	WF Blacktail Deer upper	94	0.25	0.58	2.59	0.1186	0.1183
52	Harry Canyon	82	0.88	1.28	1.41	0.2784	0.2901
53	Ratfink EOS	86	0.02	0.15	6.40	0.0324	0.0294
	Average	83.10	0.62	0.89	2.98	0.19	0.19

Table B2. Calculation of mean, standard deviation, skew, sample 95% confidence intervals and bootstrap 95% confidence intervals for % woody use data at 35 MIM DMAs. Average for all DMAs is shown at the bottom of the table.

Number	DMA	N	Mean	Stdev	Skew	Sample 95% CI	Bootstrap 95% CI
1	Burr Creek	23	13.6	12	3.06	4.9182	4.5455
2	campaign creek	23	14.5	12	2.60	5.1141	5.0000
3	Black Horse Butte	32	11.3	5	3.73	1.7582	1.6129
4	Corral Creek EOS	22	15.7	11	1.92	4.7955	4.2857
5	Crippen Creek EOS	36	23.7	22	1.64	7.2270	6.9514
6	Elk Cr Dwnst outside Exclosure	21	46.3	27	-0.10	11.7657	11.3781
7	Ferris Creek EOS	25	33.3	21	0.62	8.3976	8.3333
8	Fire Creek	36	14.0	13.5	3.35	4.4881	4.5714
9	Harry cny	26	24.1	21	1.58	6.9541	8.0011
10	Indian Creek EOS	20	40.5	23	0.39	10.5423	11.0526
11	Little Malheur R	36	28.3	29	1.29	9.6917	9.7214
12	Lower Big Creek	36	10.6	3	5.92	1.1200	1.1429
13	Crane Creek	13	26.7	24	1.15	13.5045	12.5000
14	Lower Indian Jack Cr	36	10.6	3	5.92	1.1200	0.8571
15	Marks Exclosure	36	11.5	6	5.9161	1.7920	1.3714
16	N Fk Mill Creek	11	12.0	6	3.16	3.9199	3.0000
17	Dixie Cr Meadow1	54	32.6	24	0.58	6.4265	6.0472
18	North Fork Malheur River	24	34.8	31	0.98	12.7306	12.3967
19	Pearl Creek Exclosure	36	16.9	15	2.24	5.0673	4.8643
20	Pinto Creek	31	13.2	10	3.22	3.4627	3.5299
21	Reservoir Creek	36	10.6	3	5.92	1.1200	0.8571
22	Ratfink Creek	124	14.1	9	2.05	1.5650	1.4634
23	Rio Bonito	36	36.9	29	0.65	9.7675	9.7214
24	Rock Creek	36	55.7	20	-0.23	6.5418	6.2857
25	SF Beaver	36	10.1	3	2.37	1.0231	1.0714
26	Summit Creek	32	34.5	29	0.67	10.0737	10.0000
27	Taylor Cr Lower	25	17.2	14	2.53	4.6381	5.4236
28	WF Mill Creek	14			0.68	12.1240	10.7692
29	WF Blacktail Deer upper	59	12.1	6	2.6251	1.4267	1.3793
30	Berar Creek lower	36	10.6	3	5.9161	1.1200	0.8571
31	Eighteenmile	36	11.1	5	3.9889	1.5604	1.4286
32	MARKS	36	12.9	11	4.5912	3.6444	3.4286
33	Pearl Creek	19	11.1	5	4.2426	2.1777	1.6667
34	SF Flatwillow	98	10.6	3	5.5045	0.6926	0.7216
35	Big Creek Upper	36	20	21	2.1256	6.8880	6.8571
	AVERAGE	35	21	14	2.65	5.40	5.23

APPENDIX C – BLANK DATA FORMS

Multiple Indicator Monitoring Field Data Sheet—Part 1, Site Information

Allotment		Forest/District		Ranger District/ Field Office	Observers	Date
Sample Point Spacing (M)	Starting Distance (M)	Slope Class*	Substrate Class **	Stubble Height Recorded in (I) inches or (C) centimeters	Plant Region	Subwatershed (6 th Field HUC)
Downstream Marker			Upstream Marker		Reference Marker	
Latitude		Longitude	Latitude		Longitude	Latitude
Zone		UTM		UTM		UTM
Woody Species						
1. Are woody plants supposed to be present at this site? (Y/N)				2. Are there any hydric woody plants present? (Y/N)		
3. Are all age classes of hydric woody plants present? (Y/N)				Comments:		
DMA Site Selection Criteria						
1. Was the riparian complex selected by an interdisciplinary team? (Y, N, or N/A)				2. Is the DMA in a riparian complex that represents management activity and is accessible to the activity? (Y, N, or N/A)		
3. Is the DMA in the riparian complex most sensitive to management? (Y, N, or N/A)				4. Is the DMA impervious to disturbance? (Y, N, or N/A)		
5. Will the DMA site respond to management? (Y, N, or N/A)				6. If the stream is over 4 percent, gradient, does it have a well-developed floodplain? (Y, N, or N/A)		
7. Is the DMA a livestock or activity concentration area? (Y, N, or N/A)				8. Is the DMA compounded by several management activities? (Y, N, or N/A)		

Narrative:

*Slope class: less than 0.5%, 0.5 to 2%, 2 to 4%, >4%, and >10%

** Substrate class: bd(boulder), cb (cobble), gr(gravel), cons (consolidated sand/silt/clay), nonc (nonconsolidated sand/silt/clay), br (bedrock)

Page ____ of ____

***F-fracture, SP-slump, SF-slough, E-eroding, or A-absent**
Note: Usually eight Field Data Sheets, Part 2 are needed for each monitoring site.

Multiple Indicator Monitoring Field Data Sheet—Part 3, Substrate

DMA:		Allotment:				Pasture:						
Stream:				Date:			Used Gravelometer (Y or N)?					
Plot No.	Pebble (mm)										Pool (P) Riffles (R)	Indicate which pebbles were estimated*
	1	2	3	4	5	6	7	8	9	10		
2												
4												
6												
8												
10												
12												
14												
16												
18												
20												
22												
24												
26												
28												
30												
32												
34												
36												
38												
40												
42												
44												
46												

Multiple Indicator Monitoring Field Data Sheet—Part 4, Residual Pool Depth and Pool Frequency

[illegible]

Multiple Indicator Monitoring Field Data Sheet—Part 5, Comments

[illegible]

APPENDIX D – MIM DATA ANALYSIS EXAMPLES

The following examples are provided to help users understand how to proceed when data they have collected do not meet the rules of statistics for generating a confidence interval; basically, when the data distribution is non-normal and/or when the samples collected on the greenline are not independent (spatially autocorrelated).

PART 1: Non-normal distribution

From the “Instructions” tab, data are uploaded to the Data Analysis Module. This example from the Bluebucket Creek DMA collected in 2010, represents a typical scenario involving historically collected data. The following macros are executed in the numbered order indicated in figure D1.

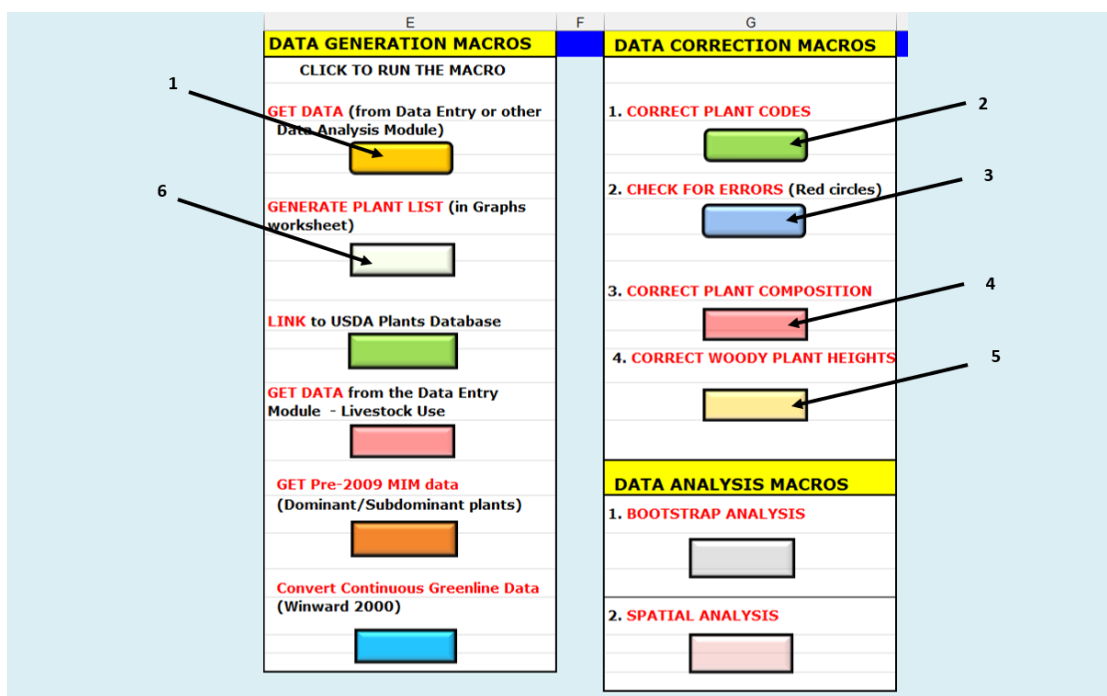


Figure D1. The order in which macros are executed to upload (1) and check for errors in the data from an historical DMA (2,3,4,5), and finally to generate the plant list (6).

Running these macros in the order indicated is necessary to process the data so that all is ready for running the data analysis macros – Spatial autocorrelation and Bootstrap analysis. For data that do not fit a normal probability distribution, the bootstrap analysis is required to generate the appropriate confidence interval (CI). To view the data distributions and determine if they

are non-normally distributed select the “Link to short-term data distributions” as shown in figure D2.

Use this link to view the data distribution

	A	B	C	D	E	F	G	H	I	J	K	L	M
1	METRIC DATA SUMMARY		DMA = Bluebucket 1			LINK TO PROPER FUNCTIONING CONDITION (PFC) ANALYSIS					LINK TO STUBBLE HEIGHT ANALYSIS		
2			Pasture = Teepee			LINK TO GRAPHS WORKSHEET					LINK TO SHORT-TERM DATA DISTRIBUTIONS		
3	SHORT-TERM INDICATORS		Date = 9/22/2010			LINK TO CORRELATION MATRIX					LINK TO SUBSTRATE ANALYSIS		
4	Stubble Height					Woody Use			Streambanks				
	Median SH for all key species (in)	Average SH for all key species (in)	Average SH for all Key species GRAZED (in)	Average SH for all Key species UNGRAZED (in)	Dom key species for SH	Avg SH of dom key species (in)	Woody species use - all woody species MEDIAN (%)	Woody species use - all woody species mean (%) (bootstrap)	Streambank alteration (%) altered	Streambank stability (%)	Streambank cover (%)	Covered - stable (%)	Covered - unstable (%)
5													
6	6.00	6.5			GLGR	7.09	10		9.97	70%	83%	70%	12%
7	n=	51	0	0	23		4	4	81	81	81	57	10
8	95% conf Int ¹		*	*	*	*	*	*	*	*	*	*	*
9	95% CI ²	1.04						15%	6%	8%	7%		

Figure D2. Summary data for the short-term indicators at Bluebucket Creek (2010). Note the “Link to short term indicators” used to examine the data distributions for Stubble height, Bank alteration, and woody riparian species use.

Note that the Data Summary table does not yet have the 95% confidence interval data populated in row 8. The purpose of this exercise is to conduct the analysis so that those important data can be generated.

Upon selection of the link, the graphic shown in figure D3 is displayed for stubble height measurements at Bluebucket Creek.

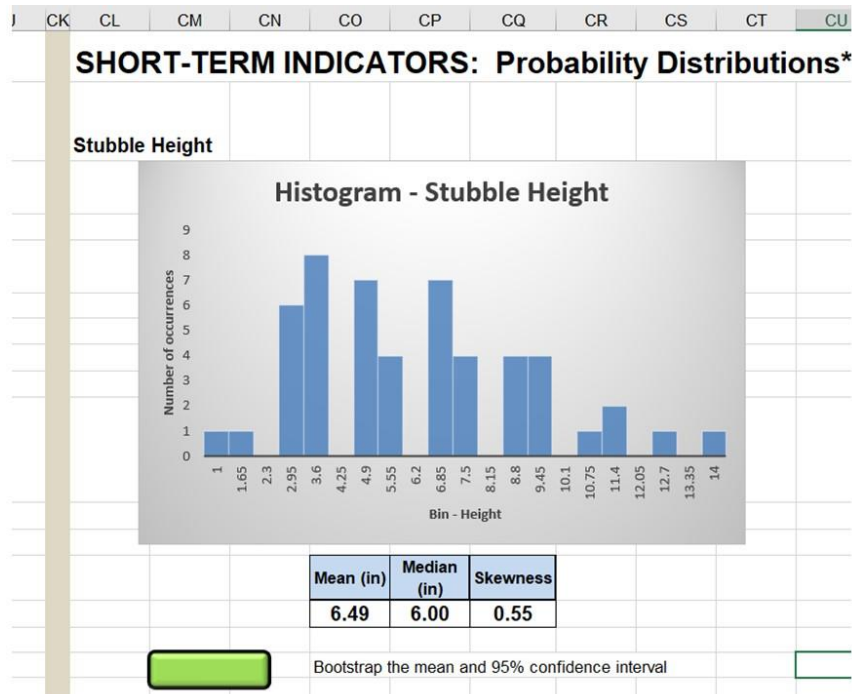


Figure D3. Probability distribution for stubble height at Bluebucket Creek in 2010. Note that the data are skewed to the right (positive skew). The skewness coefficient is .55. The data are considered non-normal if the skewness coefficient is greater than .5 or less than -.5. In this case there is a desire to bootstrap the data by selecting the green button to run that routine for stubble height.

Pressing the green button shown in figure D3 generates the bootstrap analysis for the stubble height data as shown in figure D4.

	A	B	C	D	E	F	G	H	I	ALV
1	Sample data	Boot data	Conf level	Lower bound	Upper bound	stdev	n		Random re-samples	
2	7	For the mean	95%	5.705	7.255	0.407	1000			
3	3	For the median		5	7.000	0.783	1000			
4	5	Bootstrap CI	For the mean	0.77	For the median	1.00				
5	8	Sample data	Row numbers (N)	52	51					
6	5		Sample mean	6.49		Sample median	6.00		Stubble height	☐
7	6		Sample SD	3.02		Bootstrap median	6.00		Bank alteration	☐
8	13		Sample 95% CI	0.83					Woody ri sp use	☐
9	9		Sample skew	0.55					All	☐
10	11									
11	5									
12	14									
13	9		Sample Mean	6.49		Sample Median	6.00			
14	7		Bootstrap mean	6.48		Bootstrap median	6.00			
15	10		Sample SD	3.02		Sample SD	3.02			
16	8		Bootstrap SD	0.41		Bootstrap SD	0.78			
17	3		Sample 95% CI	0.83		Sample 95% CI	0.83			
18	9		Bootstrap CI	0.77		Bootstrap CI	1.00			
19	5		Sample Skew	0.55		Sample Skew	0.55			
20	5		Bootstrap Skew	-0.07		Bootstrap Skew	-0.15			
21	7									
22	2									
23	3									
24	4									
25	1									
26	4									
27										
28										
29										
30										
31										
32										
33										
34										
35										
36										
37										
38										
39										
40										
41										
42										
43										
44										
45										
46										
47										
48										
49										
50										
51										
52										
53										
54										
55										
56										
57										
58										
59										
60										
61										
62										
63										
64										
65										
66										
67										
68										
69										
70										
71										
72										
73										
74										
75										
76										
77										
78										
79										
80										
81										
82										
83										
84										
85										
86										
87										
88										
89										
90										
91										
92										
93										
94										
95										
96										
97										
98										
99										
100										

Figure D4. Results of bootstrap analysis for stubble height at Bluebucket Creek showing both the sample mean/median/CI and the bootstrap mean/median/CI. ME is the margin of error, or the distance between the mean/median and the lower or upper bounds of the confidence interval.

The mean value for the sample (6.49 inches) changed little as a result of the bootstrap re-sampling (6.48 inches). Fortunately, the confidence interval changed in the right direction, from .83 inches for the skewed sample data, to .77 inches for the bootstrapped data. Thus, the final values are - mean of 6.38 inches plus and minus .77 inches or a range of 5.71 for the lower bound of the confidence interval and 7.26 for the upper bound. Note that bootstrap analyses are also available for bank alteration and woody riparian species use.

PART 2: Spatial autocorrelation.

Now that the data have been bootstrapped to comply with the rule for normality, all is not finished. It is also necessary to test for independence. Are the data spatially independent, or is there spatial autocorrelation? As described in Appendix A, adequate spacing of sampling units (e.g., quadrats) is needed to treat a systematic sample as if it were random. The placement of quadrats along the greenline by small distances practically ensures that adjacent sampling units will be spatially autocorrelated. This will result in an underestimation of the standard error, and therefore the confidence interval. So, it is necessary to test for spatial autocorrelation before finally deciding on a confidence interval to apply. This is done in the "Spatial" tab. As shown in figure D5, the instructions table, stubble height can be selected to execute the analysis for that indicator.

	AF	AG	AH	AI	AJ	AK	AL	AM	AN	AO	AP	AQ	AR
1													
2													
3													
4													
5													
6													
7													
8													
9													
10													
11													
12													
13													
14													
15													
16													
17													
18													
19													
20													
21													

LIST OF INDICATORS - SELECT INDICATOR

Ecological status

Wetland Rating

Veg Stability

Average Woody Use per sample point

Bank Stability Rating

Bank Alteration Rating (hits per sample)

Stubble Height (average per sample)

Greenline-greenline width

INSTRUCTIONS

1. Click on the "SELECT INDICATOR" button to copy data for the indicator of interest into the data cells.

2. The indicator is displayed in cell A3. Correlation coefficients for each column are displayed in row 3.

3. Mean, standard deviation, and 95% confidence interval are then displayed in the table K2 to S6.

4. The Correlogram shows how correlation coefficient is changing as Sample point spacing increases (adjacent Sample points, then every other Sample point, then every third Sample point, etc.). If the correlation coefficient is decreasing as Sample point spacing increases, this may indicate spatial autocorrelation. If spatial autocorrelation is present then:

Look at the interpolated distance for various values of r in the table p31 to s37. The Sample point spacing in the box at cell r29 is the interpolated distance at $r < .3$ - indicating an insignificant correlation where spatial autocorrelation is expected to be absent.

Figure D5. Indicators table with instructions for running macros to test for spatial autocorrelation.

By clicking on the button, the following table is provided (figure D6).

	A	B	C	D	E	F	G	H	I
1	DMA: Bluebucket		Link to INSTRUCTIONS			Select indicator	See column AU for significance test		
2	INDICATOR:		Adjacent sample points	Every other sample point	Every third sample point	Every fourth sample point	Every fifth sample point		SAMPLE spacing:
3	Stubble height	CORRELATION COEF.	0.3353	-0.7941	0.3155	0.5873	-0.1215	(Left bank of stream)	
4			0.4479	0.0542	0.0811	-0.0494	-0.5154	(Right bank of stream)	
5		N	48	47	46	45	44		
6	Sample point		Stubble height	Stubble height	Stubble height	Stubble height	Stubble height	DATA: EVERY OTHER PLOT	DATA: EVERY 3rd PLOT
7	1	7.0							
8	2		7						
9	3			7				7.0	
10	4				7				
11	5					7			
12	6						7		
13	7								
14	8	3.0							
15	9	5.0	3						
16	10		5	3					5
17	11	8.0		5	3				
18	12	5.0	8		5	3			
19	13	6.0	5	8		5	3		5
20	14		6	5	8		5	5.0	
21	15			6	5	8			
22	16				6	5	8		
23	17					6	5		
24	18						6		
25	19								
26	20								
27	21	13.0							

Figure D6. Correlation table for stubble height at Bluebucket Creek showing correlation coefficients for left and right banks separately in rows 3 and 4.

Note that correlation coefficients are highest with respect to adjacent sample points and are much smaller for every other sample point, which are located twice as far from each other than adjacent sample points. There are 80 sample points, so much of the data is not shown in this graphic and note also that a number of sample points had no recording of stubble height.

To the right of the correlation table are the graphs of spatial autocorrelation showing how correlations between adjacent, every other, every third, etc. sample point change. This is displayed in figure D7.

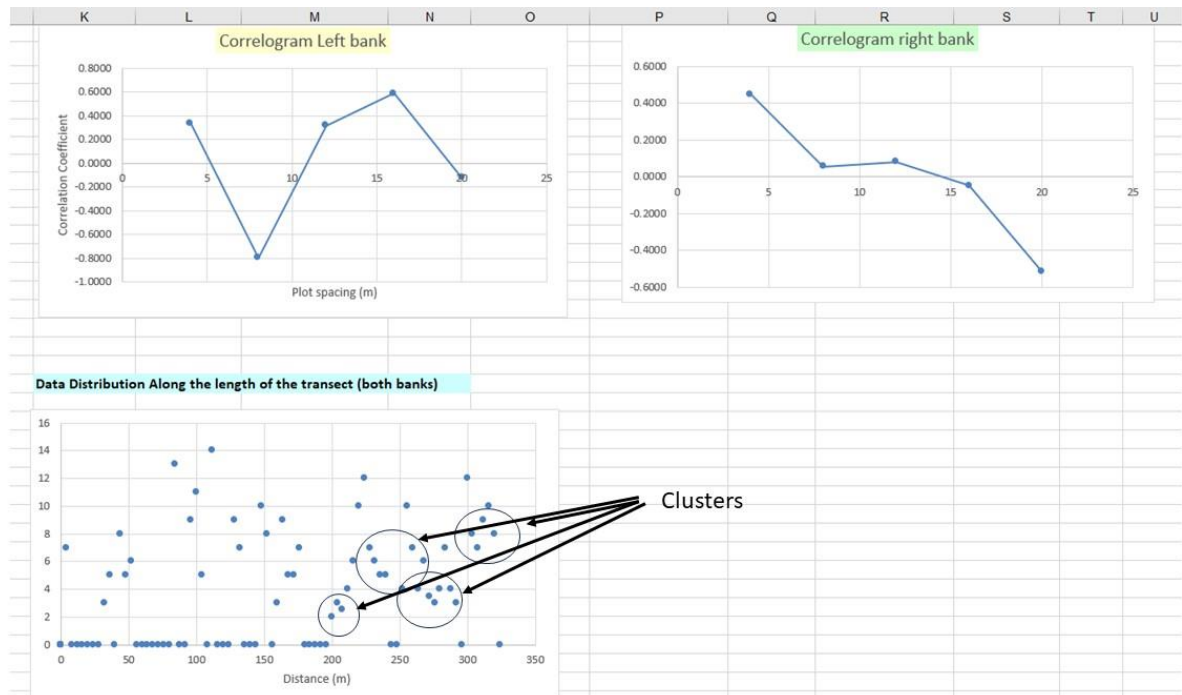


Figure D7. Correlograms for stubble height at Bluebucket Creek including correlograms for both left and right banks and the data distribution graph of stubble heights at varying distance along the greenline.

Note that in this figure, the right bank is clearly indicating spatial autocorrelation since correlation coefficients are declining with distance between sample points. The left bank appears more erratic possibly because of the number of zeros (not stubble height measurement) at a number of the sample points on this bank (as shown in the distribution graph). Note also that clustering is clearly evident in the right bank (right side of the distribution graph). These graphs help to indicate the presence of spatial autocorrelation.

Further to the right in the “Spatial” tab are the test summaries for all of the indicators, including stubble height. For Bluebucket Creek, this is shown below in figure D8.

	BH	BI	BJ	BK	BL	BM	BN	BO	BP	BQ	BR
Spatial autocorrelation - adjacent sample points						Spatial autocorrelation - Every other sample point					
Metric/Indicator	Correlation Coefficients		p values				Correlation Coefficients		p values		
	Left bank	Right bank	Left bank	Right bank			Left bank	Right bank	Left bank	Right bank	
Ecological status	-0.60		0.00				-0.06		0.69		
Wetland rating	0.07	0.49	0.67	0.00			-0.08	0.31	0.64	0.06	
Veg Stability	0.27	0.41	0.10	0.01			0.06	0.21	0.69	0.21	
Woody Rip Spec Use											
Bank stability	0.11	0.42	0.49	0.01			0.10	0.19	0.52	0.25	
Bank alteration	-0.05	0.12	0.76	0.49			-0.04	0.08	0.80	0.64	
Stubble height	0.34	0.45	0.03	0.00			-0.79	0.05	0.00	0.75	
GGW	0.45	0.33	0.00	0.04			0.10	0.20	0.54	0.22	

Figure D8. Spatial autocorrelation tables for both adjacent and every other sample points for all indicators. Note that Woody riparian species use was not measured at this DMA.

As shown in these tables, the correlation coefficients for stubble height are: .34 and .45 respectively for left and right banks of adjacent sample points, and -.79 and .05 respectively for left and right banks of every other sample point. Using the t-score test of significance, these produced p values of .03 and .00 for adjacent sample points, and .00 and .75 for every other sample point. The correlation is considered significant statistically when the p value is less than .05. Thus, only every other sample point on the right bank satisfies the test for independence.

The “Spatial” tab also provides the mean, sample size, and confidence interval for adjacent and every other sample points, for both left and right banks as shown in figure D9.

	BH	BI	BJ	BK	BL	BM	BN
33							
34							
35	LEFT BANK	Metric value		CI		N	
36	Metric/Indicator	Every SP	Every other SP	Every SP	Every other SP	Every SP	Every other SP
37	Ecological status	44.24	46.65	7.87	11.54	37	17
38	Wetland rating	70.64	74.78	8.57	11.50	37	17
39	Veg Stability	4.76	5.01	0.83	1.20	37	17
40	Woody Rip Spec Use	30.00	30.00	22.63		3	1
41	Bank stability	29.73	29.41	14.93	22.33	37	17
42	Bank alteration	0.76	0.86	0.09	0.12	37	17
43	Stubble height	4.14	4.29	1.10	1.68	14	7
44	GGW	6.20	6.05	1.16	1.73	36	17
45							
46							
47	RIGHT BANK	Metric value		CI		N	
48	Metric/Indicator	Every SP	Every other SP	Every SP	Every other SP	Every SP	Every other SP
49	Ecological status	60.57	59.25	8.33	12.07	38	18
50	Wetland rating	81.28	74.72	8.67	15.02	38	18
51	Veg Stability	6.75	6.43	0.85	1.36	38	18
52	Woody Rip Spec Use					0	0
53	Bank stability	50.00	50.00	16.11	23.77	38	18
54	Bank alteration	0.46	0.40	0.09	0.13	38	18
55	Stubble height	4.91	4.62	1.08	0.69	28	13
56	GGW						
57							

Figure D9. Mean, Confidence interval (CI), and sample size (N) for adjacent (every) sample point, and every other sample point for both left and right banks.

Note on this table, that every other sample point on the right bank has a mean stubble height of 6.7 inches and a confidence interval of 1.53 inches from 15 samples. This reduced number of samples produces a much wider range of the CI than that of the bootstrapped data which resulted in a mean of 6.38 inches and a CI of 0.77 inches.

The “Spatial” tab also produces a summary table for the test results indicating which scenarios (sample sets for both, left, and right banks) have spatial autocorrelation as shown in figure D10.

STUBBLE HEIGHT					Values if not autocorrelated			Values regardless of autocorrelation			
Both banks	r	t score	p	Significant?	Metric value	95% CI	N	Metric value	95% CI	N	Both banks
Adjacent samples	0.54	3.86	0.00	Y				15.6	1.17	54	Adjacent samples
Every other Sample	0.36	1.95	0.06	N	15.7	1.66	24	15.7	1.66	24	Every other Sample
Every third Sample	0.15	0.75	0.46	N	15.9	2.1	17	15.9	2.06	17	Every third Sample
Left bank											Left bank
Adjacent samples	0.54	3.86	0.00	y				16.1	1.62	35	Adjacent samples
Every other Sample	0.30	1.91	0.06	N	15.8	2.5	15	15.8	2.45	15	Every other Sample
Right bank											Right bank
Adjacent samples	0.52	3.04	0.01	y				14.7	1.43	19	Adjacent samples
Every other Sample	0.36	1.95	0.06	N	15.6	1.9	9	15.6	1.87	9	Every other Sample
Values associated with highest N, not autocorrelated:					15.7	1.66	24				

Figure D10. Table of test results showing which scenarios have spatial autocorrelation and the statistical values associated with each.

In this table for stubble height, note that every other sample on both banks did not indicate spatial autocorrelation and this scenario had the highest sample size for scenarios showing “N” in the “Significant?” column. This is a result of the fact that spatial autocorrelation was not observed on either the left or right banks for every other sample point.

These results are then sent to a table representing all of the indicators as shown in figure D11. As indicated, the results are automatically sent to the “Data_summary” tab.

RESULTS having no spatial autocorrelation SENT TO DATA SUMMARY TAB			
INDICATOR/METRIC	CONFIDENCE INT*	METRIC VALUE	N
Ecological status	4.21	98.01	39
Wetland rating	5.36	78.47	68
Veg Stability	0.25	7.98	39
Woody Rip Spec Use	0.12	10.06	16
Bank stability (%)	6.9%	91.0%	67
Bank alteration (%)	4%	10%	67
Stubble Height	1.66	15.69	24
GGW (m)	0.31	3.73	66

Figure D11. Results table showing the statistical values that are sent to the Data Summary tab. These are the best scenarios having little likelihood of having spatial autocorrelation.

The question may be asked, what if there was no scenario in which spatial autocorrelation was NOT indicated? In this case, the system automatically reverts to using the 95% confidence interval for user variation as shown on the “Data Summary” tab for each indicator and referred to as “95% C²”. Thus, both the “95% C²” and “95% C¹” values will be the same on that tab. The metric value will be represented by the mean of all samples.

STEP-BY-STEP DATA ANALYSIS PROCEDURE

To help users understand how to conduct MIM data analysis, the following list of steps is provided. This will allow for a complete analysis of the data and derive a confidence interval that satisfies the rules of statistics. This applies to both new and historical data uploaded to the module.

Step 1. Upload data using the “GET DATA” macro on the Instructions tab.

Step 2. Correct the data using the correction macros on the Instructions tab. Start with the “CORRECT PLANT CODES” routine followed by “CHECK FOR ERRORS”, then “CORRECT PLANT COMPOSITION” and “CORRECT WOODY PLANT HEIGHTS”.

Step 3. Run the “GENERATE PLANT LIST” macro to populate the plant metrics in the module.

Step 4. Run both “DATA ANALYSIS MACROS” to process all the indicators through the bootstrap and spatial analysis routines.

Step 5. Analyze the data distributions using the link to “Short-term data distributions” on the Data_summary tab. Check to see if the skew is greater than .5. If so, go to the “Boot” tab to examine the metric value and confidence intervals generated by the bootstrap analysis. This involves any of the 3 short-term indicators (See figure D4).

Step 6. Examine the test results for Spatial autocorrelation in the summary tables depicted in figure D10. It is particularly helpful to examine the test results for the various sub-sets or scenarios tested in the “Spatial” tab starting at cell BV14. This table shows which statistical values will be chosen by the system to be represented on the Data Summary tab.

Step 7. If there is a question about the spatial autocorrelation for any one of the indicators, then run the macro for that indicator using the table at the top of column AP.

Step 8. If a macro was run for any indicator(s), examine the graphics in the correlograms and spatial distribution graphs (as in the example of figure D7) to see if there is a declining trend in

the correlation coefficient as samples get farther apart (i.e. adjacent, every other, every third, etc.). Also, examine the spatial distribution to see if there is clustering of the data. These observations can help determine where along the greenline spatial autocorrelation is occurring, such as along a portion of one of the streambanks.

APPENDIX E – DMA MODIFICATION CALIBRATION

For several valid reasons (see Appendix B in the MIM technical reference), it may be necessary to shift the location of the DMA (within the same riparian complex) or to expand the length of the original DMA (e.g. from 110 m to 150 m). This appendix describes the procedure for calibrating the MIM indicators to permit trend analysis in situations where the original DMA has been shifted or expanded in length.

To calibrate the results for the new DMA, the approach can use statistical methods to ensure consistency between the datasets. For example, data were collected historically on a 110-m DMA. Now, with the desire to increase sample point spacing from the previous 2.5 m to the new 3.75 m minimum spacing required under the updated protocol, the DMA expands to 150 m in length adding 40 m of stream that previously were not within the DMA. The following discussion describes the approach.

1. Comparison of Means and Variances:

- **Means:** Compare the average metric indicator measurements between the two DMAs. A confidence interval test (for comparing the two means) will help assess whether there is a significant difference in the metric value such as mean stubble height. The confidence interval test can be interpreted as illustrated in figure E1.
- **Variances:** Perform an F-test to compare the variance in measurements between the two DMAs. If the variances are similar, you can proceed to combine the datasets or apply other comparisons confidently. Use the Statistical Function “F-Test Two-Sample for Variances”. The output can be interpreted as described in figure E2.

2. Resampling:

- One way to directly compare the new DMA to the historical one is to "resample" the new DMA to the same length as the old one (i.e., take the first 110 m of data from your 150-m transect). Then, apply statistical tests to determine if the two subsets are consistent. The typical test used in the MIM applies the 95% confidence interval (CI). If the CI of the mean of one DMA overlaps with the CI of the mean from the other DMA, then we can confidently conclude that the two transects are not statistically different.

3. Scaling or Normalization:

- If differences exist, a good approach is to apply scaling or normalization. This could involve adjusting the new data based on the ratio of means between the two transects, making them more directly comparable. Thus, if they are significantly different, then there is a need to apply an adjustment (or scaling) to the new DMA based on the ratio of the metric value at one to the same value at the other. For example, if bank stability is

84% at the new DMA and only 79% at the historic DMA, with a CI of 2%, then the two DMAs are significantly different. In this case, the ratio 84:79 or a factor of 1.06 could then be applied to all the historical measurements to compare to values obtained at the new DMA. Simply multiply the bank stability values of those historical measurements by 106%. Or on the other hand, all future measurements at the new DMA could be multiplied by 79/84 (or 94%) to make the measurements there comparable to the historical measurements at the original DMA.

4. Example:

Use the MIM Statistical Analysis Module to calibrate a changed DMA to a previously used historical DMA. This module accommodates comparison of multiple DMAs. The following describes the calibration steps:

- Step 1: Upload data from the historical DMA Data Analysis Module into the Statistical Analysis Module using the Get Data macro on the Get Data tab. Make sure the Pasture name includes a reference to the name of the original DMA (e.g. 110-m DMA).
- Step 2: Upload data from the new DMA Data Analysis Module into the Statistical Analysis Module using the Get Data macro on the Get Data tab. Make sure the Pasture name includes a reference to the name of the new DMA (e.g. 150-m DMA).
- Step 3: Go to the “Comp” tab and compare the results from the two DMAs. Table E1 gives an example of a table of comparisons from the “Comp” tab for Crooked Creek. If any of the indicators of interest appear to be significantly different, prepare a bar graph like the one in figure E1.
- Step 4: If the results from the original and the relocated or extended DMA are statistically significantly different, then calculate a calibration factor. Divide the mean of the original DMA by the mean of the new DMA to obtain this calibration factor (see Part 3: Scaling and Normalization, above).
- Step 5 (Optional). Verify that the conclusions of the confidence interval test are valid by running an F-test for equal variances on the same data. Open the “F-test” tab (in the Statistical Analysis Module) and copy the relevant data from the “Get Data” tab into columns A and B of the F-test tab. Follow the instructions for running the test.

Table E1. Comparison of data collected at the same time of the original 110-meter DMA and the new 150-meter DMA on Crooked Creek. The 95% confidence interval is used to determine if the indicators are statistically different.

<i>DMA</i>	Average SH for all key species (cm)	Average SH for all Key species GRAZED (cm)	Average SH for all Key species UNGRAZED (cm)	Dom key species for SH	Avg SH of dom key species (cm)
150 m DMAcrooked creek	6.58	4.45	9.64	CAAQ	5.45
110 m DMAcrooked creek	6.58	4.92	9.61	CAAQ	6.19
95% Confidence Interval	1.1	1.0	1.2		1.1
<i>DMA</i>	Woody species use - all woody species mean (%)	Streambank alteration (% altered)	Streambank stability(%)	Streambank cover (%)	
150 m DMAcrooked creek	36	9%	85%	100%	
110 m DMAcrooked creek	36	9%	85%	97%	
95% Confidence Interval	10.1	3.2%	8.0%	8.5%	

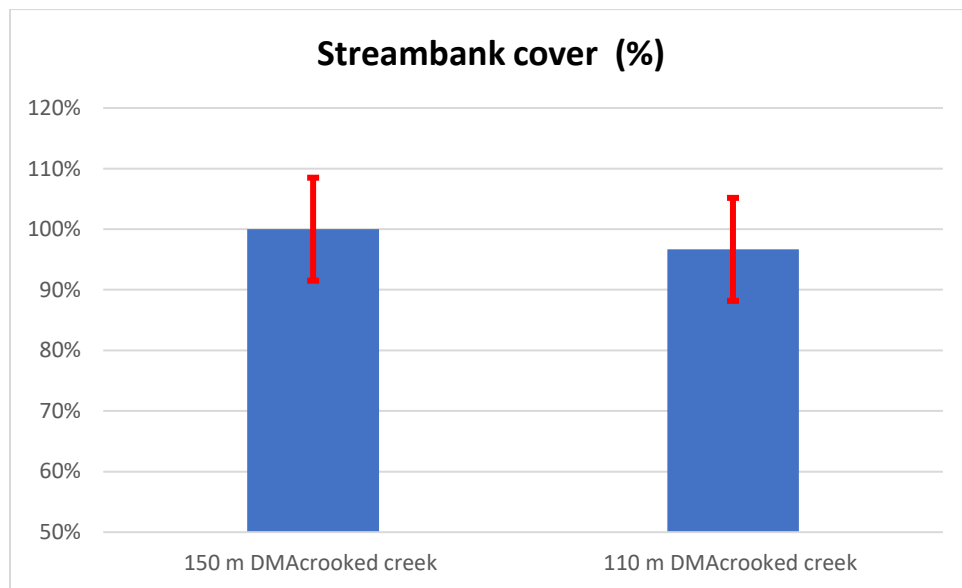


Figure E1. Graph of percent streambank cover with error bars representing the 95% confidence interval taken from the data in Table E1. Note that the error bars are strongly overlapping suggesting that there is no evidence of a statistically significant difference.

F-Test Two-Sample for Variances		
	<i>150 m</i>	<i>110 m</i>
Mean	5.45098	6.189189
Variance	11.01255	11.43544
Observations	51	37
df	50	36
F	0.96302	
P(F<=f) one-tail	0.445067	

Figure E2. Example of an F-test for stubble height comparing two overlapping DMAs. This tests the probability that the observed F-statistic is greater than or equal to a certain critical value under the null hypothesis. If this p-value is small (typically less than 0.05), it suggests that the variances are significantly different, and you reject the null hypothesis that the variances are equal. In this case we do NOT reject the null hypothesis and conclude that they are equal.

Example using Step 4 – there is a significant difference between a metric at the new DMA as compared to the original DMA.

Suppose a DMA has been altered, for example by channel adjustments, beaver dam installations, or other natural changes that no longer make the DMA effective for monitoring. A new DMA has been randomly located elsewhere in the riparian complex. Data were collected at the new DMA and compared with the original DMA resulting in the comparison of metric indicators as shown in Table E2.

As indicated in Table E2, Greenline Stability rating was significantly different between the two DMAs. The following calibration is performed on this indicator: Winward Greenline Stability rating: Original DMA/New DMA = 114%. All future ratings are multiplied by 1.14 to make them comparable to the original DMA.

No adjustment is made for any other indicator because there is no evidence of a significant difference between the two DMA.

Table E2. Comparison of metric indicators at the original and new DMAs evaluated for calibration. A test of significance was performed using the 95% CI values from the Data Analysis Modules for these two DMAs. Results indicated that all indicators were not significantly different except for Winward Greenline Stability Rating.

DMA	Average SH for all key species (in)	Woody species use - all woody species MEDIAN (%)	Streambank alteration (%) altered)	Streambank stability(%)	Streambank cover (%)
Original	19.2	11.9%	20%	90%	96%
New	18.0	13.8%	28%	86%	87%
Difference	1.2	2%	8%	4%	9%
95% CI	1.40	6%	7%	5%	5%
Significantly different?	N	N	N	N	N
DMA	Greenline ecological status rating	Site wetland rating	Winward greenline stability rating	Greenline-greenline width (m)	
Original	77	91	7.76	3.55	
New	68	85	6.78	3.45	
Difference	8.8	5.5	1.0	0.1	
95% CI	5.75	3	0.16	0.23	
Significantly different?	N	N	Y	N	

APPENDIX F. Simplified explanations of statistics used in the Data Analysis Module

A. BOOTSTRAP ANALYSIS

For stubble height and other metrics that calculate the average (mean) from the data collected at any DMA, there is a certain amount of uncertainty associated with it. In other words, is this the true mean? To address this uncertainty, the system calculates a **95% confidence interval (CI)**. A confidence interval gives a range of values that likely contains the true average of the data.

However, the stubble height, and in particular the streambank alteration and riparian woody species browse data may not be **normally distributed**, meaning they do not follow the typical bell-shaped curve in their frequency distributions. You can see such a curve by clicking on the “Short-term data distributions” link in the Data summary tab. Some users have asked, why don’t we just use the median rather than the mean when the data are not normally distributed? Here’s why:

1. The Mean Uses More Information

- The mean takes into account every value in the dataset, making it a more comprehensive measure of central tendency.
- The median only considers the middle value (or the average of two middle values), ignoring the magnitude of all other data points.

2. The Mean Is More Efficient for Normally Distributed Data

- If the data follow a normal distribution, the mean is the best estimate of the central value because it minimizes statistical error (it has the lowest variance as an estimator).
- The median, while useful for skewed distributions, is statistically less efficient in normal distributions because it does not use all data points effectively.

3. The Mean Is More Useful for Statistical Tests

- Many statistical tests (like the 95% confidence interval) are based on the mean because of its desirable mathematical properties.
- The median does not work as well in these models because it does not account for all values in the dataset and does not follow the same probability properties as the mean.

4. The Mean Allows for Easier Comparisons and Calculations

- The mean is used in many fundamental statistical concepts, like variance and standard deviation, which describe data spread.

- The median does not work well with these measures because it does not factor in each value's contribution.

When the data do not fit a normal distribution, we can use a statistical technique to normalize the distribution, called “bootstrapping”. Bootstrapping is a statistical technique that involves repeatedly resampling the data (with replacement) to create many simulated datasets. From these simulations, we can calculate a more reliable estimate of the 95% confidence interval and the mean. This approach does not assume any particular data shape, making it a better fit for the analysis. In other words, it does not depend on the data being normally distributed.

By using bootstrapping, the system obtains a 95% confidence interval that more accurately reflects the true range of uncertainty in the data. This ensures that the results are as reliable and informative as possible.

For stubble height, when you run the bootstrap analysis in the “Boot” tab, it will calculate the mean and 95% confidence interval for all plants recorded for stubble height on the DMA tab. In other words, all plants that were measured, and this result will be reflected on the Data Summary tab. When you run the bootstrap analysis in the “STHT” tab, it will calculate the mean and 95% confidence interval only for those plant species that you select (in Column G) for the analysis. The Data summary tab displays results for all plants and also the single most dominant key species and the latest version now also reports the results from the STHT tab. This will automatically display the results from the STHT tab in the 2025 version of the Statistical Analysis module. Another way to force the system to use only the selected key species on the STHT tab, is to delete from the raw data on the DMA tab, columns I, J, and K any plants that were not selected for calculation of mean stubble height. For example, a user measures stubble height for plants: POPR, CANE2, AGST, JUAR, and JUEN. The resultant data indicated that 90% of the measurements occurred on CANE2, POPR, and JUAR, so these 3 key species were selected for the analysis on the STHT tab. But you can also delete the data for AGST, and JUEN on the DMA tab to show the same results for just the selected species on the Data Summary tab. However, the STHT tab was created so that none of the data from the original recording has to be deleted in case there is some reason to refer back to them in the future.

Here is a simple example of **resampling with replacement**, which is the key idea behind bootstrapping:

Imagine you collected five stubble height measurements along the greenline:

Original data:

12,15,18,14,20

To perform bootstrapping, we randomly select **five** values from this dataset, **allowing repeats** (replacement). Here's one possible resample:

Resample 1:

15,12,12,18,20

Since we replace each selected value before picking the next one, some values may appear multiple times (like **12** here), while others may not appear at all (like **14** in this resample).

If we repeat this process many times (e.g., 1,000 times as is done in the module), we create many simulated datasets. Each resample has its own mean, and from all these means, we can determine the **95% confidence interval**—the range where the true mean is likely to fall.

This method works well because it doesn't assume the data follow a specific distribution, making it a powerful tool for estimating confidence intervals when normality is in question.

B. Spatial Autocorrelation

It is critical to understand that using confidence intervals to identify what range of values our true mean falls within, the data themselves must be random and independent. Independence in statistics refers to the lack of a relationship or association between two or more variables. In other words, if two variables are independent, then the value of one variable does not affect the value of the other variable.

Independence is an important concept in statistics because it allows making certain assumptions about the behavior of random variables like the MIM monitoring indicators. For example, if we assume that the data are independent, we can use simpler statistical models, like 95% confidence intervals, to analyze the data. However, if two variables are not independent, we need to take into account their relationship when analyzing the data, and that can make the analysis much more complicated. In the application of the MIM, we have found non-independence related to spatial autocorrelation in some of the monitoring indicators.

The MIM approach to analyzing spatial autocorrelation using Pearson correlation coefficients and a correlogram, as seen in the Spatial tab of the Data Analysis Module, is a well-established method in spatial statistics (**Dormann, C. F. et al. 2007, Legendre, P., & Legendre, L. 2012**). Here's a simplified explanation of how it works:

Understanding Spatial Autocorrelation

Spatial autocorrelation measures how similar observations are to each other based on their spatial proximity. If nearby measurements tend to be similar (e.g., stream widths at close locations are similar), there is **positive spatial autocorrelation**. If nearby measurements tend to be dissimilar, there is **negative spatial autocorrelation**. Here is an example.

The MIM Approach: Step by Step

1. **Collect Data:** You measure a MIM indicator at regular intervals along a stream transect (e.g., every 3.75 meters).
2. **Calculate Correlations at Different Lags:**
 - Compute Pearson's correlation coefficient between GGW widths at different distances (lags).
 - For **lag 1(adjacent plots)**, correlate each measurement with the next one.
 - For **lag 2 (every other plot)**, correlate each measurement with the one two plots away.

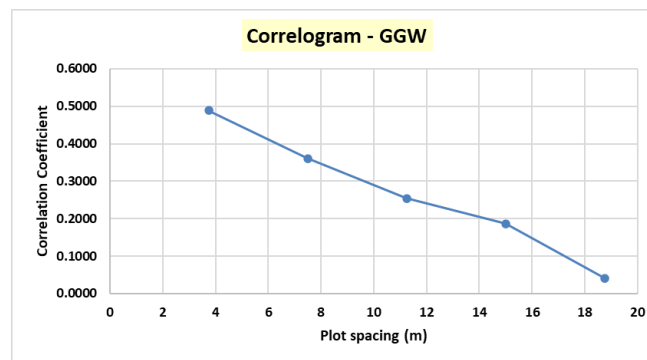
- For **lag 3 (every third plot)**, correlate each measurement with the one three plots away, and so on.

3. Create a Correlogram:

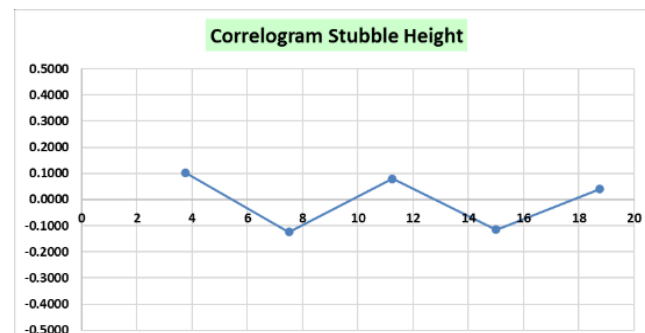
- Plot the correlation coefficients against the lag distances in a correlogram.
- The correlogram shows how the similarity between measurements decreases (or increases) with distance.

Interpreting the Correlogram

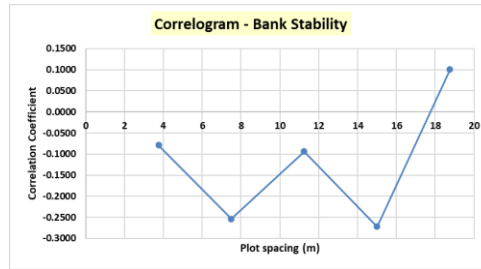
- If correlation is high at short lags and decreases as distance increases, it indicates **strong spatial autocorrelation** (e.g., stream widths change gradually rather than randomly).



- If correlations fluctuate around zero, it suggests little or no spatial pattern.



- If correlations become negative at certain lags, it may indicate a repeating pattern or cyclic variation in the data, but not spatially autocorrelated. Negative correlations tend to indicate that nearby values are dissimilar.



Why It Matters

- **Identifying Spatial Dependence:** Helps determine if nearby stream widths are related.
- **Guiding Statistical Analysis:** Traditional statistical methods assume independence of observations. If autocorrelation exists, adjustments (such as using segment means or spatial lags) may be needed. In the Data Analysis Module, alternative portions of the DMA (segment means) and spatial lags are examined – i.e. adjacent samples on just one side (left or right), every other sample, every other sample on just one side of the stream, etc.).

At what level of correlation (or at what correlation coefficient) indicates that there is spatial autocorrelation in the data?

When you compute Pearson's correlation coefficient (r) between two variables, you often want to test whether this correlation is **statistically significant**, whether it is likely due to a real relationship or just random chance.

Step-by-Step Explanation

1. State the Hypotheses

- **Null Hypothesis (H0):** There is no real correlation; $r=0$
- **Alternative Hypothesis (HA):** There is a real correlation; $r \neq 0$.

2. Compute the t-Score

The t-score for Pearson's correlation is calculated using the formula:

$$t = \frac{r\sqrt{n-2}}{\sqrt{1-r^2}}$$

where:

- r = Pearson correlation coefficient
- n = number of data points (sample size)

3. Find the Critical Value

- The t-score follows a **t-distribution** with $n-2$ degrees of freedom.
- The computed t-score is compared to a critical value provided in EXCEL at the chosen **significance level** (e.g., $\alpha=0.05$ for a 95% confidence level).

4. Make a Decision

- If the absolute value of the computed **t-score** is **greater** than the critical value from the table, you **reject the null hypothesis**, meaning the correlation is statistically significant.
- If the t-score is **less** than the critical value, you **fail to reject** the null hypothesis, meaning there is no significant correlation.

All of this is done for you in the Data Analysis Module – Spatial tab.

Example

Suppose you have **40 measurements** ($n=40$) along the left side of the DMA (GGW is collected only along the left side) and find a Pearson correlation of **0.50** on adjacent measurements.

1. Compute the t-score: Done in the Module and provided in the t-score table:

t-SCORE TABLE				
Adjacent Sample points	r	t score	p value	Significant?
Left side	0.500	3.512	0.001	y
Right side	na			
Every other Sample point				
Left side	-0.1082	0.662	0.512	N
Right side	na			
Every third Sample point				
Left side	-0.2410	1.511	0.139	N
Right side	na			

2. Look up the critical value for **df = 38** at $\alpha=0.05$ (two-tailed test). From a t-table, this t-Score is 3.51. As indicated above, this critical value is calculated in the module and displayed in the t-Score table.
3. The p value for this correlation coefficient ($r=.50$) is .001 and since this is less than .05, the correlation is **statistically significant** at the 0.05 level.

Typically, correlation coefficients for MIM data usually involving 80 samples per DMA are significant when the Pearson correlation coefficient is greater than 0.30.

Why This Matters

- This test ensures that the correlation you observe is **not just random noise**.
- It helps determine **which spatial lags in the correlogram show significant spatial autocorrelation**.

How does the MIM module address situations when spatial autocorrelation is prevalent in the data?

The module examines various scenarios to determine if any are not spatially autocorrelated and then presents the results in a summary table at the far right-hand side of the Spatial tab. The following is an example of a spatial autocorrelation test in which spatial autocorrelation is prevalent in the data (r values are high and this comes from historical data when 80 GGW samples were collected):

GGW					Values if not autocorrelated		
Both banks	r	t score	p	Significant?	Metric value	95% CI	N
Adjacent samples	0.51	3.56	0.00	Y			
Every other Sample	0.43	2.93	0.01	Y			
Every third Sample	0.44	3.02	0.00	Y			
Left bank							
Adjacent samples	0.26	1.65	0.11	N	4.96	0.6	39
Every other Sample	0.43	2.93	0.01	y			
Right bank							
Adjacent samples	0.51	3.56	0.00	y			
Every other Sample	0.43	2.87	0.01	y			
Values associated with highest N, not autocorrelated:					4.96	0.64	39

The scenarios include both banks, left bank only, right bank only, and adjacent samples, every other sample (spaced further apart), and every third sample. Note that for the t-Score tests in these scenarios, only adjacent samples on the left bank were not significant. Then this scenario with no spatial autocorrelation and the highest sample size (n value) was chosen as displayed in red at the bottom of the table.

How do these results (a GGW of 5.0) taken from just part of the sample (39 measurements), compare to the entire sample (79 measurements)?

The module provides a table of values immediately to the left of the table shown above that gives the values for all scenarios. Here is an example for the GGW data above:

Values regardless of autocorrelation			
Metric value	95% CI	N	Both banks
5.1	0.41	79	Adjacent samples
4.9	0.53	38	Every other Sample
5.2	0.67	26	Every third Sample
			Left bank
5.0	0.64	39	Adjacent samples
4.4	0.83	18	Every other Sample
			Right bank
5.3	0.51	39	Adjacent samples
5.4	0.64	19	Every other Sample

Note that the metric value for GGW using all 79 samples was 5.1. This compares with our chosen scenario having a GGW of 5.0 for adjacent samples on the left bank having only a sub-set of the data (n=39). Although this represents a smaller sample and not all of the data collected, its metric value is somewhat comparable to the GGW for the full sample, and therefore a useful metric given that its data set contains samples that are independent. However, also note that the 95% confidence interval is larger for the chosen scenario (.64 compared to .41). This smaller sample size will typically equate to a larger 95% confidence interval, a trade-off that is necessary to avoid spatial autocorrelation.

C. Advantages and disadvantages of using the 95% confidence interval for observer variation

The Data Summary tab displays two kinds of confidence intervals as shown in the following graphic.

19		Substrate:				
		Percent fines	D16 particle size (mm)	D50 particle size (mm)	D84 particle size (mm)	Total number pools
20						
21		52%	0.7	3.93	13	14
22	n=	333	333	333	333	28
23	95% conf Int ¹	7.69%	*	*	*	*
24	95% CI ²	10%				
25	¹ 95% conf Int: 95% confidence interval based the data at this DMA					
26	² 95% CI: confidence interval from observer variation tests (see Ch 3 in the Data Instructions Guide					

The first 95% confidence interval (CI) is based on the data, the second is the confidence interval based on tests of observer variation as described in Chapter 3. The first CI is always preferable. Using the second CI has its advantages and disadvantages.

Advantages:

1. Reflects Measurement Precision:

- The CI on the mean difference directly quantifies observer consistency, making it a better estimate of **measurement precision** rather than variability in the actual measured data.

2. Separates Measurement Error from Natural Variability:

- The CI on the data includes both **natural variability** and **measurement error**. The CI on the observer differences isolates the **error due to measurement** rather than environmental variability.

3. Better for Quality Control:

- If the goal is to assess how well different observers apply the measurement protocol, the CI on the mean difference is more useful than the CI on the overall data, which includes biological and environmental factors.

Disadvantages:

1. May Underestimate Total Variability:

- The precision of the technique (observer repeatability) does not capture the full range of variation in the DMA data. If those data varies greatly due to environmental factors, using only the observer CI might **underestimate** uncertainty in future measurements.

2. Not Directly Applicable for Future Predictions:

- The CI on observer differences does not provide a range for expected future metric values. If the goal is to predict future observations, the CI on the overall data is more informative.

3. Does Not Capture Bias:

- The CI on observer differences assesses precision but does not account for **systematic bias** (e.g., if all observers consistently underestimate or overestimate the value of the indicator).

Conclusion:

If the goal is to **assess the precision of the measurement technique**, the 95% CI on observer differences is appropriate. However, if the goal is to **estimate the range of metric values for the indicator in the future**, then the 95% CI on the full dataset (including all sources of variation) is more appropriate. A combination of both measures may be the best approach, depending on the monitoring objectives.

References

Dormann, C. F. et al. (2007). "Methods to Account for Spatial Autocorrelation in the Analysis of Species Distribution Data." *Ecography*, 30(5), 609–628.

Legendre, P., & Legendre, L. (2012). *Numerical Ecology* (3rd English ed.). Elsevier.