

Statistik i sammanställningar av sort- och ogräsförsök

Seriesammanställningarna som redovisas på SLU Fältforsks websida är redovisningar av statistiska analyser av serier av fältförsök. Dessa redovisningar i Excel kallar vi även *resultatblanketter*. Med *serie* menar vi ett antal försök som vanligen har samma försöksupplägg och som ska analyseras tillsammans. Här beskrivs, kortfattat, hur analysen går till och vad de statistiska resultaten som redovisas på resultatblanketterna betyder.

Databearbetning

Export av data från NFTS till Sas

Först analyseras de enskilda försöken i Nordic Field Trial System (NFTS). I det systemet redovisas också resultat från statistiska analyser av de enskilda försöken.

Serisammanställningen görs på ledvisa observationer och ledvisa medelvärden som har beräknats i de enskilda försöken. Från NFTS exporteras de ledvisa observationerna och de ledvist beräknade medelvärdena till en eller flera xml-filer som sammanlagt innehåller all nödvändig information om seriens försök. Dessa xml-filer läses sedan in i programmet Sas, där beräkningarna för seriesammanställningarna görs. Slutligen skrivs resultaten ut i en excelfil.

Urval av sortförsök och sorter

I flerårssammanställningar inkluderas bara sorter som har testats det senaste året och dessutom minst ett ytterligare år i perioden.

Före analys av bestånds- och sjukdomsgraderingar exkluderas vanligen sortförsök som inte uppvisat tillräcklig variation mellan sorterna. Detta urval kallar vi *gradval*. Den undre gränsen är vanligen 5 %. Den övre gränsen är vanligen 95 % för beståndsgraderingar och 50 % för sjukdomsgraderingar. Försök som innehåller antingen bara graderingar som är lägre än den undre gränsen, eller bara graderingar som är högre än den övre gränsen, anses ointressanta och exkluderas därför före den statistiska analysen. Resultatblanketten sammanfattar därmed bara försök som behållits efter att detta urval gjorts. I ekologiska försöksserier används dock inte gradval, utan samtliga försök är intressanta och inkluderas i sammanställningen.

Odlingsegenskaper och sjukdomar graderas ofta vid flera tillfällen i samma försök. När en responsvariabel har observerats vid flera tidpunkter i ett försök används den

tidpunkt som uppvisat störst varians mellan leden. Därmed används tidpunkten då skillnaderna mellan leden varit som störst i försöket. Det här urvalet av tidpunkter görs inte bara för graderingar, utan för alla responsvariabler som observerats vid flera tillfällen i sortförsöken. De utvalda tidpunkterna kan variera mellan försöken.

Urval av ogräsförsök

I ogräsförsöken görs effektgraderingar vid upprepade tidpunkter. Om antingen antalet plantor eller tätheten varit alltför liten i det obehandlade ledet analyseras inte effektgraderingarna. Reglerna är närmare bestämt följande. Om det senast registrerade antalet plantor per kvadratmeter har varit mindre än 5 i det obehandlade ledet, så exkluderas försöket från sammanställningen av ogräset ifråga. Likaså exkluderas försöket om den senaste registrerade tätheten av ogräset har varit mindre än 1 % i det obehandlade ledet. Kraven avser ledvisa observationer eller medelvärden. Om uppgifterna om antalet plantor och tätheten i det obehandlade ledet helt saknas förutsätts ogräset ha förekommit i tillräcklig mängd, varvid försöket behålls och inkluderas i redovisningen av effektgraderingen på resultatblanketten.

Statistisk analys

Linjära blandade modeller

De ledvisa observationerna och medelvärdena analyseras i Sas med proceduren *mixed*. Denna procedur anpassar *linjära blandade modeller*. På engelska kallas dessa *linear mixed-effects models*, eller kortfattat *mixed models*. I en linjär blandad modell är effekterna av de förklarande faktorerna antingen *fixa* eller *slumpmässiga*. Fixa effekter är okända konstanter, som måste skattas. Slumpmässiga effekter är normalfördelade med väntevärde noll och okänd varians, som måste skattas. Den specifika modellen kan beskrivas genom att man anger vilka faktorer som är fixa och vilka som är slumpmässiga. Observationerna, dvs. responsvariablerna, antas alltid vara normalfördelade. Proceduren *mixed* anpassar modellen genom att skatta de okända fixa effekterna och varianserna. En iterativ numerisk algoritm används för att anpassa modellen. Ibland händer det att analysen inte konvergerar, vanligen om antalet observationer är litet och modellen är komplex, eller om variationen i observationerna är liten. I så fall saknas resultat, för den responsvariabeln, på resultatblanketten.

Statistiska modeller för sortförsök

Vid analys av ettåriga serier av sortförsök innehåller den statistiska modellen fixa effekter av sorter och slumpmässiga effekter av försök. Vid analys av fleråriga serier används i första hand en modell med fixa effekter av sorter och slumpmässiga effekter av år, samspel mellan sorter och år, samt försök. Om inte den modellen konvergerar används en enklare modell. Den enklare modellen innehåller fixa effekter av sorter och slumpmässiga effekter av försök. Även i vallförsök används dessa två modeller, där år avser det år försöket skördats. I norrländska

vallförsök förkommer, för analys av torrsubstansskörd, mer komplicerade modeller. Dessa innehåller dessutom slumpmässiga effekter av vallår och flera samspel med vallår.

Statistiska modeller för ogräsförsök

Den statistiska modellen för analys av effektgraderingar i ettåriga serier av ogräsförsök innehåller fixa effekter av behandlingar och slumpmässiga effekter av försök. Samma modell används i ettåriga serier för analys av marktäckning strax innan första vårbehandling. För fleråriga sammanställningar av effektgradering används en modell med fixa effekter av behandlingar och slumpmässiga effekter av år, samspel mellan behandlingar och år, samt försök. För variabler av typerna ”antal plantor/m²” och ”% marktäckning” beräknas medelvärden.

Statistisk redovisning

På resultatblanketten redovisas, för varje utförd statistisk analys, i separata kolumner: *i)* antal, *ii)* medelvärde, *iii)* relativtal, och *iv)* stjärnor som anger signifikans. För sjukdomar i sortförsök redovisas dock inga relativtal. Medelvärdena är så kallade *minsta-kvadrat-medelvärden*. Dessa är skattade marginella medelvärden. När serien är helt balanserad, så att samtliga försöksled finns i samtliga försök, är dessa medelvärden lika med vanliga aritmetiska medelvärden. När försöksserien inte är balanserad har medelvärdena justerats så att leden ska kunna jämföras med varandra. I kolumnen med antal redovisas hur många försök ledet ingick i. *Relativtalet* anger hur stort ledets medelvärde är i procent av referensledets medelvärde. I sortförsöken är referensledet antingen en specifik sort, som alltså fungerar som *mätare*, eller en så kallad *syntetisk mätare*. Det senare betyder att sorternas medelvärden relaterats, inte till medelvärdet för en enda utvald sort, utan till medelvärdet av några utvalda sorter. Den syntetiska mätaren kan bestå av två, tre eller fyra utvalda sorter. I ogräsförsöken är referensledet antingen det obehandlade ledet (negativ kontroll), eller en standardbehandling (positiv kontroll). I kolumnen efter relativtalet markerar ”ref” mätarledet.

När ledet är statistiskt skilt från referensledet markeras detta med en, två eller tre stjärnor. Dessa avser test för skillnaden mellan det aktuella ledet och referensledet. I sortförsöken är den statistiska nollhypotesen att det inte finns någon skillnad mellan sorten och mätaren. Det gäller oavsett om mätaren är ett av leden eller en syntetisk mätare. I ogräsförsöken är den statistiska nollhypotesen att det inte finns någon skillnad mellan ogräsbehandlingen och kontrollen. Testen är inte justerade för multipla jämförelser. Stjärnorna anger hur stort *sannolikhetsvärdet*, även kallat *p-värdet*, är, enligt följande tabell:

Sannolikhetsvärde (p)	Redovisning på resultatblanketten
$0,05 \leq p$	
$0,01 \leq p < 0,05$	*
$0,001 \leq p < 0,01$	**
$p < 0,001$	***

Sannolikhetsvärdet är sannolikheten att av en slump få en så stor skillnad som man har fått, mellan ledets och referensledets medelvärden, eller en ännu större skillnad, om nollhypotesen är sann. Vid en redovisad stjärna är skillnaden signifikant på 5 % -nivån, vid två stjärnor är skillnaden signifikant på 1 % -nivån, och vid tre stjärnor är skillnaden signifikant på 0,1 % -nivån. Om skillnaden inte är signifikant redovisas ingen stjärna.

Stjärnor redovisas inte heller om sannolikhetsvärdet för det övergripande *F-testet* är större än 0,05. Detta sannolikhetsvärde redovisas i nedre delen av resultatblanketten, på raden där det står *Sannolikhetsvärde*. Det övergripande testet är ett test av huruvida det alls finns några skillnader mellan leden. Den statistiska nollhypotesen är att det inte finns några skillnader mellan leden, dvs. att de alla har samma förväntade värde. Om sannolikhetsvärdet är mindre än 0,05 har nollhypotesen förkastats. Då finns det signifikanta skillnader mellan leden. När det finns signifikanta skillnader mellan leden anger stjärnorna vilka av leden som är signifikant skilda från referensledet.

På raderna i resultatblankettens nedre del redovisas även ett totalt medelvärde beräknat baserat på den statistiska modellen, samt *variationskoefficienten* (CV) i procent. Förkortningen, CV, kommer av engelskans *coefficient of variation*. Variationskoefficienten har beräknats med formeln:

$$CV(\%) = \frac{\sqrt{MSE}}{m} \cdot 100$$

där *MSE* är residualvariansen, och *m* är medelvärdet enligt modellen. Residualvariansen mäter den oförklarade variationen i modellen, dvs. den variation som inte kan förklaras av faktorerna i modellen. Medelvärdet är det totala medelvärde som redovisas på resultatblanketten. Vid analys av enskilda försök används ofta CV som ett mått på försöksplatsens jämnhet. Ju högre CV, desto ojämnare försöksplats, och tvärtom. I seriesammanställningarna finns däremot ingen direkt koppling mellan försökens jämnhet och storleken på CV. Däremot ger CV information om hur väl modellen beskriver observationerna. Om CV är litet förklarar modellen mycket, och tvärtom. CV kan dock vara stort inte bara på grund av hög residualvarians, utan också på grund av lågt medelvärde. Om CV är mycket stort, större än 33 %, så är det inte rimligt att observationerna är normalfördelade. Då är signifikanttesten, dvs. sannolikhetsvärdet och stjärnorna, approximativa.

På resultatblanketterna redovisas även *minsta signifikanta skillnaden*. Denna kallas även LSD, efter engelskans *least significant difference*. LSD kan användas för att jämföra alla par av led med varandra. Om skillnaden mellan två medelvärden är större än LSD är skillnaden signifikant på 5%-nivån. Detta är ett mycket praktiskt mått, men tyvärr fungerar det bara om serien är balanserad, dvs. när alla led inkluderats i alla försök. När serien inte är balanserad finns det inte ett enda LSD som duger till att jämföra alla par av led. På resultatblanketten redovisas därför alltid ett minsta LSD och ett största. När försöket är balanserat är dessa lika stora, men annars är de oftast inte det. Om försöket är obalanserat ska minsta och största LSD tolkas så här:

- Om skillnaden mellan medelvärdena för två led är större än det största LSD är skillnaden signifikant.
- Om skillnaden är mindre än det minsta LSD är skillnaden inte signifikant.
- Om skillnaden är större än det minsta LSD men mindre än det största LSD, så vet vi inte om skillnaden är signifikant.

LSD kan användas för jämförelser med mätaren om mätaren är ett av leden, men inte om mätaren är syntetisk. Stjärnorna, däremot, fungerar för jämförelser med mätaren oavsett om mätaren är ett av leden eller en syntetisk mätare.

Minsta och största LSD redovisas bara om det övergripande F-testet är signifikant, dvs. om sannolikhetsvärdet är mindre än 0,05. Denna praxis, att bara använda LSD om det övergripande testet är signifikant, kallas *Fisher's protected LSD*.