

АЛЕКСАНДЪР ЕФРЕМОВ

Идентификация на многомерни системи

Линеен подход за оценяване

Монография

2013

Рецензенти: проф. д.т.н. Никола Маджаров
проф. д-р Емил Гарипов

Ефремов, Александър
Идентификация на многомерни системи. Линеен подход за оценяване.
Монография

© 2013 Александър Ефремов, автор
© Александър Ефремов, художник на корицата



Технически университет – София



Факултет Автоматика

URL: <http://anp.tu-sofia.bg/aefremov/index.htm>
e-mail: aefremov@gmail.com

ISBN 978-954-9489-34-7

© Издателство Д А Р – Р Х

DAR. PH

Велико Търново,
ул. Цар Самуил, № 7, етаж 1
тел.: 062/63 99 14
URL: <http://www.darvt.com>
e-mail: darhristova@abv.bg

Александър Ефремов

Идентификация на МНОГОМЕРНИ системи

Линеен подход
за оценяване

DAP-PX

2013

Трудът може да послужи на студенти, докторанти и специалисти, занимаващи се с изследване, разработване и използване на алгоритми за извличане на информация от данни.

Предполага се, че читателят има известни познания по линейна алгебра, математически анализ, статистика и динамични системи. Познания по идентификация на едномерни системи, числени методи, а също иконометрия, финанси и др. биха допринесли за по-бързото навлизане в материята.

Елементи от теорията, засегната в монографията, се преподава от автора по дисциплините: „Моделиране и оптимизация на процеси“ и „Многосвързани системи за автоматизация“ във факултет „Автоматика“ и „Идентификация на многомерни системи“, част I и II във Факултет „Приложна математика и информатика“ на ТУ – София.

БЛАГОДАРНОСТ

Благодаря проф. Никола Маджаров и проф. Емил Гарипов от ТУ – София за знанията, които са основа на работата ми в индустрията, както и за полезните съвети при рецензирането на ръкописа.

Благодаря на доц. Асен Тодоров от ТУ – София за безценната подкрепа през годините.

Благодаря на Христо Хаджичонев от „Експириън“ за средата, която ми осигури – без нея този труд нямаше да го има.

Благодаря на Даря за разбирането.

АЛЕКСАНДЪР ЕФРЕМОВ е гл. ас. д-р в ТУ – София, факултет „Автоматика“, катедра „Автоматизация на непрекъснатите производства“. Завършил е специализация „Системи и управление“. Работил е в петролната индустрия (Gallus Ltd.), спътниковата комуникация (SkyGate BG, cera RaySat Inc.), пазарната индустрия (Retail-Analytics Ltd.), а в момента в сферата на финансите (Experian Ltd.). Занимава се с разработване на алгоритми за извличане на информация от данни и оптимизация на стратегии.

ИДЕНТИФИКАЦИЯ НА МНОГОМЕРНИ СИСТЕМИ

В днешно време човечеството разполага с огромно количество данни. С нарастване на възможностите за извличане на информация от данните, става решаема задачата за описание на системи с все по-голяма размерност. Често предварителната информация за такива системи е недостатъчна за построяване на модел, а понякога липсва. Това са основни предпоставки за развитието на идентификацията на многомерни системи.

Трудът засяга етапите на идентификацията, като акцентът е изцяло върху многомерните системи. Изложени са техники, приложими както в техническата област, така и в икономиката, финансите, медицината, социологията и др. Теорията е изложена основно за динамични системи, но са представени и дейности, типични за изграждането на статични модели. Засегнати са характерните особености и проблеми при изграждането на многомерни модели. Представени са техни решения, а също и числено устойчиви реализации на методите за оценяване на параметри, както и ефективни решения за големи набори от данни.

В монографията са обобщени представянията на многомерните регресионни модели, чиито параметри може да се отделят (поякога условно) от факторите. Това са представяния с матрица и с вектор на параметрите. Те не са еквивалентни от гледна точка на задачата на идентификацията. Затова са описани съответните реализации на методи за оценяване на параметри, различията при избора на структура на модела, както и практическата приложимост на представянията с матрица и с вектор на параметрите.

В отделен труд предстои излагането на задачата за оценяване на параметри и избор на структура на нелинейни модели, които не може или не е удачно да се представят в линейно параметризиран вид. Също така, ще бъдат разгледани рекурсивни методи за оценяване на параметри както за линейно параметризирани, така и за нелинейни по параметри ММО модели.

А. Ефремов
София, 2013

Съдържание

Увод	1
1 Моделиране на многомерни системи	5
1.1 Въведение	5
1.1.1 Извличане на информация от данните	7
1.1.2 Приложение на моделите	8
1.1.3 Примери за многомерни системи	16
1.1.4 Неопределеност	24
1.1.5 Подходи за моделиране	26
1.2 Аналитично моделиране	30
1.2.1 Граници на системата	31
1.2.2 Първоначален модел	31
1.2.3 Преобразуване от частни в обикновени ДУ	32
1.2.4 Линеаризация	32
1.2.5 Намалвяване на реда на модела	33
1.2.6 Дискретизация	33
1.2.7 Връзка с експерименталния подход	33
1.3 Експериментално моделиране	35
1.3.1 Граници на системата	36
1.3.2 Клас на модела	37
1.3.3 Експеримент, предварителна обработка и анализ	37
1.3.4 Оценяване на параметри и избор на структура	38
1.3.5 Валидация на модела	39
1.3.6 Връзка с аналитичния подход	39
1.4 Автоматизирана идентификация	41
2 Етапи на идентификацията на ММО системи	43
2.1 Клас на модела	43
2.1.1 Избор на клас на модела	44
2.1.2 Класификации на моделите	44
2.1.3 Регресионни модели	51
2.1.4 Декомпозиции на ММО моделите	75

2.2	Експеримент, предварителна обработка и анализ на данни	79
2.2.1	Планиране на експеримент	80
2.2.2	Предварителна обработка на данни	88
2.2.3	Анализ на данни	121
2.3	Оценяване на параметри и избор на структура на модела	132
2.3.1	Оценяване на параметри	133
2.3.2	Целева функция	133
2.3.3	Критерий за оптималност	136
2.3.4	Методи за оценяване на параметри	136
2.3.5	Еднократно и итеративно оценяване	155
2.3.6	Избор на структурата на модела	156
2.3.7	Поетапно моделиране	159
2.4	Валидация на модела	161
2.4.1	Разпределение на остатъка	162
2.4.2	Модел със свободен член	162
2.4.3	Обобщени показатели	163
2.4.4	Частични показатели на качеството	171
2.4.5	Преоразмеряване и универсалност на модела	174
2.4.6	Модифицирани показатели на качеството	175
2.4.7	Кросвалидация	178
3	Оценяване на параметри с линеен подход. Структура на модела	183
3.1	Параметри и структура на модела	183
3.1.1	Оценяване на линейно параметризирани модели	185
3.1.2	Оценяване на един ММО или множество МISO модели	186
3.2	Методи за модел с матрица на параметрите	188
3.2.1	Модел в общ вид с матрица на параметрите	189
3.2.2	Метод на най-малките квадрати	191
3.2.3	Метод на претеглените най-малки квадрати	206
3.2.4	Метод на обобщените най-малки квадрати	214
3.2.5	Метод на разширените най-малки квадрати	215
3.2.6	Метод на инструменталните променливи	224
3.2.7	Робастен метод на най-малките квадрати	227
3.3	Методи за модел с вектор на параметрите	233
3.3.1	Модел в общ вид с вектор на параметрите	234
3.3.2	Метод на най-малките квадрати	237
3.3.3	Метод на претеглените най-малки квадрати	239
3.3.4	Метод на обобщените най-малки квадрати	241

3.3.5	Метод на разширените най-малки квадрати	241
3.3.6	Метод на инструменталните променливи	246
3.3.7	Робастен метод на най-малките квадрати	247
3.4	Числено устойчиви реализации на оценителите	251
3.4.1	Числени проблеми при оценяването на параметри	252
3.4.2	Линейно зависими фактори	254
3.4.3	LS с регуляризация на Тихонов	257
3.4.4	LS с QR декомпозиция	260
3.4.5	LS с SVD декомпозиция	263
3.5	Избор на структурата на модела	268
3.5.1	Оптимизиране на структурата	269
3.5.2	Структурни параметри на многомерни модели	270
3.5.3	Стъпкови методи	272
3.5.4	Регресия с пълно множество от модели	276
3.5.5	Права стъпкова регресия	277
3.5.6	Обратна стъпкова регресия	282
3.5.7	Комбинирана стъпкова регресия	285
3.5.8	Групова стъпкова регресия	291
Заклучение		293
Приложение. Умножения на тензор и матрица		297
Означения		303
Съкращения		315
Библиография		316

Увод

Моделирането е познавателен процес, чиято цел е изучаването на изследваните обекти. Това се постига с изграждането на модели – приближения, отразяващи определени аспекти на обектите. Моделите са фундаментални за човешкото познание. От раждането си човек се учи да интерпретира потока от сетивни данни от околния свят, като с времето формира все по-сложни и по-прецизни представи за света. Така той изгражда своята познавателна система [63], с която осъзнава света. Например, ако човек попадне в стая, в която не е влизал, с бърз поглед той мигновено „подрежда“ конкретната реалност. Не е нужно да се вира във всеки предмет – неговата познавателна система услужливо попълва всички липси в потока от сетивни данни и така се получава една завършена реалност (дори не е нужно постоянно да се уверява, че стената зад гърба му е все още там – благодарение на представите, в съзнанието му стената е точно на мястото си...). В този смисъл, за да може човек да поддържа своята ежедневна реалност, се нуждае от възприятията дотолкова, че да захрани познавателната си система с данни, а същинското осъзнаване се получава с помощта на извикваните представи. Тъй като осъзнаваме света, като боравим с представи, може да се каже, че изучаването му е процес на моделиране, а продуктът от тази дейност са представите – модели, опростени заместители на различни аспекти от околния свят.

Може би поради толкова важното значение на моделите за човека, те намират приложение в много области от неговата дейност – и в науката, и в изкуството. Областта, засегната в монографията, е изграждането на математически модели, като по-точно е изложен единият от двата подхода за моделиране, наречен идентификация. Този подход е експериментален. При него, на базата на достъпни данни за изследвания обект, се формира математическо описание на тези негови аспекти, които данните отразяват. Подходът рядко намира приложение в чист вид, тъй като данните съдържат, освен информация за обекта, и неопределености. Ако сяко се разчита на данните, е възможно да се получи модел, който не отразява достоверно описваните аспекти на обекта. Затова винаги има стремеж идентификацията да се комбинира с другия

подход – аналитичното моделиране. Докато при идентификацията източникът на информация за обекта са данните, при аналитичния подход информацията се извлича от предварителните знания за обекта (основните принципи и закони от съответната област). Този подход също не е желателно да се прилага самостоятелно, защото предварителното знание също е описание на реалността, а това е предпоставка за получаване на неточен модел. Тъй като допусканията грешки при използването на двата източника на информация имат различна природа, комбинирането на подходите за моделиране позволява тези грешки да се потиснат. Затова, въпреки че трудът засяга дейностите, свързани с идентификацията, ще бъде отделено внимание и на използването на предварителна информация, характерна за аналитичния подход.

В Първа глава е засегнат въпросът за извличане на информация от данните, като обикновено, в определен етап на този процес, се построява модел. Представени са различни цели, за които се използват моделите, както и някои примери на обекти, за които изложената теория намира приложение. В тази част на монографията се излагат накратко етапите на аналитичното и експерименталното моделиране. Първа глава завършва с въпроса за нарастващата нужда от автоматизиране на идентификацията. Когато изследваната система е с много голяма размерност (в смисъл на брой входове и/или изходи), автоматизираното изграждане на модел може да се окаже единственият път за изпълнение на идентификацията.

Втора глава подробно засяга етапите на идентификацията. Разглеждани са различни класификации на моделите, най-вече с цел уточняване на класовете модели, които се разглеждат по-подробно в монографията. Обобщени са възможните представяния в общ вид на многомерните регресионни модели, при които параметрите може да се отделят от регресорите. Те се разделят на две групи, а именно – представяния с матрица и с вектор на параметрите. Подробно са изследвани техните свойства. Направени са заключения за тяхното приложение както в рамките на процеса на идентификация на многомерни системи, така и от гледна точка на някои изисквания към моделите, налагани в различни приложни области. Подробно е засегната темата за събирането, обработката и първоначалния анализ на данните – етап, който може да отнеме основната част от времето за изграждане на модел и който може да е решаващ за изхода от идентификацията. Представен е и същинският етап по изграждането на модела, като тук целта е обзор на методите за оценяване. Разглеждат се оценители на линейни и нелинейни модели,

засягани в монографията, а също и изграждане на оценители за произволни, но известни разпределения на остатъка. Освен това е засегнат и въпросът за отчитането на предварителни знания за стойностите на параметрите. По този начин е извършен преглед на методите за оценяване на параметри, в основата на които е методът на Бейс. Общо е описана дейността по избора на структурата на модела. Последната тема във Втора глава се отнася за проверката на достоверността на модела. Често в литературата този етап се нарича накратко валидация. Това е установен термин, който се използва в монографията, заедно с по-точното и коректно наименование: проверка на достоверността на модела, въведено в българската литература в [10]. Изложени са обобщени и частични показатели на качеството. Разгледан е и въпросът за евентуалното преоразмеряване на описанието и са описани начини за избягване на този нежелан ефект. Накрая са представени еднократната и многократната кросвалидация.

В последната глава подробно е изложено оценяването на параметрите и изборът на структурата на многомерни модели. Изведени са конкретни варианти на двете общи описания на моделите. За тези варианти са изведени и конкретни оценители на параметрите. В основата на оценителите, разгледани в труда, е заложено, че моделите или са линейни по отношение на параметрите, или се приемат за такива на етапа на оценяването на параметрите им. Също така са представени структурните параметри на многомерните модели, както и начините за тяхното определяне. Подробно са описани идеите на различни стъпкови методи за избор на множество от значими (като група) фактори и са изложени конкретни алгоритми.

Монографията завършва със заключение, в което се акцентира на приложимостта на теорията в нетехнически области, особено където моделите са с много голяма размерност. Основно внимание е отделено на приносните моменти в труда, а именно: обобщението на описанията на многомерните регресионни модели; описанието на някои оценители на параметри на многомерни модели, където представянето с тензори е удачно; автоматизирането на процеса на идентификация и построяването на многомерни динамични модели на системи с много голям брой входове и изходи.

В отделно приложение са описани две умножения на тензор и матрица, едното от които е предложено в монографията и се използва за опростяване на характерни действия между тензор и матрици, извършвани в някои от многомерните оценители.

Практическите разработки на автора са свързани с анализи, извършени за компанията „Галус“ [96] и проектирането на алгоритми, внедрени в софтуерни продукти на компаниите „Ритейл Аналитикс“ [134] и „Експириън“ [87]. По-конкретно, част от теорията отразява функционалността на софтуера за моделиране MDS на „Експириън“ и разработените от автора автоматизирани системи за изграждане на динамични модели на хипермаркети с много голяма размерност (от порядъка на милиони входове и изходи) по поръчка на „Ритейл Аналитикс“.

За представяне на теорията са дадени примери, изпълнени в средата на Matlab. Тъй като практическата работа в основата на труда е разработване на алгоритми без използването на готови функции, съответно такива не са дискутирани в монографията. В този смисъл цел на монографията е изследването на принципите в идентификацията, а не представяне на конкретни среди за моделиране, изследване на възможностите на SAS, R, приложимостта на System Identification Toolbox на Matlab, и т.н. за изграждането на многомерни модели. По тази причина стремежът в изложението е да се изясни идеята на разглежданите методи и алгоритми, както и конкретни реализации в подробности, позволяващи на читателя да се запознае практически с проектирането на алгоритми за решаване на реални задачи.

Глава 1

Моделиране на многомерни системи

1.1 Въведение

Системата най-общо е множество от взаимосвързани елементи. Основните ѝ свойства се получават от взаимодействието между елементите, а не от действието им поотделно. От гледна точка на моделирането, понятието система или обект на изследване е тази част от реалния свят, която трябва да бъде описана. В някои източници се прави разлика между обект и система, като системата включва само тази част от реалния обект, която подлежи на моделиране. Например в кредитната индустрия, обектът може да е кандидат за кредит, а системата да включва само връзката между възрастта, семейното положение, дохода и др. характеристики и вероятността кандидатът да е „добър“ кредитополучател (а други аспекти на обекта, като гръдна обиколка и облекло, не се включват в изследваната система). В монографията понятията система и обект (ако не е уточнено друго) ще се използват за означаване именно на аспектите от света, които трябва да бъдат описани. Особеност на системата е, че тя си взаимодейства с околния свят. За да се формира нейно описание, първо тя трябва да се изолира от заобикалящата я среда. Тази задача невинаги е лесна – има приложни области, където отделянето ѝ е твърде условно.

В резултат от моделирането се получава модел – приближение на споменатите аспекти на реалния обект. В текста модел е математическото описание на системата. Основни свойства на модела са крайност, опростеност, приближеност и достоверност. Под крайност се има предвид, че моделът отразява само тези свойства на обектите, които са важни за приложението му. Опростеността е свойство, което често се търси, тъй като по-лесно се борава с модел с по-проста структура. Моделът е приближение на системата, а основна цел на моделирането е неговата достоверност да е задоволителна от гледна точка на целта, за която ще се използва.

В долните разглеждания моделът отразява връзката между множество от входни величини и друго множество от изходни величини. В монографията се използва и понятието сигнал, когато се има предвид конкретно поведение на обекта. Сигналят е функция на една или няколко независими променливи [128]. Най-често той се разглежда като функция на времето, например температура в определена точка от пространството, която се изменя с времето. От друга страна, сигналят, носещ информация за червения цвят, в изображение зависи от две независими променливи, които са пространствени.

Понякога е възможно описанието на системата да бъде получено аналитично, като за целта се прилагат основните принципи в съответната област на познанието (физика, химия, биология, икономика и т.н.). Този подход за получаване на модели има предимството, че се базира на конкретните закономерности в изучаваните процеси. Но често за това са нужни специализирани и задълбочени познания, които понякога липсват. Също така аналитичното определяне на модела може да е свързано със значително време и усилия, което оскъпява и затруднява неговото формиране. Получените по този начин модели обикновено са твърде сложни. От друга страна, невинаги е възможно аналитично да се уточни класът, подходящата структура на модела, значимите фактори, влияещи на изхода му, и др. В такъв случай се прилага експерименталният подход, наречен още идентификация, в основата на който е обработката на входно-изходни данни за изследваната система.

Често терминът идентификация се свързва с експерименталното моделиране на динамични системи в техническата област [113, 139]. Но тъй като етапите при изграждането на модели в области като финанси, пазарни системи, медицина, социология и др. са същите като в техниката, в монографията идентификацията, като понятие, съвпада с експерименталното моделиране изобщо, независимо от свойствата на системата или приложната област. Обикновено този подход се комбинира с първия, като ползването на априорната информация помага за избор на класа на модела, неговата структура, за планирането, провеждането на коректни експерименти, адекватната обработка и анализ на получените данни, избора на подходящ оценител на параметри, въвеждането на ограничения при определянето на модела и т.н.

Големите възможности на съвременните информационни системи позволяват лесно да се събират, съхраняват и пренасят данни за много величини. В резултат на това в днешно време съществува огромно количество данни, което е основната предпоставка за развитието на експе-

рименталния подход. Успешното му прилагане при решаването на много практически задачи е подпомогнато и от това, че съвременните компютърни системи са много удобно средство за реализация на процеса на идентификация.

Разширяването на областта, в която експерименталното моделиране намира приложение, е свързано и с отчитане на спецификата на изследваните явления, при което възникват и допълнителни предизвикателства при неговото прилагане в практиката.

Наличието на данни за голямо количество величини позволява да се изследват много явления, които в миналото е било немислимо да бъдат изучавани, а събития, които са разглеждани като независими, сега, с помощта на този подход за моделиране, може да се интерпретират като особености на поведението на една обща система. Това позволява границите на изучаваните системи да се разширят и така да се отчетат повече аспекти от поведението на изследваните обекти. Едно от основните предизвикателства пред идентификацията е взаимната свързаност на множеството процеси, протичащи в реалните системи. В резултат на това се получават многомерни описания, които са значително по-прецизни от отделните модели, описващи конкретни подсистеми.

1.1.1 Извличане на информация от данните

В началото на индустриалната революция основният фактор, влияещ на себестойността на продуктите и услугите (а оттам и на живота на хората от гледна точка на съвременната цивилизация), са били въглищата. След това е настъпила т.нар. „ера“ на нефта, а след нея – на електричеството. След появата на компютрите и развитието на информационните системи е настъпила ерата на изчислителната мощ. В днешно време доминиращият фактор, влияещ на развитието на цивилизацията, е свързан с данните и по-точно с възможността за извличане на полезна информация от тях.

Поради натрупаното голямо количество данни в различни области на човешката дейност възниква понятието Data mining, което може да се преведе като извличане на информация от данните (ИИД). Понятията данни и информация не са еднозначни. Данните са сведения за изследвания обект и може да се разглеждат като стойности (не задължително числови) на величини. Информацията, от друга страна, е закономерната смислена съставка в данните. В този смисъл данните са носител на информация, а информацията е носител на знания, като знанието от своя страна е осъзната информация за обекта.

ИИД [43] представлява търсене на закономерности в набори от данни, като моделирането и по-конкретно идентификацията е основната част, а понякога е и единствената дейност. Съществуват редки изключения, при които извличането на информация се постига без явно моделиране, например ако се търси степента на корелация между известни сигнали, то не е нужно да се изгражда модел.

Обикновено ИИД се асоциира с боравенето с големи масиви от данни. Поради често значителната размерност на масивите, както и възможността да се конкретизират стъпките при изпълнението на идентификацията възниква стремежът моделирането да се автоматизира. Така, в днешно време, по естествен начин се стига до автоматизиране на процеса на идентификация в области като пазарни системи, финанси и др.

Има голямо припокриване между статистиката и идентификацията. Повечето методи в експерименталния подход могат да се разглеждат и като статистически. Поради това идентификацията се нарича освен експериментално и статистическо моделиране. Все пак има разлики в прилагането на методите. Традиционните статистически методи често се свързват с активното участие на изследователя, който потвърждава/отхвърля правилността на модела, а размерността на системата (брой входове и изходи) е сравнително малка. Така изследователят лесно може да анализира значимостта на величините, участващи в изследването. Човешката намеса при изпълнението на статистическите методи подпомага откриването на особени случаи, неотчетена информация, грешки в данните и т.н., но затруднява автоматизирането на моделирането. От друга страна, методите в ИИД се използват за голям обем от данни и автоматизираното изпълнение на предварително уточнената методология води до значително по-бързо достигане до крайно решение.

1.1.2 Приложение на моделите

Може да се каже, че моделът е есенцията, извлечената от данните информация за системата. Съществуват различни задачи, които се нуждаят от тази информация. Използването на модела вместо реалната система при решаването на някои задачи улеснява, а понякога е и задължителна стъпка към достигането до търсеното решение. По-долу са разгледани основните приложения на моделите.

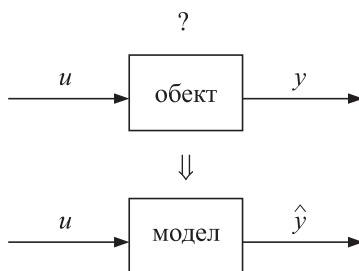
1.1.2.1 Анализ

Анализът на системата обикновено се извършва с помощта на модел. Изследването на взаимовръзките между величините може да се постиг-

не, като се изгради достоверен модел на базата на налични данни. След това фокусът се измества от системата и анализът се извършва с помощта на модела.

На Фигура 1.1 е представена идеята за използването на модел при анализа на система. С u е означен входът, с y – изходът на системата, а с $\hat{y}$ – изходът на модела. Тъй като системите, разглеждани в монографията, са многомерни, u , y и $\hat{y}$ съдържат съответно стойностите на всички входове, изходи на системата и изходи на модела.

Добре построенят модел обикновено е по-удобен за изследване от реалния обект. В много случаи не е желателно, а понякога е и невъзможно поведението на обекта да се анализира директно. Например не е допустимо да се експериментира с обект, ако той е хипермаркет, тъй като промяната на въздействията (цени, намаления и т.н.), с цел анализ на взаимовръзките между величините, може да доведе до значителни парични загуби. Други затруднения може да са: голяма продължителност на експеримента, вероятност обектът да изпадне в нежелано състояние, и др. Освен това предимство на анализа с помощта на модел е възможността да се изследва евентуалното поведение на системата при въздействия, които не са наблюдавани до момента. Това има отношение и към оптимизацията на стратегии, която е разглеждана в точка 1.1.2.6.



Фигура 1.1. Анализ

1.1.2.2 Класификация

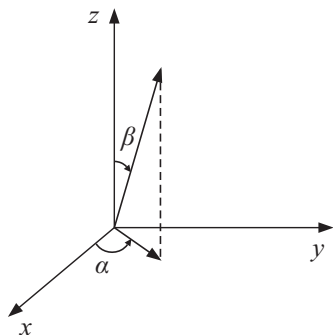
Друга форма на извличане на информация с помощта на модели е изолирането на различни режими на работа на системата. Например, ако моделът е дърво на решенията, а приложението – оценка на кредитен риск, с такъв модел банките може да идентифицират сегменти от населението, във всеки от които индивидите имат сходно поведение (в смисъл на изпълнение на задълженията по кредит), което се отличава от това в другите сегменти. Ако кредитоискател е с висок доход, с чисто шофьорско досие и притежава нов автомобил на лизинг, може да се очаква, че той би бил по-добър от средния.

От гледна точка на автоматиката това приложение може да се използва за изолиране на работни области, във всяка от които изследвана-

та система има конкретно (различно за отделните области) поведение. Впоследствие тази информация може да се използва при прилагането на многомоделния подход за управление на обекти.

1.1.2.3 Оценяване

Оценяването на недостъпни за измерване сигнали е друго приложение на модела. Съществуват алгоритми, с които е възможно да се оценяват такива сигнали, като се отчитат наличните неопределености (видовете неопределеност са дискутирани в точка 1.1.4). Например в навигационните системи за определяне на абсолютното завъртане на тяло в пространството, изразено с азимут α и тилт β (Фигура 1.2), се използва филтър на Калман [115, 24, 103]. Сигналите α_k и β_k , където k е дискретен момент от времето, са неизмерими и за тяхното оценяване се използват данните от жирокопи, инклинометри, акселерометри и др.

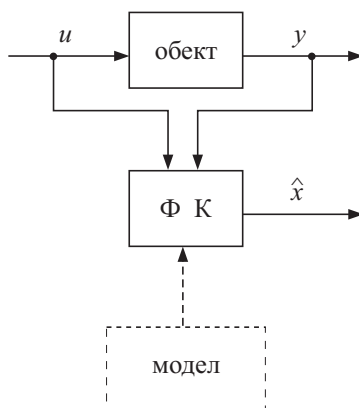


Фигура 1.2. Ъглите азимут (α) и тилт (β), описващи абсолютно завъртане в тримерното пространство

Най-общо филтърът на Калман е оценител, който при определени условия определя най-вероятното поведение на неизмерими величини (обикновено състояния и/или параметри) при наличието на неопределеност в данните и в системата. Този оптимизационен метод е наречен филтър поради това, че потиска (филтрира) влиянието на споменатите неопределености върху оценките. Първоначално той е създаден за оценяване на състоянията x_k (Фигура 1.3) на динамични системи и затова се нарича още оптимален стохастичен наблюдател на състоянието, но също се използва и за оценка на параметрите на модела [23, 104]. Почти

винаги се прилага рекурсивният му вариант, като оценените състояния може да се използват за целите на управлението, за мониторинг и др., а ако с него се оценяват параметри, той намира приложение в адаптивните системи и идентификацията. В [122, 88, 31] е засегнат въпросът за оценяване на параметри и състояния.

При работата си филтърът на Калман използва два източника на информация. Единият източник е потокът от входно-изходни данни, а другият е моделът (представата за обекта). Работата му може да се разглежда като комбинация от две дейности, наречени предиктор и ко-



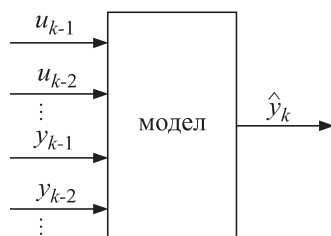
Фигура 1.3. Оценяване на състояния с филтър на Калман (ФК)

ректор. Едно от основните приложения на филтъра е за предсказване на сигнали, в т.ч. и неизмерими.

1.1.2.4 Предсказване

Предсказването на поведението на системата е важно при решаването на задачи от най-различни области като климатология, автоматика, финанси, медицина и др.

Основно в монографията се разглеждат динамични системи, а тъй като данните са дискретни стойности (в случая във времето), моделите също са дискретни. Използването на модели за прогнозиране на изхода на една динамична система предполага, че всички фактори, нужни за изчисляване на изхода, в съответния момент от време са налице. Когато предсказването се извършва за целите на управлението, изходът y в даден момент зависи от предисторията на управлявания сигнал u , но не и от стойността (стойностите) му в текущия k -ти момент, както е показано на Фигура 1.4. Причината е, че в k -тия момент първо се измерва y_k (вектор от текущите стойности на изходите), а след това, в зависимост от измерените стойности, регулаторът формира управлението u_k (отново вектор). По тази причина не е възможно y_k да зависи от u_k . В монографията, когато системите са динамични, е при-



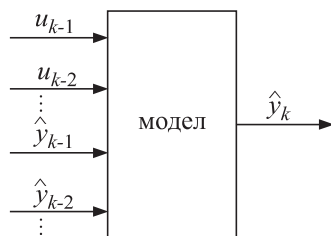
Фигура 1.4. Предсказване

ета тази постановка и затова компонентите на u_k не са фактори при прогнозирането на y_k .

На Фигура 1.4 е дадено предсказване на изхода с един такт, което се извършва с т.нар. едностъпкови екстраполатори. В някои случаи предсказването се извършва не за един, а за $h > 1$ такта, например при предсказващото управление. За целта е необходимо да се формират бъдещите стойности на управлението $u_k, u_{k+1}, \dots, u_{k+h-2}$. Възможно е така да се проектира моделът, че директно да предсказва y_{k+h-1} (предсказването с един такт означава да се формира $\hat{y}_k$ и съответно предсказване с h такта означава, че се формира $\hat{y}_{k+h-1}$). Алтернативата е на базата на предисторията да се формира предсказаният изход $\hat{y}_k$. След това, вместо все още неизвестния вектор y_k се използва $\hat{y}_k$, за да се определи $\hat{y}_{k+1}$. Така общо h пъти се извършва едностъпкова екстраполация на изхода.

1.1.2.5 Симулация

Друго приложение на модела е симулирането на поведението на обекта. Тази дейност е важна при проектирането на системи за управление, например за определяне на свойствата на тези системи преди практическата им реализация. Това е важно, особено когато е нежелателно или недопустимо да се експериментира с реалния обект. Така със симулации може да се спестят средства, да се предотвратят аварии и т.н. Освен при синтеза на системи за управление, симулациите често се използват и за оптимизация, за отчитане на повреди и др. Със симулации може да се анализират сложни алгоритми, чиито свойства е трудно или невъзможно да се изследват по аналитичен път. Пример за това е анализът на филтър на Калман, изведен за нелинеен модел (това е т.нар. разширен филтър на Калман). За разлика от линейния филтър, чиито свойства като устойчивост, сходимост и др., може да се анализират по аналитичен път, разширеният филтър, особено при сложни нелинейни модели, е удачно да се анализира с Монте-Карло симулации [33, 84]. При този подход моделът на обекта се използва за генериране на множество реализации на (симулирани) входно-изходните величини. В реализациите се въвежда неопределеност, и с помощта на статистически методи се проверяват хипотези, свързани със свойствата на филтъра.



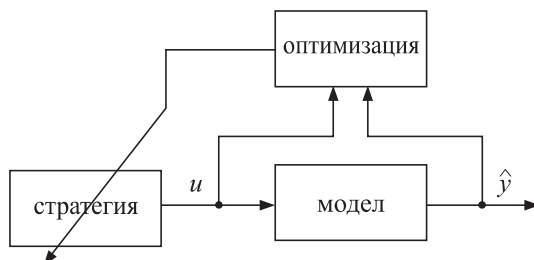
Фигура 1.5. Симулация

Основната разлика между симулацията и предсказването е, че при първата само входните сигнали са налични (при това често са генерирани, а не измерени), а вместо наблюдаваните изходи на системата участват предишни стойности на изхода на модела, както е показано на Фигура 1.5.

1.1.2.6 Оптимизация

Модел може да се използва и при синтеза на оптимални системи за управление. При изграждането им, описанието на обекта е ограничение, което трябва да се отчете по време на синтеза на системата за управление или служи като заместител на обекта, когато се предвижда поведението му при конкретни въздействия.

На Фигура 1.6 е представена примерна схема за това, как с помощта на модел може да се оптимизират стратегии. Тук моделът предоставя информация за вероятното поведение на обекта, която е нужна за определянето на правилото (стратегията), по което се формира въздействието u_k .



Фигура 1.6. Оптимизация

Например във финансите се оптимизира големината на предлаганите кредити с цел максимизиране на печалбата при различни ограничения (кредитни лимити, степен на риска [46] и др.). В пазарните системи се оптимизират цените, разположението на продуктите в магазина, рекламната дейност и промоциите, като ограниченията са свързани със складовите наличности, пространството, заемано от изложените продукти, както и допълнителни ограничения, произтичащи от бизнес логиката (съотношения между цени на сходни продукти, цени на различни разфасовки на даден продукт и т.н.).

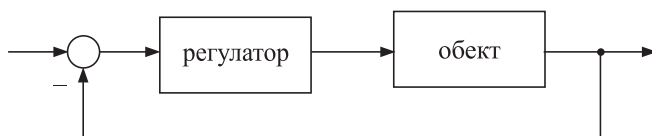
Моделите за класификация също се използват при търсенето на оптимални стратегии. Например агенция, предлагаща коли на лизинг, мо-

же да използва класификационен модел, с който да изолира потребителски сегмент, на който да предложи подходящи промоции, и по този начин да увеличи печалбата си при предварително зададен риск.

Предимства на използването на модел вместо реалната система са избягването на парични загуби, свързани с реални експерименти, както и драстичното намаляване на времето за оптимизация. Също така е възможно многократно изменение на величините, които се оптимизират, както и проиграването на недопустими за реалната система условия на работа.

1.1.2.7 Управление

В предишните точки е споменато приложението на модела при синтез на система за управление, както и това, че заместването на обекта с модел на този етап може да спести средства, време, да се избегнат злополуки и т.н.

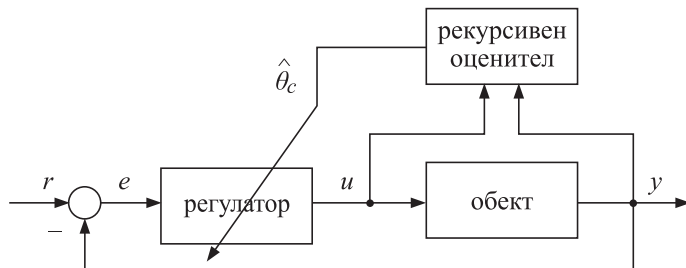


Фигура 1.7. Управление

На Фигура 1.7 е дадена класическа схема на система за управление. Почти винаги за извеждането на регулатор е необходим модел на обекта. В много случаи моделът е линеен и отразява поведението на системата в конкретна работна област. По този начин за номиналния режим на работа се определя съответен линеен регулатор. Примери за така формирани регулатори са регулаторът по зададени полюси и линейният квадратичен регулатор. Задачата се усложнява, когато нелинейността в обекта не може да се пренебрегне. Едно решение е използването на многомоделно управление, като се формират множество линейни модели, описващи обекта в конкретни работни области. Друг вариант е директно да се търси нелинеен регулатор, като за неговия синтез се използва нелинеен модел. Пример за това е нелинейният квадратичен регулатор. Причината да се избягва работата с нелинейни модели, е, че в общия случай няма унифицирано решение за регулатора, както и за анализа на поведението на системата за управление.

Въпреки трудностите при боравенето с нелинейни модели за целите на управлението такива често се използват за предсказване на сигнали с цел формиране на оптимални въздействия в области като финанси, социология и медицина.

1.1.2.8 Адаптация

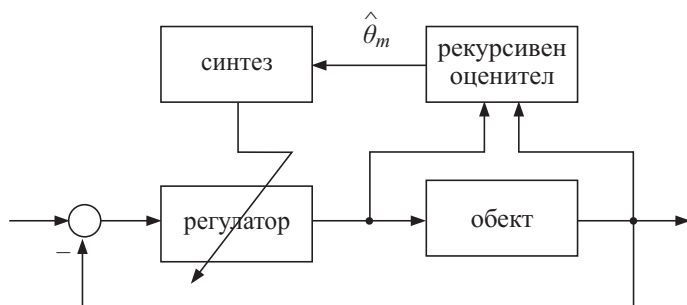


Фигура 1.8. Неявен самонастройващ се регулатор

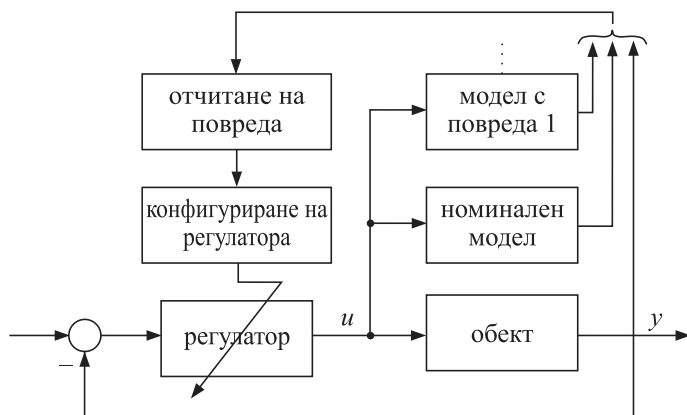
Адаптивните системи се изграждат за случаите, когато поведението на обекта се изменя в големи граници и когато други подходи, например робастни системи [28, 27], не осигуряват желано поведение на системата за управление. Друго приложение на адаптивните системи е, когато началната неопределеност в модела е недопустимо голяма и е необходимо описанието да се доуточни. Този случай може да възникне, когато първоначалният модел е получен по аналитичен път и с цел опростяване на описанието са извършени допускания, с които се пренебрегват някои от свойствата на обекта. В споменатите случаи има нужда моделът да се актуализира в реално време.

Възможно е директно да се оценяват параметрите на регулатора, например при адаптивните системи с еталонен модел и неявните самонастройващи се регулатори (Фигура 1.8). Въпреки това често при изграждането на адаптивни системи се оценяват параметрите на модела. На Фигура 1.9 е представен такъв случай.

Системите с понижена чувствителност към повреди са широк клас адаптивни системи. За изграждането им рекурсивно се оценяват параметри, състояния или други величини. В някои случаи се използва представяне на системата с множество от линейни модели [49, 53, 54] (Фигура 1.10). При този вид системи за управление стремежът е да се отчете физиката на повредите и те да бъдат включени в модела по естествения за тях начин: къде влияят (в датчиците, регулиращите органи или



Фигура 1.9. Явен самонастройващ се регулатор

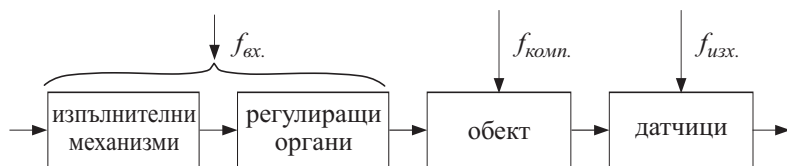


Фигура 1.10. Вариант за отчитане на повреди

по други канали, както е показано на Фигура 1.11) и как влияят (адитивно, мултипликативно или по друг начин). За тази цел се изгражда разширен модел, в който нестационарността е представена с величини, отразяващи вероятните повреди.

1.1.3 Примери за многомерни системи

Системите, разглеждани в монографията, са изцяло многомерни (с много входове и изходи). До момента бяха давани примери за системи от пазарния и финансовия сектор. По-долу тези и други системи са описани от гледна точка на многомерността. Целта тук не е задълбоченото им



Фигура 1.11. Видове повреди: входни, изходни и компонентни, както и канали на въздействие

изучаване, а по-скоро да се представят различни приложни области, в които теорията, описана в монографията, се прилага за изграждането на модел.

1.1.3.1 Технически системи

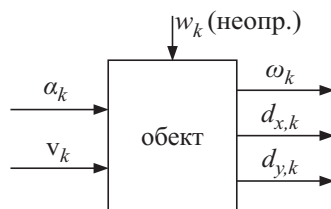


Фигура 1.12. Модули на вятърен електрогенератор (вятърна турбина)

На Фигура 1.12 е представен вятърен електрогенератор, който преобразува вятърната енергия в електричество. Най-общо има два вида турбини – хоризонтални и вертикални.

Хоризонталните турбини улавят вятърната енергия с помощта на витла, монтирани на ротор. Така енергията се преобразува във въртливо движение, което от своя страна – в електричество. Модулите на този тип обекти са дадени на Фигура 1.12, а величини, свързани с аеродинамичния модул, са показани на Фигура 1.13.

Входовете на аеродинамичния модул са ъгълът на позициониране на витлата α_k и скоростта на вятъра v_k . Изходите са ъгловата скорост



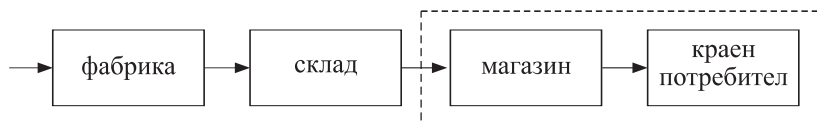
Фигура 1.13. Аеродинамичен модул на вятърна турбина

на ротора ω_k (това е основната реакция, която се управлява), както и отместването на върха на колоната по напречната $d_{x,k}$ и по надлъжната ос $d_{y,k}$ на ротора.

Освен промяната на ъгъла на витлата често се използват допълнителни въздействия, подавани от балансираща система от тежести. Тези въздействия не са включени в схемата. Отместванията по осите участват в модела, като основната им цел е избягване на повреди на съоръжението, причинени от вятъра.

Тези обекти са нелинейни, но в околност на номиналния им режим на работа може да се представят с модели като тези, разглеждани в монографията. Тъй като турбините са с повече от един вход и изход, теорията в Трета глава е подходяща за тяхното изграждане.

1.1.3.2 Пазарни системи



Фигура 1.14. Обобщена схема на пазарна система

На Фигура 1.14 е дадена примерна схема на пазарна система, която включва производството (от гледна точка на пазарното търсене), складовете, магазините и крайните потребители. За планиране на производството е необходимо да се предсказва търсенето на произвежданите стоки в зависимост от тяхната цена и характеристики, както и от цената на суровините, енергията, икономическите, социалните условия и др. По този начин могат да се предвидят оптималните количества суровини и енергия, постъпващи във фабриката, а също и времето за производство, качеството на продукцията и др.

Фабриката е многомерен обект. Поради голямата сложност на такива системи, моделирането им не се извършва само с помощта на аналитичния подход. Естествено е на определено ниво да се използва априорната информация при изграждането на модела, но като цяло експерименталното моделиране води до получаването на модел с необходимата за поставените цели точност.

По отношение на склада е необходимо да се предсказва търсенето на стоките, като се отчита предлагането (от фабриките), ограниченията на

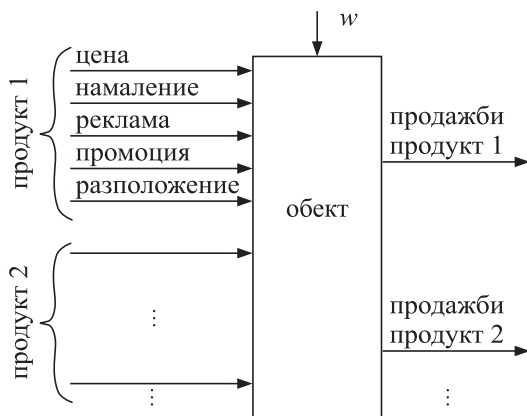
пространството в склада и т.н. Оптимизирането на количеството и разпределението на складовите наличности, откриването на оптимален път при събиране на поръчките и други дейности, свързани с управлението на склада, изискват наличието на модел.

Търсенето е основният процес, определящ поведението на търговците, които въздействат върху пазара с цените, рекламите, промоциите, намаленията, разположението на продуктите и др. За изучаване на тази част от пазарната система се използва т.нар. анализ на „потребителската кошница“ (Market Basket Analysis) [132]. От една страна, той изучава поведението на клиентите – какво купуват и как реагират на различните въздействия на търговците, а от друга, се анализират връзките между продуктите – кои продукти може да се комбинират и кои са взаимозаменяеми (конкурентни). Като резултат от анализа се извлича информация за целеви групи от клиенти, които биха купили определен продукт или комбинация от продукти. Следващата стъпка в тази насока е формирането на индивидуални въздействия.

С помощта на т.нар. последователен анализ (Sequence Analysis), се отчита изменението на сигналите във времето, като предисторията за даден клиент позволява по-точно да се предвиди неговото поведение.

Основна разлика между пазарния сектор и техническата област е, че в техниката априорната информация често е достатъчна за пълноценното прилагане на аналитичния подход, т.е., достатъчна е за построяването на модел, чиито параметри евентуално трябва да се доуточнят по експериментален път. При пазарните системи, от друга страна, предварителната информация не е достатъчна за това.

Причината за липса на задълбочени познания, на базата на които да се изведе модел, е значителната неопределеност, породена от взаимодействието на системата с околната среда. Съществуват много, и то значими фактори, за които липсват данни и следователно няма как те да се отчетат от модела. Примери за такива фактори са влиянието на конкуренцията, метеорологичните условия, някои социални, икономически събития, дори отклоняването на пътния трафик поради ремонт на пътна артерия може да повлияе осезаемо на продажбите, както и появата на конкурентен магазин в близост до изследвания търговски обект. Въпреки това априорната информация, съдържаща се в бизнес логиката, може да изиграе решаваща роля за адекватното прилагане на статистическите методи и анализа още на ниво подготовка на данните за същинското моделиране, а също може да помогне и за избора на класа, на подходящата структура на модела и за неговата валидация.



Фигура 1.15. Магазин като система за прогнозиране на търсенето

Примерите в монографията са най-вече върху последната част от обобщената схема, дадена на Фигура 1.14. Моделите, описващи този аспект от пазара (Фигура 1.15), се характеризират със значителен брой входове. Например в съвременните хипермаркети се предлагат стотици хиляди продукти и съответно толкова са изходните величини (продажбите). Входните въздействия са значително повече от изходните – има търговски вериги, които отчитат по няколко десетки въздействия за всеки продукт, например: цена, намаление, видове промоции, реклами, позициониране на продукта в магазина, брой изложени продукти и т.н.). Така задачата за идентификация се изпълнява по начина, характерен за ИИД – масивите от сурови данни са от порядъка на десетки гигабайти, особено за вериги от хипермаркети, и затова идентификацията е автоматизирана.

Често данните представляват агрегирани седмични стойности, което позволява моделите да отчитат динамиката на пазара. Пример за динамика в поведението на продажбите е ефектът от прилагането на промоция. В началото на промоционалния период продажбите нарастват, но поради ефекта от презапасяване на клиентите с продукта продажбите намаляват, като те може да спаднат под базовите продажби независимо от промоционалната цена и рекламите. Така подходите, които широко се използват в техниката (където по правило обектите са динамични), намират приложение и в областта на пазарните системи. В източниците

[48, 95, 68, 69, 99, 100, 50, 105] са описани опити в тази насока, а в [17] е представена автоматизирана обобщена методология за построяване на динамични модели на търсенето.

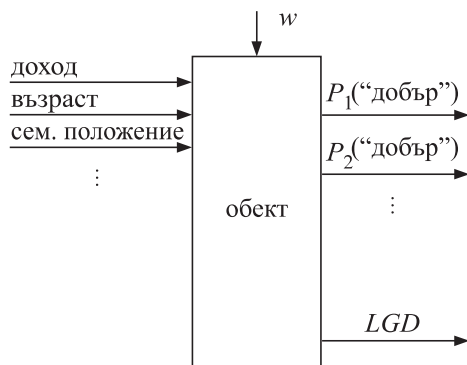
1.1.3.3 Финансови системи

Едни от важните проблеми на банките е свързан с кредитния риск. Вземането на адекватни решения в тази област до голяма степен определя коректното функциониране изобщо на финансовите институции.

При вземане на решение дали дадена банка да отпусне кредит, поведението на всеки кредитоискател се прогнозира. В миналото това се е извършвало от човек. В днешно време намесата на служител на банката в процеса на оценяване е силно ограничена. Тук моделите (най-вече линейни и логистични, както и дървета на решенията) се използват за анализ и прогноза на поведението на кредитоискателите или кредитополучателите, а също и за оптимизация на стратегиите на банките. Модели участват при оценката на риска, който банката поема по даден кредит, каква е вероятността кредитополучател да просрочи плащане и в какъв размер ще е загубата за банката, а също и какви мерки да се предприемат, ако длъжникът стане рисков.

Експерименталният подход е пътят за формиране на модел в тази област, като предпоставката е наличието на значителни масиви от данни, които банките и кредитните бюра събират. Както в пазарните системи, така и във финансите априорната информация е недостатъчна за прилагането на аналитичното моделиране (във финансите предварителните знания са още по-оскъдни). Тази информация се свежда до очаквани посоки на зависимостите между факторите и изходите на модела, предварителна информация за стратегиите на банките, реализирани по време на събирането на данните, както и спецификата на предлаганите кредити. Това би помогнало за правилното тълкуване и обработка на масивите от данни. Едно от важните приложения на априорната информация е по време на валидацията на модела. Ако моделът няма логично поведение от гледна точка на натрупания опит, то правилната стъпка е да се анализират причините. Ако те се дължат на обективни фактори, резултатът от анализа е натрупването на повече знания, но може да се окаже, че нелогичното поведение се дължи на случайни фактори и тогава е необходимо идентификацията да продължи. По тази причина моде-

ли като невронните мрежи, чиито параметри нямат физически смисъл, трудно намират приложение в кредитната индустрия. Все пак има региони в света (като Северна Америка), където финансовите институции проявяват интерес към такива модели и са по-отворени към изменения в установените методологии (за сравнение често в Англия моделите на поведението на кредитоискателите все още са линейни, въпреки че както теоретично, така и като точност за предпочитане са логистичните модели).



Фигура 1.16. Кандидат за кредит

Примерни характеристики на кредитополучателите (Фигура 1.16) са възраст, доход, образование, продължителност на стажа в настоящата работа, брой зависими членове на семейството, жилищен статус (живее под наем, в собствено жилище или с родители), бил ли е неплатежоспособен, просрочвал ли е задължения и колко пъти и т.н. Това са потенциални входове на модела. Един от изходите е вероятността кредитополучателят да е „лош“ или „добър“ клиент според определени правила на банката. Друга величина, от която банката се интересува, е загубата, която тя би претърпяла, ако клиентът ѝ изпадне в несъстоятелност (LGD – Loss Given Default). Изходите на модела може да са повече, ако с него се прогнозира поведението на кандидатите за различни продукти на банката (кредити, кредитни карти и др.) или според различни стратегии, както е показано на Фигура 1.16.

За момента данните в този сектор рядко отразяват динамиката в поведението на хората [58]. Причината е, че голяма част от характеристиките на кредитоискателите се събират само в момента на кандидатстването. Въпреки това има предпоставки за използването на динамични

модели. Например в момента няма практика да се следи за евентуална промяна в семейното положение на хората, получили кредит. Може да се очаква, че ако един клиент е семеен, но в даден момент се разведе, вероятността да се влоши поведението му, нараства. Така, ако факторите се актуализират във времето, би могло по-точно да се предвиди бъдещото изменение на степента на риска. Най-общо хора с добри други показатели, въпреки сътресения в живота като развод, може с времето отново да станат нискорискови, а при такива с нисък доход, живеещи под наем, с деца, на които ще изплащат издръжка след развод, поведението трайно да се влоши.

1.1.3.4 Примери от медицината

С модели е възможно да се оцени вероятността даден пациент да страда от определено заболяване, при известни показатели от медицински изследвания. Постановката на тази задача е много сходна с оценката на кредитния риск, където се изчислява вероятността даден човек да е добър/лош кредитоискател. Също на ниво анализ моделите се използват за откриване на зависимости между симптоми на базата на статистически данни. Впоследствие тези зависимости трябва да се обяснят или да се отхвърлят (т.е. да се приеме, че се дължат на неопределеността в данните). Резултатът от такъв анализ се използва от болнични лаборатории за допълнително назначаване на изследвания (проверка за допълнителни симптоми, с цел по-точна диагностика), както и за намаляване на разходите за реактиви.

1.1.3.5 Социални системи

В тази област неопределеността в системите е толкова голяма, че понякога не е ясно дали дадена величина е вход или е изход на системата, т.е. липсва ясна представа за причинно-следствената връзка. Например, ако се изследва връзката между вътрешната увереност и дохода, то, от една страна, неувереният човек по-трудно би убедил потенциален работодател (както и себе си), че е подходящ за добре платена позиция. В този смисъл степента на увереност влияе на дохода. От друга страна, увеличаването на дохода може да доведе до нарастване на вътрешното усещане за сигурност и увереност, т.е. доходът може да се разглежда като фактор, влияещ на усещането за увереност. На практика както доходът, така и чувството за увереност зависят от други фактори. Така по естествен път, следвайки наличните знания за социалните системи, се достига до една от основните им характеристики – многомерността.

1.1.3.6 Примери от психологията

Моделите, използвани във финансите и медицината за предсказване на вероятност, намират приложение и в психологията при извличане на информация от данни от психологически тестове. Например в криминалистиката в зависимост от отговорите на такъв тест може да се оцени вероятността за укриване на информация. Така е възможно да се изолират лицата, които евентуално дават грешни показания.

1.1.4 Неопределеност

В идентификацията от данните се извлича информация, изразена с помощта на модел, който описва системата с определена точност. Процесът на моделиране е свързан с наличието на неопределеност и за правилното му протичане е необходимо тя да бъде отчетена по подходящ начин. Най-общо, изхождайки от това, към какво се асоциира неопределеността, може да се обособят два вида: неопределеност в данните и неопределеност в системата.

1.1.4.1 Неопределеност в данните

Освен че данните са носител на информация, характерно за тях е, че съдържат и случайна съставка. Неопределеността в тях се дължи на шума от измерване на изходните величини, влиянието на околната среда, спецификата на събирането на данните и др. В някои случаи за шума от измерване има информация. Тя се предоставя от производителите на измервателната апаратура. Понякога за влиянието на смущения от околната среда върху изхода също може да се набави информация. Колкото по-запознат е изследователят с неопределеността в данните, толкова по-адекватно би могъл да обработи и анализира данните, преди да пристъпи към изграждането на модел.

1.1.4.2 Неопределеност в системата

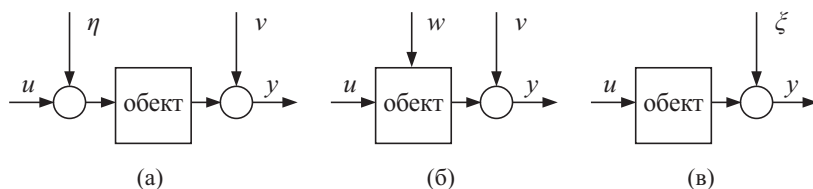
Тъй като моделът отразява с определена точност поведението на системата, това означава, че в системата има аспекти, които не се отчитат от модела. Например, за да се опише напълно една динамична система, е необходимо да се използват диференциални уравнения от безкраен ред, което е практически невъзможно. Освен това може да са нужни частни производни, нелинейни зависимости и др., които при опростяването на модела са пренебрегнати. Понякога е удачно величини, чието влияние

върху изхода е твърде сложно, да не се моделират, а да се разглеждат като смущения от околната среда. Неопределеността при тези примери не е свързана с данните, а отразява понякога изкуственото стесняване на границите на системата.

1.1.4.3 Отчитане на неопределеността в идентификацията

За да се включи неопределеността в постановката на задачата на идентификацията, се въвеждат допълнителни случайни сигнали (при динамичните системи се наричат още стохастични, тъй като са функция на времето). Често те са фиктивни, тъй като една и съща величина (смущение) може да се интерпретира като неопределеност в системата, а може да се приведе и към нейния вход или изход и да се разглежда като неопределеност в данните. Такова привеждане се извършва най-вече с цел задачата да се представи във вид, за който има готови решения. Често срещан случай е цялата неопределеност да се приведе към изхода и да се разглежда като един случаен сигнал, отразяващ цялата неопределеност.

На Фигура 1.17 са дадени различни представяния на обект и сигналите, имащи отношение към неговото поведение.

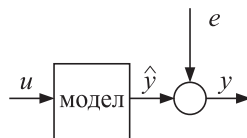


Фигура 1.17. Представяне на неопределености в системата

Както беше споменато, в монографията е прието, че сигналите са векторни. Входното въздействие u е полезен сигнал, който в автоматиката се използва за управление на обекта. В по-общ смисъл част от входовете може да са измерими величини, влияещи на изхода. На етапа на идентификацията няма да се прави разлика между управляващи въздействия и измерими смущения – тези сигнали са известни (наблюдавани) входни въздействия за обекта. Изходът y е измерим и съдържа както реакцията на системата, така и шума от измерване и смущенията от околната среда.

Сигналите η , w , ν и ξ са смущения, тъй като смущават детерминраното поведение на системата. Величината η е входно, а w е системно смущение. Обикновено в автоматиката η отразява разликата между управляващия сигнал, изчислен от регулатора, и реалното въздействие върху обекта, формирано например от регулиращ орган, както и смущения от околната среда, навлизащи по канала на управлението. С w се представя влиянието на околната среда, когато то се отчита по отделен канал, а също и неточности в модела. Изходното смущение ν отразява шума в измерването, зависещ от датчиците и смущенията от околната среда, влияещи директно на изхода. Величината ξ включва освен ν , също и външни влияния по различни канали, когато те са приведени към изхода.

Случайният сигнал e (Фигура 1.18) отчита всички стохастични влияния върху поведението на системата, при това с отчитане на конкретния модел. В този смисъл при последното представяне обектът се разглежда като сума от изхода на модела и един обобщен сигнал, съдържащ неотчетеното от модела, но наблюдавано в данните поведение на системата. Затова e се нарича обобщена остатъчна грешка или просто остатък.



Фигура 1.18. Обобщена грешка, изход на модела и на системата

Начинът, по който се отразяват неопределеностите, зависи основно от априорната информация за тях и от техниките, които се използват за оценяване на параметрите. Както се вижда, сигналите, с които се отчита неопределеността, са условни (невинаги отговарят на конкретни физически величини).

1.1.5 Подходи за моделиране

В моделирането се използват два типа информация за системата. Това са априорната и апостериорната информация. Априорна е предварителната информация, налична преди провеждането на експеримент. Например в модела на технологичен обект като ректификационната колона, при който с използване на енергия един материал се разделя на фракции, залягат законите за запазване на материята и на енергията, изразени чрез уравненията на материалния и енергийния баланс. При извеждането на модела се отчитат и конструктивните особености на обекта, химическите реакции и т.н. На този етап не се използват данни

от експеримент, а моделирането се извършва единствено на базата на априорна информация.

Има системи като техническите, където в много случаи предварителната информация е достатъчна за получаването на модел. Но в области като социология, финанси и т.н. априорната информация често не е достатъчна дори за добиването на груба представа за структурата на модела. В такива случаи идва на помощ апостериорната информацията, която се получава след провеждането на експеримент и се съдържа в сметите данни. Описването на взаимовръзки между величини, които не са свързани с точните науки, по правило се извършва именно с помощта на апостериорна информация.

Според типа на използваната информация съществуват два подхода за моделиране. Това са аналитичният и експерименталният подход, които се споменават в предишните теми. Много често моделирането е комбинация от двата подхода, като стремежът е това обединяване да се извършва винаги, когато е възможно. Така са се обособили три принципа за изграждане на модели. Това са принципите на бялата, сивата и черната кутия.

1.1.5.1 Принцип на бялата кутия

Принципът на бялата (прозрачната) кутия е всъщност прилагането на чисто аналитичния подход. Моделите, изградени по този начин, изцяло се определят от известните закони и принципи от областта, към която спада реалната система. Тези модели се наричат аналитични, физически или детерминирани. Предимствата на принципа на бялата кутия са:

- моделът има физически смисъл, т.е. всички параметри, както и структурата му, отразяват физически величини и известни закономерности в реалната система;
- моделът е универсален, тъй като се изхожда от общите принципи, валидни за процесите от изследвания вид;
- не е необходимо да се провежда експеримент със системата.

Макар че чисто аналитичният подход изглежда много ефективен и гарантира достоверността на модела, реално той е приложим за относително прости обекти. Недостатъците на принципа на бялата кутия са:

- необходимо е задълбочено познание за системата и средата, с която тя си взаимодейства. Това означава, че такива модели може да се получат само за добре изучени процеси.
- моделът често е сложен, а това затруднява използването му. Опростяването му, от друга страна, може да отдалечи чувствително поведението му от това на системата.

- моделът не се валидира с наблюдаване на реалното поведение на системата, а се разчита, че аналитично изведените и впоследствие опростени зависимости са достатъчно точни;
- изисква се много време за получаване на модела, особено когато изследователят тепърва трябва да изучи (аналитично) процесите;
- изграждането на модела не може да се автоматизира.

1.1.5.2 Принцип на черната кутия

Този принцип се основава на чисто експерименталното моделиране. Структурата и параметрите на моделите, изградени по този начин, изцяло се определят на базата на експериментални данни. Моделите, получени с използване на този принцип, се наричат експериментални, формални или статистически. Предимствата на принципа на черната кутия са:

- не са нужни задълбочени знания за системата и средата;
- обикновено моделите са опростени и лесни за употреба;
- моделът се валидира с данни за поведението на конкретната система;
- времето за получаване на модела е значително по-малко в сравнение с това на аналитичното моделиране;
- изграждането на модела може да се автоматизира.

Въпреки че чисто експерименталният подход изглежда да е приложим навсякъде, където има налични данни за изследваната система, реално той е подходящ, когато се описва поведението на конкретен обект в определен режим, а не се търси универсалност на модела. Недостатъците на принципа на черната кутия са:

- необходимо е провеждането на експеримент, който осигурява достатъчно информативни данни;
- възможно е наличието на „лъжливи“ връзки в модела, между наблюдаваните величини;
- ограничена достоверност на модела в рамките на условията, в които е проведен експериментът;
- понякога не е възможно интерпретирането на параметрите и структурата модела.

1.1.5.3 Принцип на сивата кутия

Много рядко в практиката се срещат модели, получени на базата на горните два принципа, приложени поотделно. Обикновено се прилага

принципът на сивата (полупрозрачната) кутия, при който двата подхода за моделиране се обединяват. Възможни са различни варианти за комбиниране на подходите, като съчетаването им зависи от спецификата на задачата.

Принципът на сивата кутия има предимствата и на бялата и черната кутия. От една страна, априорната информация внася яснота, като обяснява поведението на системата, а областта, в която моделът е адекватен, нараства. От друга страна, апостериорната информация приближава поведението на модела до това на конкретната система, и то когато тя функционира в желания режим на работа. Също така използването ѝ често води до опростяване на модела.

Таблица 1.1. Сравнение на принципите на бялата, сивата и черната кутия

	бяла кутия	сива кутия	черна кутия
източник на информация	теоретични принципи, закони	предварителна информация и данни от експеримент	експериментални данни
предимства	физически смисъл; универсалност; не изисква експеримент		опростен модел; близост до изследвания процес; не изисква познание на областта; по-кратко време
недостатъци	не е потвърден с експеримент; детайлно познание на областта; сложен модел; достъпното знание определя точността; изисква много време		липса на физически смисъл; осигуряване на информативни данни; вероятност за лъжливи взаимовръзки; данните определят точността; ограничена достоверност в рамките на експеримента
приложение	анализи; проектиране; по-прости обекти; добре изучени области		достъпни за измерване величини; сложни обекти; неизучени

Важно предимство на комбиниания подход е свързано с разработването на модели за практически задачи, където времето за намиране на достоверен модел е ограничено. Това условие може в голяма степен да предопредели цялостното изпълнение на задачата на моделиране-

то. Много често при проектите от индустрията или бизнеса, свързани с изграждане на модел, освен ограниченото време за изпълнение, няма достатъчно априорна информация. Също така, понякога не е допустимо провеждането на подходящи за моделирането експерименти. В тези случаи правилното балансиране между двата подхода може да се окаже единственият път за успешното изпълнение на проектите.

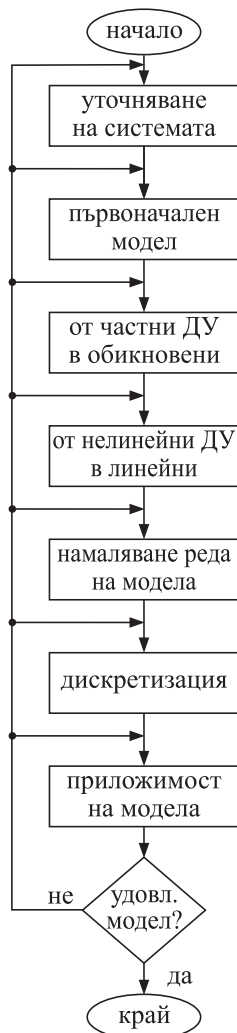
В Таблица 1.1 са обобщени особеностите на трите принципа. Непопълнените полета за сивата кутия зависят от балансирането на двата крайни принципа.

Модели, получени с преобладаващ експериментален подход, са се доказали като успешни в много приложения области като икономика, медицина, финанси, социология, психология и др., а моделите, определени главно с аналитичния подход, са широко разпространени в техниката.

Въпреки че монографията засяга главно приложението на черната кутия, важно е да се акцентира на комбинацията с методите от бялата кутия. Затова в следващата тема е засегнато аналитичното моделиране.

1.2 Аналитично моделиране

Целта на изложението не е пълно представяне на подхода, а запознаване с етапите му дотолкова, че да се определят връзките му с експерименталния подход. Познаването на етапите и допусканията, които се извър-



Фигура 1.19. Етапи на аналитичния подход

шват, позволява да се определят местата, където информацията в налични данни би подпомогнала получаването на модела.

В долните разглеждания е прието, че системите, които се изучават, са динамични. Тогава най-общо етапите на аналитичното моделиране са дадени на Фигура 1.19. Характерно за моделирането изобщо е, че то е итеративен процес, което е отразено на фигурата. Някои от етапите може да не участват при изграждането на модела, например ако от априорната информация директно се получат обикновени диференциални уравнения (ДУ), то не се преминава през третия етап. Също така последователността на прилагане на етапите може да е различна. Моделирането се прекратява, ако описанието, получено от даден етап, е приемливо за поставените цели, например, ако се търси непрекъснат модел, очевидно той не се дискретизира.

Тъй като при чисто аналитичния подход моделът не се валидира с данни за обекта, на всеки етап трябва да се следи за степента на влошаване на качеството на описанието. Този анализ значително усложнява моделирането и това е основна причина за комбинирането на аналитичния подход с експерименталния.

1.2.1 Граници на системата

На този етап изследваната система се изолира от околната среда и, ако е необходимо, се обособяват подсистеми, което улеснява извеждането на модела.

При някои задачи конкретизирането на системата е очевидно и не изисква задълбочаване в проблема, но има области, където изолирането ѝ от средата е твърде условно. Може да се каже, че на този етап се определя „прозорецът“, през който изследователят изучава реалните явления, които са от интерес. Тук се определят входовете и изходите, причинно-следствените връзки, очакваните смущения и техните статистически характеристики и т.н. Грешките, допуснати тук, са най-скъпо струващи, тъй като промяната на границите на системата може да обезсмисли много дейности, извършени междувременно в следващите етапи.

1.2.2 Първоначален модел

След уточняване на системата се пристъпва към намиране на модел (обикновено система от уравнения), който описва връзката между входовете и изходите. За целта се използват известните закони от съот-

ветната област. На този етап, с натрупване на знание за величините и уточняване на връзките между тях, е възможно някои граници на системата да се променят. Може да възникне необходимост нови величини да се включат, а други да отпаднат. При достатъчно априорна информация, обикновено в края на този етап, се получава сложен модел, в който параметрите и сигналите имат физически смисъл. Следващите дейности са насочени към опростяване на модела, а степента на опростяването му зависи от целта, за която се изгражда.

1.2.3 Преобразуване от частни в обикновени ДУ

Ако първоначалният модел съдържа частни производни, един начин за получаване на по-удобен модел е частните ДУ да се преобразуват в обикновени [13]. Този процес е наречен редукция. Примери за обекти, чиито начални модели са с частни производни, са ректификационните колони, резервоарите, изсушителите, тръбопроводите и др. Общото при тях е, че сигналите зависят, освен от времето, и от пространствени променливи. Това означава, че моделите се състоят от частни производни по съответните променливи. Процесите в тръбопровод обикновено се представят с частни производни по една променлива, описваща изменението по надлъжната ос на тръбопровода, а при изсушителите например, производните са по трите пространствени координати.

1.2.4 Линеаризация

Линейната теория намира широко приложение в практиката. Също така в много случаи нелинейните описания може да се линеаризират и получените линейни модели да се използват успешно. Затова винаги, когато е възможно, се правят опити крайният модел да е линеен.

Ако моделът съдържа нелинейни ДУ, на този етап те се заменят с линейни, като обектът се описва в околност на избран режим на работа. Това се постига, като нелинейните функции се разлагат в ред на Тейлър, в околност на работната точка. С отчитане на свободните членове и на тези от първи ред се получава линейна апроксимация на нелинейния модел. Ако областта на нормалната работа на обекта е изразено нелинейна, може да се формира набор от линейни модели, отговарящи на различни работни точки (описанието е почасти линейно), а ако линеаризацията не е удачна, се работи директно с нелинейния модел.

1.2.5 Намаляване на реда на модела

Понякога сложността на модела може допълнително да се намали, ако редът на описанията на отделните входно-изходни канали се редуцира. От гледна точка на управлението, редът на модела влияе на реда на регулатора, на наблюдателя на състоянието, на компенсиращи звена и други елементи, участващи в системата за управление. Така понижаването на реда на модела води до опростяване на структурата на други елементи от по-мощабни системи.

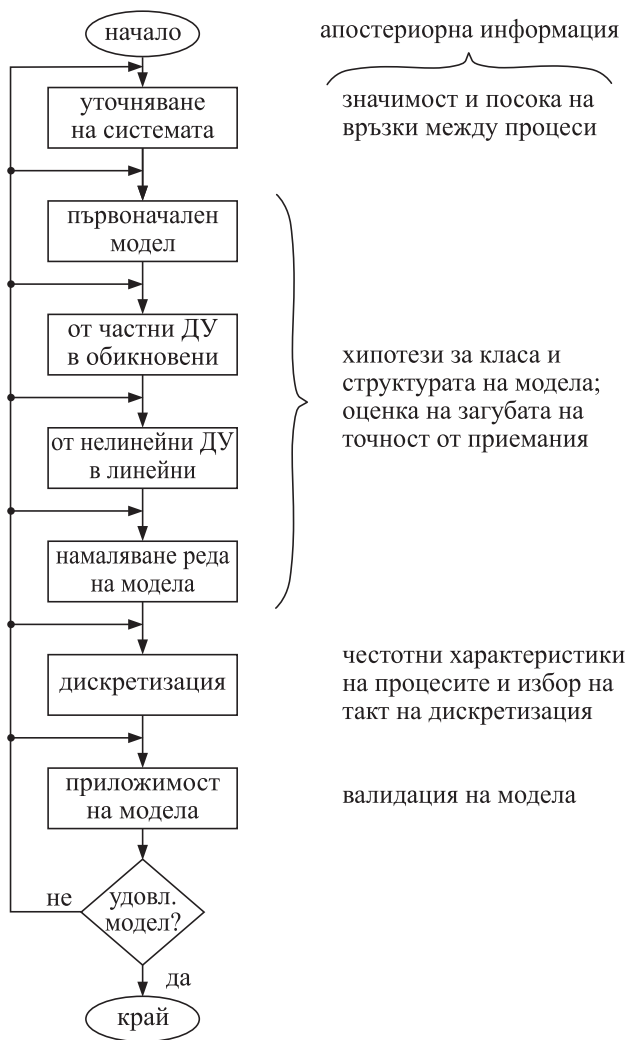
1.2.6 Дискретизация

Дискретните модели намират голямо приложение както за управление, така и за анализ, предсказване и т.н. Основна причина е, че снемането на данни се извършва в дискретни моменти от времето. Като резултат и оценителите на параметри (използвани в идентификацията) формират дискретни описания на системата. Много съвременни регулатори също имат дискретно действие, а това важи и за методите за предсказване на изхода, за оценка на състояния и т.н.

1.2.7 Връзка с експерименталния подход

На Фигура 1.20 отново са представени етапите на аналитичното моделиране, както и възможните места, където експерименталният подход би могъл да повиши достоверността на модела или да опрости неговото получаване. Може би най-значим ефект от използването на данните е валидацията на модела. На фигурата това се случва след дискретизацията на аналитичния модел. Въпреки това моделът, чиято достоверност се проверява, може да е резултат от някой от предходните етапи. Според информативността на данните те може да се използват и за отделянето на системата от средата, например чрез проверка на хипотези за това, кои величини имат отношение към системата и кои може да се изключат от границите ѝ. Също е възможно да се подпомогне определянето на структурата на модела, например откриване на силни връзки или липса на зависимости между входове и изходи, проверка на хипотези за вида на модела, избор на реда на звената, описващи отделните връзки в системата и др.

Като следствие от комбинацията на двата подхода под понятието аналитично моделиране често се включват и елементи от експерименталния подход. Обратното също е валидно. Идентификацията не е пъл-



Фигура 1.20. Аналитичен подход и възможности за комбиниране с експерименталния подход

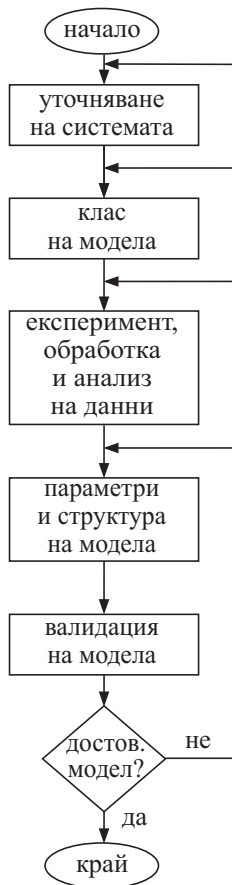
ноценна без предварителната информация за обекта (което е тема на следващата точка) и затова често при разглеждането ѝ се включват и елементи от аналитичния подход. При това положение според дейностите, които доминират, се определя за кой подход ще става дума.

По тази причина моделирането, засегнато в монографията, не е чисто експериментално – с понятията идентификация и експериментално моделиране ще се описва изграждането на модел, когато информация за системата се съдържа основно (но не задължително изцяло) във входно-изходните данни.

1.3 Експериментално моделиране

В изложението по-долу се представят накратко етапите на идентификацията на системи, като се разглежда и връзката с аналитичния подход.

Обикновено идентификацията се свързва с експерименталното моделиране в техническата област, където често системата може предварително да се дефинира. Поради това в литературата от тази област (например [10]) като първи етап се разглежда изборът на клас на модела – дейност, за която е нужно границите на системата да са уточнени. От друга страна, в много задачи извън техниката системата не е обособена в началото на изследването [14, 32]. Още повече че в процеса на идентификация може да се наложи границата ѝ да се измести с цел получаване на по-точно описание. Тъй като в монографията се разглеждат и нетехнически обекти, като първи етап се приема обо-



Фигура 1.21. Етапи на идентификацията

собирането на системата. Етапите в процеса на идентификация, който също е итеративен, са дадени на Фигура 1.21.

Някои дейности може да не участват при изграждането на модела, например ако данните са налице в началото на изследването. Това е често срещан случай при пазарните и финансовите системи, където, освен нежеланите ефекти от провеждането на активен експеримент, времето за събиране на данните продължава с години. Освен това последователността на някои етапи може да се промени. Ако предварителната информация не е достатъчна за избор на класа на модела, то трябва да се използва основният източник на информация – данните, които трябва да се съберат, обработят и анализират, а след това вниманието да се насочи към избор на клас на модела. В този смисъл, както и при аналитичното моделиране, последователността на дейностите е ориентировъчна.

1.3.1 Граници на системата

Този етап е характерен за области, където изследваната многомерна система предварително не е ясно дефинирана. Тук изследователят може да се натъкне на множество величини, които са взаимно свързани, като силата на връзките е различна. Водещата задача е същата като описаната в точка 1.2.1, но пътят за разрешаването ѝ е различен. За уточняването на системата е желателно да се използва както априорна, така и апостериорна информация, тъй като грешното определяне на системата може да обезсмисли следващите дейности.

Ако има предварителни сведения за реалните явления (априорна информация) и системата условно се отдели от средата, преди да се пристъпи към провеждането на експеримент, това може значително да опрости моделирането, както и вероятността от грешни предположения (дължащи се на статистически изводи) за структурата на системата като цяло, а оттам и да спести много усилия и време.

Когато областта не е добре изследвана, източник на информация са данни от експеримент. Тук идентификацията до голяма степен съвпада с извличането на информация от данни. Правилото е първоначално да се използва набор от данни за голям брой величини. Например при идентификацията на финансови системи величините са стотици или хиляди, а в пазарните системи, каквито са хипермаркетите, потенциалните входно-изходни сигнали може да са десетки милиони. От тях трябва да се изолира подмножество, което да участва в следващите етапи на моделирането.

1.3.2 Клас на модела

Предишният етап обикновено осигурява груба представа за системата и средата. Следващата стъпка е акцентът да се измести към търсения модел. Целта на този етап е да се конкретизират особеностите на връзките между сигналите, например дали ще се търси динамика в поведението, възможно ли е описанието да е линейно, или се налага използването на нелинейни модели, дали може системата да се приеме за стационарна и т.н. Резултатът от този етап определя като цяло стратегията на по-нататъшните действия, като какъв експеримент ще се проведе, каква обработка на данните е подходяща, какви методи за оценка на параметрите и структурата на модела ще се прилагат, както и как ще се валидира моделът. Естествено, ако в процеса на уточняване на модела се получи аналитично описание, което е удовлетворително като сложност и достоверност, то моделирането се прекратява, а подходът очевидно е аналитичен. Обикновено, ако се стигне до модел на този етап, той е сложен и е необходимо опростяване. В много случаи е нужно да се доуточнят структурата и параметрите му, а в области, където априорната информация е оскъдна, е възможно единствено да се вземе решение за най-общите свойства на бъдещия модел (класа му).

1.3.3 Експеримент, предварителна обработка и анализ

Ако няма налични данни за системата, в началото на този етап се планира и провежда експеримент, а след това данните се обработват и анализират. Целта е да се формира набор от данни, подходящ за следващия етап – изграждане на конкретен модел. Смисълът на експеримента е да се осигурят колкото е възможно по-информативни данни (виж точка 2.2.1). Експериментите може да са активни и пасивни. Предимството на активните е, че с допълнително въздействие върху системата може да се осигурят подходящи данни за идентификацията, докато при пасивните поведението на системата не се смущава целенасочено и данните може да не са информативни. От друга страна, в много случаи изкуственото смущаване на системата не е желателно и тогава единственият изход е да се проведе пасивен експеримент. Също така, типично за някои области е данните да са натрупани предварително. В кредитната индустрия например експериментът, който е пасивен, продължава с години. Тогава първата дейност, с която се стартира, е обработката на наличните до момента данни.

Понякога не се отделя нужното внимание на този етап, тъй като резултатът от него все още не е модел. Но често, особено при многомерните системи, и то с голям брой входове и изходи, този етап е решаващ за качеството на крайния модел. Като пример при моделирането на поведението на клиентите на банка времето, което се отделя за подготовката на данните, може да надхвърли 80% от цялото време за изграждане на модела. Важно е данните да се приведат в удобен вид за цифровата обработка в следващия етап (например нечислови величини да се кодират с числови значения, да се попълнят липсващите стойности, особени случаи да се кодират, да се обработят нетипичните стойности и др.). Също така е удачно още на този етап да се избегнат евентуални числени проблеми, които са характерни за идентификацията на многомерни системи.

Освен обработка тук се извършва и предварителен анализ на данните. Той включва статистически, честотен, корелационен анализ; изследване на потенциалната значимост на факторите (поотделно и на групи, при описанието на изходите); анализиране на нелинейни трансформации на сигналите за получаване на линейно параметризиран модел и др. Дори може да се стигне и до грубо моделиране при попълването на липсващи стойности. Тук може данните да се сегментират (например, ако отразяват различно поведение на системата) и за всеки сегмент да се търси отделен модел.

1.3.4 Оценяване на параметри и избор на структура

Класът на модела е необходим както за определяне на оценителя на параметри, така и за метода, с който се избира структурата на модела. Ако описанието може да се представи в линеен (по параметри) вид, то се използват методите, разгледани в настоящия труд. Когато параметрите на модела са предварително определени, то с метода на Бейс е възможно тази априорна информация да се обедини с апостериорната (съдържаща се в данните). Обикновено оценителите са директно приложими, без да са необходими допълнителни извеждания, специфични за избраното описание, за разлика от нелинейните модели, където оценителите често изискват извеждане на допълнителни величини.

За линейния случай структурата на модела се уточнява с добре отработени, автоматизирани методи, докато изборът на структурата на нелинейните модели се автоматизира по-трудно. При многомерните системи, освен структурните параметри, характерни за едномерните описания, например степени на полиноми и чисти закъснения, се проверява

и значимостта на входно-изходните величини. Така в резултат от този етап се получава модел с евентуално намален брой, но значими входове и изходи.

1.3.5 Валидация на модела

Полученият модел се проверява дали изпълнява изискванията за степента на достоверност, която е удовлетворителна от гледна точка на приложението му. Ако резултатът не е задоволителен, процесът на моделиране продължава от някой предишен етап. Желателно е да се започне от проверки на дейностите от предходния етап, като евентуално някои стъпки да се преповторят при различни условия, например промяна на структурата и евентуално на оценителя. Ако промяната не подобри достатъчно модела, то е необходимо да се продължи от предишен етап, а именно отново да се преосмисли обработката на данните. В най-лошия случай границите на системата се изменят (с добавянето на нови величини, за по-добро описание на обекта). Този итеративен процес продължава, докато моделът не удовлетвори изискванията, наложени от приложението, за което се създава.

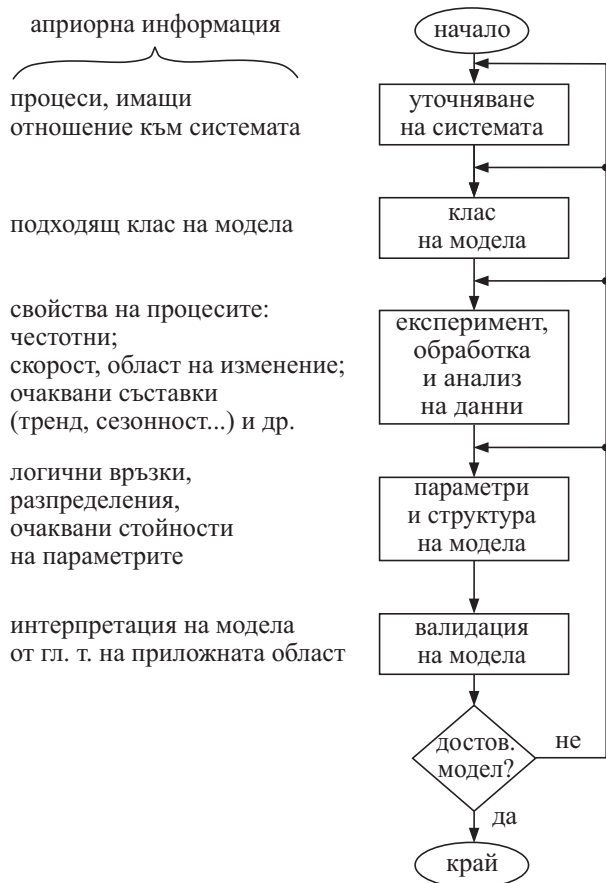
1.3.6 Връзка с аналитичния подход

Тази, вече засегната, тема е представена графично на Фигура 1.22. Тъй като при чисто експерименталния подход достоверността на модела не може да се потвърди с помощта на априорна информация, има опасност моделът да отразява лъжливо особеностите на системата. Затова е важно идентификацията да се комбинира с аналитичния подход. Със следващия пример от финансите се представя участието на предварителните знания в идентификацията.

При изграждане на системите за оценяване на кредитен риск, основно се използва експерименталният подход. Но дори и в този случай, при избора на фактори, които да участват в модела, се извършва проверка за некоректна асоциация между входовете (факторите) и изхода (вероятността кредитоискателят да е добър платец). Например на базата на предварителните знания е известно, че факторът ‘възраст’ е положително корелиран с изхода. Това правило се залага в процеса на моделиране и ако според модела връзката е обратна, то е необходимо да се преразгледат дейностите в предходните етапи, като може да се стигне и до изключване на фактора ‘възраст’ от модела.

Също така може предварително да се извърши (статистическа) сегментация на данните, например според фактори като ‘доход’, ‘семеино

положение' и др., като отново се търси логика в сегментацията. След това може да се окаже, че в определени сегменти от населението възрастта е значим фактор, а в други – не. Така с търсенето на бизнес логика моделът става по-ясен и обоснован, а вероятността от допускане на грешки, дължащи се на вариацията в данните, намалява.



Фигура 1.22. Експериментален подход и възможности за комбиниране с аналитичния подход

1.4 Автоматизирана идентификация

Всеки етап на идентификацията обикновено е свързан с вземане на решения за следващите стъпки, които трябва да се предприемат. Когато системата е с малка размерност, е обоснована активната намеса на експерт. Но с развитието на компютърните технологии и информационните системи става възможно прилагането на идентификацията и за системи с много голямата размерност на входно-изходните сигнали. Така се стига до необходимостта идентификацията да се автоматизира [81, 79, 73].

Резултатът от автоматичното изпълнение на идентификацията зависи освен от конкретната методология (какви дейности и в каква последователност се изпълняват) и от параметрите, с които изследователят „настройва“ цялостния процес на моделиране. Примери за такива настройваеми параметри са флагове, с които се разрешава/забранява прилагането на конкретна обработка, на анализ или на метод, както и за избор на различни начални условия; прагови стойности, например на допустимата степен на мултиколинеарност в данните, на значимостта на потенциалните фактори за добавяне в (или премахване от) модела, прагове при числено устойчивото обръщане на матрици и др.; критерии за спиране на стъпкови алгоритми за определяне на структурата; максимален брой фактори в модела и т.н. Изследователят задава начални стойности на настройваемите параметри и в зависимост от крайния резултат (модел) той изменя стойностите на параметрите. От една страна, колкото повече са настройваемите параметри, толкова по-голяма е възможността да се въздейства върху процеса на моделиране. Една причина за тази човешка намеса е недостатъчната възможност на автоматизирания процес да оптимизира споменатите параметри. Този проблем теоретично е решим – възможно е с многократно изпълнение на цикъла на идентификацията даден параметър да се оптимизира. Но практически автоматизираният алгоритъм много се усложнява и достигането евентуално до желан резултат би отнело твърде много време. Друга причина за наличието на тези параметри е трудността при дефинирането на целта на оптимизацията на модела (например, когато има повече цели, между които трудно се балансира, особено в началото на моделирането).

При разработването на автоматизирания цикъл се търси компромис между броя на споменатите параметри и степента на влияние върху процеса на изграждане на модела. Част от параметрите може да останат

скрити, като прагът на значимост на сингулярните числа, определящ резултата от числено устойчивото обръщане на матрици, или да имат подходящи начални стойности и така вниманието на изследователя да се насочи към малък брой основни параметри, с които да влияе върху резултата.

Освен това за по-лесно боравене с параметрите е желателно така да се изгради автоматизираната логика, че промяната на един параметър да не изисква промяна на други параметри, т.е. настройката им да е по възможност независима.

Друг важен момент е настройваемите параметри да се интерпретират ясно, без да е нужно задълбочено познание за конкретната реализация на методите. Така автоматизираният процес може да се използва от хора, които не са експерти по идентификация, а да имат по-задълбочени познания в приложната област.

Освен задаването на начални стойности на настройваемите параметри от помощ би била и информацията за интервалите на тяхното изменение (независещи от конкретния проблем). Също е желателно, доколкото е възможно, чувствителността на процеса на моделиране да е равномерна в различните части на допустимите интервали. Така се улеснява работата на изследователя и се намалява броят на итерациите с човешка намеса.

Отделено е специално внимание на автоматизирането на идентификацията, защото често времето за моделиране е ограничено, а хората, прибегващи до автоматизирано формиране на модел, в болшинството от случаите са експерти в изследваната област, а не в математическия анализ, оптимизацията, числените методи, статистиката и въобще в областите, необходими за имплементацията на етапите от описания цикъл.

Глава 2

Етапи на идентификацията на MIMO системи

2.1 Клас на модела

В първия етап (определяне границите на системата) акцентът е реалната система и тъй като неговото провеждане е специфично за конкретната задача, той не е разгледан по-подробно.

Във втория етап вниманието се насочва към модела. Понеже математическите описания са абстрактни представяния на реалните явления, това позволява тяхното обобщаване в отделни класове. Поради това вторият етап не зависи толкова от конкретната система, колкото първия. Освен това, след като системата е конкретизирана, изследването на свойствата ѝ също може да се обобщи. Затова изложението започва с избора на клас на модела.

Системите имат характерни белези и съответно класът на модела се избира по такъв начин, че да отрази тези особености. Ако е известно, че реалните процеси протичат във времето, то и моделът трябва да е динамичен. Ако връзката между входно-изходните величини е изразено нелинейна, то и моделът трябва да спада към класа на нелинейните модели и т.н. Изборът на класа зависи и от задачата, за която се търси описание на системата. Например, ако става дума за класификация, възможен модел е дърво на решенията, но ако целта е апроксимация, то подходящо е да се стартира например с линеен модел, като при нужда описанието да се усложни.

От друга страна, неподходящо избраният клас на модела може да доведе до изместени оценки на параметрите, числени проблеми в процеса на моделиране, грешни изводи относно поведението на реалната система и др.

2.1.1 Избор на клас на модела

Най-общо изборът на класа зависи от следните фактори:

- особености на системата (както беше споменато, моделът се избира от класа на линейните/нелинейните модели, стационарните/нестационарните модели и т.н.);
- приложение на модела (например за класификация при оценка на кредитен риск или вероятност за заболяване);
- установени практики (свободата при избора на подходящ клас е силно ограничена в консервативни области като финансова индустрия, където бизнесът има определени изисквания, очаквания и разчита на установени практики);
- средства за моделиране (функционалността на използвания софтуерен продукт може значително да ограничи множеството от класове на модела);
- опит на изследователя (често специалистите в приложни области не са с достатъчно дълбоки математически познания, което също може да повлияе на избора на клас на модела);
- време за идентификация (понякога това ограничение е решаващо при избора на класа на модела).

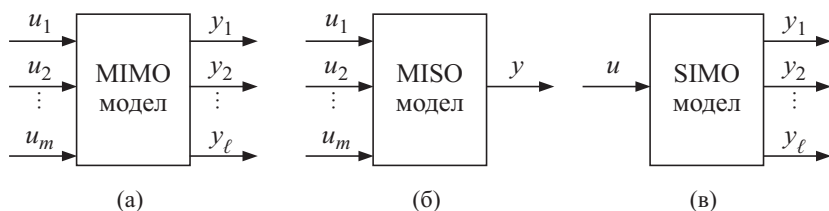
2.1.2 Класификации на моделите

По-долу са описани класификации [5, 10, 21, 30], които имат пряко отношение към моделите, засегнати в монографията. Целта на представянето им е най-вече да се очертаят границите, до които се разпростират разглежданията. Класовете модели, за които е валидна представената теория, са подчертани.

2.1.2.1 Брой входове и изходи

Според броя на входните и изходните величини, които участват в описанието, моделите може да се разделят на:

- модели с един вход и един изход (SISO – Single Input Single Output);
- модели с много входове и много изходи (MIMO – Multiple Input Multiple Output);
- модели с много входове и един изход (MISO – Multiple Input Single Output);
- модели с един вход и много изходи (SIMO – Single Input Multiple Output).



Фигура 2.1. Многомерни модели: MIMO (а), MISO (б) и SIMO (в)

SISO са частен случай на MIMO моделите и при нужда теорията в монографията лесно може да се сведе до едномерния вариант. По тази причина на SISO описанията не е отделено специално внимание. От друга страна, преходът от SISO към MIMO системи невинаги е директен, поради някои характерни особености на MIMO системите, възможните проблеми, свързани с тях и съответните решения.

Съществуват подходи за преобразуване на MIMO структури в SISO. Те са практически приложими за системи със съвсем малко входове и изходи, което силно ограничава приложението им. Използват се предимно при изграждане на системи за управление в техниката и са напълно неприложими при извличането на информация от данни за обекти с голям брой входно-изходни величини. Основният недостатък на това преобразуване е необходимостта от въвеждането на компенсатори, които да неутрализират влиянието на кръстосаните връзки (тук възниква и необходимостта да се дефинират прави (сепаратни) и кръстосани връзки в MIMO структурата). Освен това синтезът на компенсатори силно се усложнява за системи с повече от два входа и два изхода, а тъй като компенсаторите са зависими от параметрите на модела, при настъпване на нестационарност се налага тяхното преизчисляване. По тези причини преобразуването на MIMO системите в SISO е засегнато слабо в точка 2.1.4.2, най-вече за да се подчертаят предимствата от запазването на структурата на системата като MIMO и съответно изграждането на MIMO модел. При това положение не се налага да се въвеждат понятия като сепаратни и кръстосани връзки, а системата се разглежда като единно цяло, преобразуваща векторното входно въздействие във векторна реакция. Основното предимство на изложената теория е, че тя не се усложнява с нарастването на входовете и изходите. Моделите, методите за изграждането им и въобще всички дейности остават същите, като при автоматизирания случай броят на сигналите оказва влияние единствено на времето за идентификация. На Фигура 2.1 са представени структурите на многомерни модели според размерността на входовете и изходите им. Тъй като (б) и (в) са частни случаи на MIMO структурата,

за постигане на най-голяма общност описанията в монографията са за ММО модели, а споменатите частни случаи се използват в примери за онагледяване на теорията.

Понякога декомпозирането на ММО системите на подсистеми е необходима стъпка в процеса на моделиране. Такива случаи също са разгледани в текста, например декомпозирането на пазарна система на MISO или ММО подсистеми.

2.1.2.2 Разпределеност на параметрите

При някои обекти за съставянето на модел е необходимо да се отчете пространственото или друго разпределение на параметрите. В това отношение описанията се разделят на:

- модели с разпределени параметри;
- модели със съсредоточени параметри.

Моделите с разпределени параметри са типични за аналитичното моделиране, а тъй като данните са стойности на измерени или изчислени величини, то по естествен начин моделите, които са резултат от идентификацията, са със съсредоточени параметри. В монографията е разгледано изграждането на модели от втория тип.

2.1.2.3 Случайна съставка

Както стана дума, важно за идентификацията е отчитането на неопределеността. От тази гледна точка се разграничават:

- детерминирани модели;
- стохастични модели.

При детерминираните модели връзката между факторите и изходите е еднозначна. От друга страна, когато се отчита наличието на стохастични влияния в системите, зависимостите се проявяват при голям брой наблюдения. На едно и също входно въздействие в даден момент от времето отговарят множество случайно разпределени изходни наблюдения, характеризиращи се със съответна функция на разпределение. С други думи, изходите са подложени на влиянието на неизвестни фактори (источниците на неопределеност са разгледани в точка 1.1.4). Тъй като всеки изход съдържа случайна компонента, не е възможно той да се предскаже точно, а може само да се оцени с определена вероятност. Рядко неопределеностите, влияещи на поведението на системата, може да се пренебрегнат. Затова акцентът в изложението е върху стохастичните модели.

2.1.2.4 Функционална зависимост на изхода от факторите

В много приложения на моделите, например при управлението на процеси, е важен видът на функционалната връзка между входните и изходните сигнали, тъй като той определя сложността на управляващите устройства, на евентуални компенсиращи звена, оценители на състояния и други елементи на системите за управление. От тази гледна точка моделите се делят на:

- линейни модели;
- нелинейни модели.

За линейните модели, при които важи принципът на суперпозицията, теорията е обща и решенията са добре изследвани. От друга страна, изграждането и използването на нелинейните модели значително се усложнява. В монографията е разгледано както построяването на линейни, така и на клас нелинейни модели, който е споменат по-долу.

2.1.2.5 Зависимост на изхода от параметрите

Според връзката на изхода от параметрите моделите може да се разделят на:

- линейно параметризирани модели;
- нелинейни по параметри модели.

В монографията значението на термина параметризация е в смисъла, използван в [148], а именно представянето на даден модел като функция на вектора или матрицата на параметрите, които се оценяват в процеса на идентификацията. В този смисъл, ако при това представяне моделът е линеен по параметри (т.е. изходът е линейна функция на параметрите), то той се нарича линейно параметризиран модел.

Такива модели са разгледани по-подробно в точка 2.1.3.2. Въпреки че в общия случай те може да са нелинейни според предишната класификация, изграждането им не се отличава принципно от това на линейните модели. Както ще бъде показано в точка 2.1.3.2, общият вид на линейно параметризираните не се отличава от вида на линейните модели. Ето защо за идентификацията е важна не функционалната връзка между изходните величини и факторите, а между изходите и параметрите на модела. Ако изходът (а оттам и остатъкът) зависи линейно от параметрите на модела, то целевата функция, когато е сума от квадрата на остатъка, зависи квадратично от параметрите. Тогава техните оптимални оценки може директно да се определят от наличните данни. В противен случай процесът на моделиране значително се усложнява. Поради това в някои източници, например [9], под линеен модел се разбира

такъв, който е линеен по отношение на параметрите, а не по отношение на връзката между входа и изхода.

От друга страна, когато моделът е нелинеен по параметри, моделирането е свързано с използването на итеративни процедури. Има някои модели от този клас, които може да се представят, макар и условно, в линейно параметризиран вид, с което оценяването на параметрите им се свежда до многократно оценяване на това условно представяне. Но в общия случай нелинейните по параметри модели изискват оценители, които се извеждат за избраната функционална зависимост. Това е основната причина често да се търси (макар и условно) линеен модел по отношение на параметрите и ако се окаже, че това представяне не е подходящо, тогава да се прибегне до описание, което е нелинейно по параметри.

В монографията се разглеждат модели, които може да се представят като линейно параметризирани, и такива, които не е възможно да се запишат в този вид, но е подходящо да се приемат за линейни по параметри.

2.1.2.6 Зависимост на изхода от предисторията на факторите

В някои случаи, когато неопределеността е значителна, а априорна информация за евентуална динамика липсва, се използват статични зависимости, отразяващи връзката между наблюдаваните величини. Пример за това са някои икономически системи, за които се използват модели на статиката. Също така, понякога апостериорните данни предоставят информация само за статичното поведение на системата. С такива данни не е възможно да се построи динамичен модел. От друга страна, когато данните осигуряват информация за динамиката на системата, напълно естествено е да се използват динамични модели.

Според това, дали описанието отразява развитието на сигналите във времето, се разграничават:

- статични модели;
- динамични модели.

Въпреки че статичните модели са частен случай на динамичните, има дейности в етапите на идентификацията, които са характерни за статични описания на системата и са неприложими за отчитането на динамиката. Например при обработката на данни за формиране на статичен модел често срещана дейност е т.нар. категоризация на входни и/или изходни величини (описана в точка 2.2.2.9). При тази обработка акцентът не е върху развитието на сигналите във времето, а върху мо-

ментните им стойности. За статични системи това е полезна дейност, с която се потиска в определена степен неопределеността в данните, а от друга страна се запазва основната зависимост между категоризирани входни величини и изхода. Освен това може адекватно да се отчетат някои предварителни знания за системата, улеснява се тълкуването на модела и др. (виж точка 2.2.2.9). Но ако се търси динамичен модел, където стойностите на изхода зависят от предисторията, категоризирането на данните би влошило значително поведението на модела. Ето защо, въпреки че статичните модели са частен случай на динамичните, в монографията се разглеждат и характерни само за статичните описания особености на изграждането на модел.

2.1.2.7 Зависимост на поведението на модела от времето

В зависимост от това, дали описанието остава постоянно във времето, моделите може да се разделят на:

- стационарни модели;
- нестационарни модели.

Ако са взети под внимание евентуални изменения в поведението на системата, математическото описание е нестационарно. В противен случай моделът е стационарен, т.е. поведението му не се променя с времето (структурата и параметрите остават постоянни). Характерно за нерекурсивното оценяване е, че за конкретни данни се получава съответен модел, който е крайният резултат от идентификацията. В този смисъл, тъй като моделът не се актуализира с времето, може да се разглежда като стационарен. От друга страна, в монографията са описани техники, свързани и с отчитане на нестационарност в поведението на системата. От тази гледна точка е засегнато формирането и на двата класа модели.

2.1.2.8 Времева скала

Според това, дали сигналите са непрекъснати във времето или дискретни, се разграничават:

- непрекъснати модели;
- дискретни модели.

Някои реални процеси имат дискретен характер и съответно моделите, които ги описват, са дискретни. Например производството на детайл може да се опише като последователност от операции (доставка на материал, позициониране на заготовката, поставяне/смяна на инструмент, обработка, снемане на детайла и т.н. [15]). Други примери за

такива процеси обхващат доставката на продукти в склад, продажбите на продукт, изтичане на срок на годност и т.н. Освен използването на дискретни модели за описание на такива процеси често се срещат и дискретни описания на непрекъснати процеси, например свързани с изменение на температурата в пещ, завъртания и трансляции, описващи движение на кинематична система, и др. Използването на дискретни модели може да е продиктувано от целта, за която се изграждат, например, ако моделът е нужен за управление и контролерът е дискретен. Както вече беше споменато, в монографията се разглеждат само дискретни модели поради отчитането на данните в дискретни моменти от времето.

Въпреки че идентификацията използва дискретни данни, има подходи, в резултат на които директно се получава непрекъснато описание. Такъв е случаят с графоаналитичните методи, при които в зависимост от формата на преходния процес се определят параметрите на непрекъснат модел. Въпреки че тези методи може да се автоматизират, основен недостатък е нуждата от провеждането на активен експеримент. Освен това при многомерните системи с много входове експериментите (ако са допустими) може да отнемат твърде много време, тъй като трябва да се проведат множество независими експерименти, за да се снесе реакцията на всички изходи при прилагане на въздействие по всеки от входовете. Поради това тези методи, и съответно формирането на непрекъснати модели, не са засегнати в изложението.

2.1.2.9 Брой на параметрите

Броят на параметрите в модела е признак, според който се разграничават:

- параметрични модели;
- непараметрични модели.

Параметричните модели имат краен брой параметри, а непараметричните – теоретично безкраен брой. В практиката по-широко приложение намират параметричните модели. Причини за това са значително по-малкият брой на оценяваните параметри, както и възможността успешно да описват динамиката на обекти с много различно поведение. Недостатък на тези модели е, че при неподходящо избрана структура е възможно динамиката им чувствително да се различава от тази на обекта. Затова една от важните стъпки в идентификацията е изборът на структурата на модела.

2.1.3 Регресионни модели

Терминът „регресия“ е въведен от английския антрополог Франсис Галтон. С него той нарича тенденцията родителите с по-висок ръст от нормалния да имат деца с по-близък ръст до средния. Този факт Галтон нарекъл „*regression to mediocrity*“. От съвременна гледна точка това название е неподходящо [8], имайки предвид сегашния смисъл на регресионния модел, а именно – описание на връзката между множество от входни и друго множество от изходни величини. Понякога входовете се наричат въздействия, независими или описателни променливи/характеристики, а ако моделът е статичен, също се наричат фактори, регресори, предиктори и др. Изходите се наричат още: реакции, зависими или описвани променливи/характеристики, признаци и др. Въпреки че някои названия са взаимозаменяеми, важно е да се прави разлика между тях. Например фактори, регресори и предиктори в динамичен регресионен модел обикновено са изместени във времето входно-изходни величини (или техни функции). Затова е желателно, когато се набляга на зависимостта на изхода от множество величини, те да се наричат фактори, регресори или предиктори. Но когато става дума за външни сигнали, влияещи на системата и отчетени от модела, те да се наричат входни сигнали, въздействия, независими променливи и т.н.

В някои източници се прави разлика между фактор и регресор, като под регресор се има предвид променлива, която участва в модела, а фактор е реална, физическа величина. В този смисъл, ако даден фактор се трансформира, например с цел получаване на линейно параметризиран модел, то трансформираната величина е регресор, а първоначалната – фактор. Естествено, ако даден фактор участва директно в модела, той е и регресор. В монографията не се прави разлика между двете понятия, защото от гледна точка на оценяването на параметри и избора на структурата е важен типът на модела, а не пътят, по който е получен. Още повече често в литературата векторът на регресорите се означава с буквата φ (от фактори).

2.1.3.1 Общ вид

Тъй като данните са дискретни, моделите също са дискретни. Индексът на текущото наблюдение е k . Ако данните са функция на времето, то наблюдението в предишния дискретен момент е с индекс $k - 1$. Когато се търси статичен модел, подредбата на данните във времето не е определяща. Например в набор, съдържащ еднократни (не периодично отчитани) данни за пациенти, наблюденията може да са подредени не по

времето на провеждане на медицинското изследване, а по азбучен ред. В този случай индексът k отговаря на текущ пациент според въведената подредба, а не на момент от времето. Тъй като основно се разглеждат динамични системи, ако не е указано друго, k ще има смисъл на дискретен момент от времето.

В общ вид регресионният модел може да се запише като

$$\hat{y}_k = f(\varphi_k, \theta), \quad (2.1)$$

където $\hat{y}_k$ е изходът на модела, θ и φ_k са съответно вектор на параметрите и вектор на регресорите, с помощта на които се описва реакцията на обекта y_k в текущия момент (или отговарящ на k -тото наблюдение, при статична система). За разлика от едномерните системи, където броят на факторите и параметрите е еднакъв, при многомерните системи обикновено броят им е различен. Ако случаят е такъв, броят на факторите ще е z , а броят на параметрите ще е означен с p .

Тъй като задачата на идентификацията е да се намери описание на връзката между известните входно-изходни величини, то в (2.1) е нужно изходът на все още неизвестния модел да се замени с измерения изход на системата, т.е.

$$y_k = f(\varphi_k, \theta) + e_k. \quad (2.2)$$

С други думи зависимостта между изхода на системата и на модела е $y_k = \hat{y}_k + e_k$, като e_k е обобщен сигнал, който отразява шума от измерване, смущенията от околната среда и несъвпадението между регресионната функция $f(\cdot)$ и реалната връзка между факторите и изхода. За опростяване на идентификацията, обикновено се приема, че e_k е адитивен сигнал (както е записано в (2.2)).

2.1.3.2 Общ вид на линейно параметризиран модел

В много случаи с подходящи трансформации на факторите и/или на изхода регресионните модели може да се представят в линейно параметризиран вид. Това позволява един и същ оценител на параметрите да се използва за получаване както на линейни, така и на нелинейни модели (които се свеждат до линейни по параметри). В други случаи, дори да не съществуват такива трансформации, за някои регресионни модели е подходящо да се пренебрегнат нелинейните зависимости и да се приеме, че изходът е линейна функция на параметрите.

Общ вид на МІМО модел

В едномерния случай линейно параметризираният (SISO) модел може да се запише така:

$$y_k = \varphi_k^T \theta + e_k. \quad (2.3)$$

Тук θ и φ_k са вектори с еднаква размерност, а y_k и e_k са скаларни сигнали. Когато системата е с повече изходи (нека $y_k \in \mathcal{R}^\ell$), тъй като откъсно на равенство (2.3) се намира вектор, то и резултатът от произведението на факторите и параметрите също трябва да е вектор (отговарящ на изхода $\hat{y}_k$ на многомерния модел). Това означава, че горното умножение трябва да се извърши между матрица и вектор. Така възникват две представяния на линейно параметризирания МІМО регресионен модел в общ вид [74, 75]. При едното параметрите се подреждат във вектор, а факторите – в матрица с подходяща структура, докато при другото представяне факторите са във вектор, а параметрите в матрица. Първият запис на МІМО модел в общ вид е

$$y_k = \Phi_k \theta + e_k, \quad (2.4)$$

където векторът $\theta \in \mathcal{R}^p$ се състои от всички параметри на модела, а матрицата $\Phi_k \in \mathcal{R}^{\ell \times p}$ съдържа стойностите на регресорите, описващи изхода на системата в текущия момент. Другото представяне е

$$y_k = \Theta^T \varphi_k + e_k, \quad (2.5)$$

при което параметрите са подредени в матрицата $\Theta \in \mathcal{R}^{z \times \ell}$, а векторът $\varphi_k \in \mathcal{R}^z$ съдържа стойностите на регресорите. На пръв поглед няма значение как се формира $\hat{y}_k$ – и в двата случая изходът е линейна функция на параметрите и на факторите. Въпреки това (2.4) и (2.5) са свързани с различни особености, които са важни при уточняването на многомерната структура. Причината за различията се дължи на структурата на матрицата или вектора на параметрите. При оценяването на $\hat{\Theta}$ или $\hat{\theta}$, всеки елемент е оценка, която отговаря на уникален параметър в модела. Тази оценка трябва да е независима от останалите и да се среща еднократно в $\hat{\Theta}$ или $\hat{\theta}$. Също така не е допустимо матрицата/векторът на параметрите да съдържа елементи, които не се оценяват (например нулеви елементи, които нямат смисъл на параметри, понеже не са включени в структурата на модела). От друга страна, при конкретна подредба на параметрите, регресорите се групират по такъв начин, че резултатът от произведението да е векторът $\hat{y}_k$. При подредбата на факторите не се налага ограничението те да се срещат еднократно, както и се допуска въвеждането на нулеви или други елементи, ако е необходимо.

За да се наблегне на важната, за идентификацията, подредба на параметрите (в смисъл на вектор или матрица), са избрани имена на двете ММО представяния, които са свързани именно с подредбата на параметрите, а не на регресорите.

Структурата на матриците и векторите в общите записи са разгледани по-подробно във втората част на точка 2.1.3.3.

Общ вид на MISO и SISO модели

Когато моделът е с един изход, регресорите и параметрите е удобно да се групират във вектори и в този случай общото представяне е

$$y_k = \varphi_k^T \theta + e_k. \quad (2.6)$$

То не се отличава от вече показаното описание на SISO моделите. Съответно изходът, изчислен от модела, е

$$\hat{y}_k = \varphi_k^T \theta. \quad (2.7)$$

Линейна параметризация

Представянето на различни модели в общ вид дава възможност да се извеждат общи оценители на параметри, общи методи за избор на структурата и др. Освен това съществен момент е, че параметрите и регресорите в (2.4) и (2.5) са в отделни матрици и вектори, а това улеснява извеждането на съответните алгоритми.

С помощта на долния пример се показва разликата между моделите с линейна връзка между входно-изходните сигнали и y_k , линейно параметризираните, както и моделите, които не може да се приведат във вида (2.4) или (2.5).

Пример. *Регресионни модели и представяне в линейно параметризиран вид*

Тъй като тук не е от значение многомерната структура, а зависимостта на изхода от факторите и параметрите, за да не се усложнява примерът, се разглеждат MISO модели. Нека моделът е с два входа, т.е. $u_k = [u_{1,k} \ u_{2,k}]^T$. Тъй като y_k е скаларен сигнал, параметрите и регресорите може да се подредят във вектори.

Случай 1. Линейна връзка между първоначалните фактори и изхода на модела $\hat{y}_k$

Нека изходът на модела е

$$\hat{y}_k = \theta_1 y_{k-1} + \theta_2 u_{1,k-3} + \theta_3 u_{2,k-1} = \varphi_k^T \theta. \quad (2.8)$$

Тук векторите на параметрите и на регресорите са $\theta = [\theta_1 \ \theta_2 \ \theta_3]^T$ и $\varphi_k = [y_{k-1} \ u_{1,k-3} \ u_{2,k-1}]^T$.

Случай 2. Нелинейна връзка между първоначалните фактори и $\hat{y}_k$

Нека изходът на модела е

$$\hat{y}_k = \theta_1 y_{k-1}^2 + \theta_2 \ln u_{1,k-1} + \theta_3 e^{u_{2,k-1}} = \varphi_k^T \theta, \quad (2.9)$$

където $\theta = [\theta_1 \ \theta_2 \ \theta_3]^T$, $\varphi_k = [y_{k-1}^2 \ \ln u_{1,k-1} \ e^{u_{2,k-1}}]^T$. Както се вижда, моделът е линеен по параметри, а за привеждане в общ вид към първоначалните фактори се прилагат нелинейни трансформации. В резултат на това се получават новите фактори във вектора φ_k .

Случай 3. Нелинейна връзка между първоначалните фактори и $\hat{y}_k$

Нека изходът на модела е

$$\hat{y}_k = e^{\theta_1 y_{k-1} + \theta_2 u_{1,k-3} + \theta_3 u_{2,k-1}}. \quad (2.10)$$

След логаритмуване на двете страни на (2.10) се получава новата зависимост:

$$\ln \hat{y}_k = \theta_1 y_{k-1} + \theta_2 u_{1,k-3} + \theta_3 u_{2,k-1}.$$

Нека изходът на този нов модел се означаи с $\tilde{y}_k = \ln \hat{y}_k$. Така отново се получава модел с желаната линейна структура:

$$\tilde{y}_k = \theta_1 y_{k-1} + \theta_2 u_{1,k-3} + \theta_3 u_{2,k-1} = \varphi_k^T \theta,$$

където $\theta = [\theta_1 \ \theta_2 \ \theta_3]^T$ и $\varphi_k = [y_{k-1} \ u_{1,k-3} \ u_{2,k-1}]^T$. За разлика от предишния случай тук се трансформира изходът y_k , а не първоначалните фактори.

Случай 4. Нелинейна връзка между първоначалните фактори и $\hat{y}_k$

Нека изходът на модела е

$$\hat{y}_k = (\theta_1 y_{k-1} + \theta_2 u_{1,k-1})^2 + \theta_3 u_{2,k-1}. \quad (2.11)$$

Този модел, за разлика от предишните, не може да се приведе в линеен вид по отношение на параметрите. Например, ако се разкрият скобите, (2.11) добива вида:

$$\hat{y}_k = \theta_1^2 y_{k-1}^2 + 2\theta_1 \theta_2 y_{k-1} u_{1,k-1} + \theta_2^2 u_{1,k-1}^2 + \theta_3 u_{2,k-1} = \varphi_k^T \tilde{\theta}.$$

На пръв поглед моделът $\hat{y}_k = \varphi_k^T \tilde{\theta}$ е линейно параметризиран – параметрите са отделени от регресорите, а $\hat{y}_k$ зависи линейно от $\tilde{\theta}$. В случая

$$\begin{aligned} \varphi_k &= [y_{k-1}^2 \ 2y_{k-1}u_{1,k-1} \ u_{1,k-1}^2 \ u_{2,k-1}]^T, \\ \tilde{\theta} &= [\theta_1^2 \ \theta_1\theta_2 \ \theta_2^2 \ \theta_3]^T. \end{aligned} \quad (2.12)$$

Проблемът на това описание е, че ако се оцени векторът $\tilde{\theta}$, приемайки, че моделът е линеен по параметри, то при прехода от $\tilde{\theta}$ към θ се получава система от 4 уравнения с 3 неизвестни, която в общия случай няма решение. Това води до налагане на ограничението:

$$\tilde{\theta}_2 = \sqrt{\tilde{\theta}_1 \tilde{\theta}_3}.$$

Освен това $\tilde{\theta}_1$ и $\tilde{\theta}_3$ трябва да са неотрицателни според (2.12), а ограничения при оценяване на параметрите на линейно параметризираните модели не се налагат.

От друга страна, ако $\tilde{\theta}$ беше например

$$\tilde{\theta} = [\theta_1^3 \ \theta_1 \theta_2 \ \theta_3]^T, \quad (2.13)$$

въпреки че елементите на $\tilde{\theta}$ са нелинейно свързани с елементите на θ , а също са и зависими помежду си, при прехода към първоначалните параметри се решава система с 3 уравнения и 3 неизвестни. Тя има еднозначно решение, независимо от стойностите на оценките на $\tilde{\theta}$, над които не са наложени ограничения според (2.13). В такъв случай моделът $\hat{y}_k = \varphi_k^T \tilde{\theta}$ би бил коректно линейно параметризирано представяне на (2.11). $\diamond$

В примера се приема, че са известни нелинейните трансформации, с които се описва реалната система. Често обаче тази информация липсва. Тогава се налага да се използват статистически методи, с които да се определят подходящи трансформации на факторите или на изходите. Тази дейност е описана в точка 2.2.2.5.

В много случаи от практиката се извършват такива трансформации. Например в пазарните системи, където изходите са продажбите на продуктите, а входовете са цените, рекламите, промоциите и др., широко са разпространени следните модели:

- експоненциален модел [149, 144]

$$\hat{y}_k = \theta_0 e^{\theta_1 u_{1,k-1} + \dots + \theta_m u_{m,k-1}}; \quad (2.14)$$

- log-log модел [107, 108, 117, 116]

$$\ln \hat{y}_k = \theta_0 + \theta_1 \ln u_{1,k-1} + \dots + \theta_m \ln u_{m,k-1}; \quad (2.15)$$

- полулогаритмични модели [144]

$$\ln \hat{y}_k = \tilde{\theta}_0 + \theta_1 u_{1,k-1} + \dots + \theta_m u_{m,k-1}; \quad (2.16)$$

$$\hat{y}_k = \theta_0 + \theta_1 \ln u_{1,k-1} + \dots + \theta_m \ln u_{m,k-1}. \quad (2.17)$$

Лесно се вижда, че модел (2.16) се получава след логаритмуване на модел (2.14), където $\tilde{\theta}_0 = \ln \theta_0$. В [55] е използван статичен линеен (като връзка между u и y) модел за описание на пазарна система, но като цяло този модел не води до добър резултат. По тази причина в [134] са разработени методологии за изграждане и на нелинейно параметризирани модели.

Свободният член θ_0 в моделите (2.14)–(2.17) отразява факта, че описанието е построено в околност на работна точка, която не съвпада с началото на първоначалната координатна система на статичната характеристика. Регресионните модели със свободен член са разгледани подробно в точка 2.4.2.

Също така в различни методологии за изграждане на модели на поведението на кредитоискатели се използва т.нар. трансформационен анализ. С него се подбират такива нелинейни функции на факторите (характеристиките на кандидатите за кредит), че връзката между изхода (вероятността поведението им да е добро) и трансформираните фактори да се доближи максимално до линейната.

2.1.3.3 Типови линейни регресионни модели

В тази точка са представени основните типове линейни регресионни модели. Акцентът е върху динамичните модели, като е отделено внимание и на статичните описания. Тъй като дискретните динамични модели често се представят в полиномен вид, се въвежда операторът q за изместване на сигналите във времето. Така например един такт закъснение на сигнала x_k може да се представи като $x_{k-1} = q^{-1}x_k$, а съответно n такта закъснение се записват като $x_{k-n} = q^{-n}x_k$. Както се вижда при това описание, сигналите са представени във времевата област, а не в Z областта, която често се използва при описанието на дискретни системи. Поради удобството, че сигналите остават във времевата област при представянето на модела в полиномен вид, в монографията ще се използва операторът q (а не операторът z).

Ако системата по природа е нелинейна, в смисъл на връзката между входа и изхода, но моделът е линеен по отношение на параметрите, то в долните разглеждания е прието, че нелинейните трансформации вече са приложени към съответните величини.

SISO модели

Първо са представени модели с един вход и един изход, а в следващата подточка са дадени едни от най-често използваните многомерни описания, които се използват и в монографията.

Авторегресионен модел (AR)

Този модел отразява връзката между изхода y_k и неговата предистория от na предишни такта. По тази причина моделът често се използва за описание на времеви редове [3]. С отчитане на остатъка, описанието може да се запише в следния полиномен вид:

$$a(q^{-1})y_k = e_k, \quad (2.18)$$

като полиномът $a(q^{-1})$ е

$$a(q^{-1}) = 1 + a_1 q^{-1} + \dots + a_{na} q^{-na}. \quad (2.19)$$

Ако параметърът на свободния член в (2.19) не е $a_0 = 1$ (това може да се получи вследствие на аналитично моделиране на системата), то моделът лесно се привежда във вида (2.19), като двете страни на равенство (2.18) се разделят на a_0 . Така в новото представяне параметърът на свободния член става единица.

Моделът (2.18) може да се запише като:

$$(1 + a_1 q^{-1} + \dots + a_{na} q^{-na}) y_k = e_k$$

и след отчитане на оператора q и изразяване на y_k се получава:

$$y_k = -a_1 y_{k-1} - \dots - a_{na} y_{k-na} + e_k.$$

Тъй като описанието е линейно по отношение на параметрите a_i , последното равенство, представено във векторна форма (виж общия линейен вид (2.3), е

$$y_k = \varphi_k^T \theta + e_k.$$

Векторите на регресорите и на параметрите са:

$$\varphi_k = [-y_{k-1} \quad -y_{k-2} \quad \dots \quad -y_{k-na}]^T,$$

$$\theta = [a_1 \quad a_2 \quad \dots \quad a_{na}]^T.$$

Пълзяща средна стойност или пълзящо средно (МА)

При този модел изходът y_k е претеглена сума от предишни стойности на друг сигнал. В икономиката, където системите често се изучават с времеви редове, този друг сигнал обикновено се разглежда като стохас-

тичен (случаен във времето). Тук реакцията не зависи от детерминирани величини. Нека y_k се формира от nc на брой предишни стойности на e_k . Тогава полиномният вид на МА модела е

$$y_k = c(q^{-1})e_k. \quad (2.20)$$

Полиномът $c(q^{-1})$ е

$$c(q^{-1}) = 1 + c_1q^{-1} + \dots + c_{nc}q^{-nc}. \quad (2.21)$$

Моделът (2.20) може да се запише като:

$$y_k = c_1e_{k-1} + \dots + c_{nc}e_{k-nc} + e_k. \quad (2.22)$$

В общото представяне (2.3) на връзката между измерения изход и регресорите e_k има смисъл на остатъчна съставка в y_k , която не е отчетена от модела. По тази причина се приема, че изчисленият изход на МА модела е $\hat{y}_k = c_1e_{k-1} + \dots + c_{nc}e_{k-nc}$, а наличието на e_k в (2.22) осигурява участието на y_k в равенството, а не на $\hat{y}_k$. Така моделът от тип МА може да се разглежда като прозорец, който се измества (пъзли) при формирането на всяка следваща стойност на изхода, като най-старата стойност на остатъка отпада от израза, а се включва най-нова стойност. Например изходът в $k+1$ -ия такт (т.е. при $k \leftarrow k+1$ в (2.22)) се формира по следния начин:

$$y_{k+1} = c_1e_k + \dots + c_{nc}e_{k-nc+1} + e_{k+1}.$$

Отново, тъй като регресионният модел МА е линеен по отношение на параметрите, (2.22) може да се представи във вида (2.3), като векторът на регресорите в k -ия такт и векторът на параметрите са:

$$\varphi_k = [-e_{k-1} \quad -e_{k-2} \quad \dots \quad -e_{k-nc}]^T,$$

$$\theta = [\theta_1 \quad \theta_2 \quad \dots \quad \theta_{nc}]^T.$$

Авторегресионен модел с външен входен сигнал (ARX)

Тъй като AR моделите, използвани за описание на времеви редове не считат допълнителни сигнали, оказващи влияние на изхода, те не могат да бъдат твърде точни. Ако има данни за други величини, влияещи изхода, за повишаване на качеството на модела те трябва да се включват в структурата му като известни външни въздействия.

Обикновено в динамичните системи има връзка между изхода в даден момент и неговите предишни стойности, като това е отчетено в (2.18). Освен тази авторегресионна зависимост реакцията на система-

та зависи и от предисторията на външния входен сигнал u_k . Така се получава ARX моделът, който има следния полиномен вид:

$$a(q^{-1})y_k = b(q^{-1})u_k + e_k. \quad (2.23)$$

Полиномът $a(q^{-1})$ е дефиниран в (2.19), а $b(q^{-1})$ е

$$b(q^{-1}) = 0 + b_1q^{-1} + \dots + b_{nb}q^{-nb}, \quad (2.24)$$

като nb е броят на предишните стойности на u_k , участващи във формирането на y_k . Има източници, в които (2.23) се нарича ARMA модел, за да се наблегне на факта, че y_k зависи от предишни стойности на u_k в рамките на пълзящ прозорец с дължина nb такта. В монографията, от друга страна, е прието (2.23) да се нарича ARX модел, защото от гледна точка на идентификацията има разлика между случая, когато изходът не зависи от външни входни сигнали и когато зависи от u_k , стойностите на който са наблюдавани в рамките на експеримент. Това разграничение е отчетено от Люнг в [119] и в днешно време се е установило като стандарт в идентификацията на системи.

Отново, както при AR модела, може да се осигури свободният член в (2.19) да бъде $a_0 = 1$. Параметърът на свободния член b_0 в (2.24) ще се приема за нула, като обяснението е следното. Ако моделът се използва за прогнозиране и наличните данни са до $k - 1$ -я такт, те се използват за прогнозиране на изхода в следващия k -ти такт. В такъв случай прогнозата $\hat{y}_k$ не зависи от u_k , тъй като тази стойност още не е налична. А ако данните до k -тия такт са налични (в това число и u_k), прогнозирането на y_k не е нужно, тъй като е настъпил моментът, в който изходът на системата се измерва.

Ако целта на модела е управление, то u_k зависи от стойността на y_k , тъй като последователността на действията в системата за управление е следната. Първо се измерва изходът y_k , а после, ако обратната връзка е по изхода, се формира разликата между заданието и измерената изходна стойност. Тази разлика е входен сигнал за регулатора, който формира управлението $u_k = f(y_k)$. Ето защо не е възможно и при тази постановка y_k да зависи от u_k .

И в двата описани случая е необходимо параметърът на свободния член да е $b_0 = 0$.

От друга страна, ако системата е статична, т.е. времето се изключи от разглежданията, параметрите $a_1, a_2, \dots, a_{na}$ и $b_1, b_2, \dots, b_{nb}$ са нула, т.е. изходът не зависи от предисторията на сигналите, а $b_0 \neq 0$, което

значи, че k -тата стойност y_k на реакцията зависи само от стойността на входната величина за същото k -то наблюдение.

Динамичният модел (2.23) може да се развие по следния начин:

$$y_k = a_1 y_{k-1} - \dots - a_{na} q^{-na} y_{k-na} + b_1 u_{k-1} + \dots + b_{nb} q^{-nb} u_{k-nb} + e_k,$$

като във векторна форма отново се достига до общия линейен вид на регресионен модел $y_k = \varphi_k^T \theta + e_k$. Тук векторите на регресорите и параметрите са:

$$\begin{aligned} \varphi_k &= [-y_{k-1} \quad -y_{k-2} \quad \dots \quad -y_{k-na} \quad u_{k-1} \quad u_{k-2} \quad \dots \quad u_{k-nb}]^T, \\ \theta &= [a_1 \quad a_2 \quad \dots \quad a_{na} \quad b_1 \quad b_2 \quad \dots \quad b_{nb}]^T. \end{aligned}$$

Авторегресионен модел с външен входен сигнал и формиращ филтър от тип пълзящо средно (ARMAX)

Понякога ARX структурата на модела не е подходяща за описание на реалната система. Това означава, че част от детерминираното поведение на обекта, което остава немоделирано, се причислява към остатъка. Когато остатъкът съдържа немоделирани аспекти от връзката между входа и изхода, той ще се означава с $\tilde{e}_k$. Един от начините за отчитане на детерминираната съставка в $\tilde{e}_k$ е тя да се опише с помощта на допълнителен филтър (наречен формиращ филтър). На входа му постъпва фиктивен бял шум e_k , а на изхода му се формира цветният шум $\tilde{e}_k$. Един вариант за описание на оцветеността на $\tilde{e}_k$ е в ARX структурата да се въведе филтър от тип МА, т.е.

$$\tilde{e}_k = c(q^{-1})e_k.$$

Регресионният модел, който се получава е ARMAX, има вида

$$a(q^{-1})y_k = b(q^{-1})u_k + c(q^{-1})e_k.$$

Полиномите $a(q^{-1})$ и $b(q^{-1})$ имат същия вид както в ARX модела, а полиномът $c(q^{-1})$ е дефиниран в (2.21).

ARMAX моделът, записан в общ вид (2.3), се получава при следните вектори на регресорите и параметрите:

$$\begin{aligned} \varphi_k &= [-y_{k-1} \quad -y_{k-2} \quad \dots \quad -y_{k-na} \quad u_{k-1} \quad u_{k-2} \quad \dots \quad u_{k-nb} \\ &\quad e_{k-1} \quad e_{k-2} \quad \dots \quad e_{k-nc}]^T, \\ \theta &= [a_1 \quad a_2 \quad \dots \quad a_{na} \quad b_1 \quad b_2 \quad \dots \quad b_{nb} \quad c_1 \quad c_2 \quad \dots \quad c_{nc}]^T. \end{aligned}$$

Съществуват и други типове линейни регресионни модели. Част от тях са споменати по-долу.

Авторегресионен модел с външен входен сигнал и формиращ филтър от тип авторегресия (ARARX)

$$a(q^{-1})y_k = b(q^{-1})u_k + \frac{1}{d(q^{-1})}e_k. \quad (2.25)$$

Авторегресионен модел с външен входен сигнал и формиращ филтър от тип авторегресия използващо средно (ARARMAX)

$$a(q^{-1})y_k = a(q^{-1})u_k + \frac{c(q^{-1})}{d(q^{-1})}e_k.$$

Тези и други модели са частни случаи на базовия регресионен модел, който е от вида:

$$a(q^{-1})y_k = \frac{b(q^{-1})}{f(q^{-1})}u_k + \frac{c(q^{-1})}{d(q^{-1})}e_k. \quad (2.26)$$

Типовите SISO описания не се разглеждат по-нататък в монографията, но са дадени, с цел да се представи идеята за полиномното представяне на линейните регресионни модели, както и някои от често срещаните типови описания. Затова, в следващите теми, съкращенията AR, ARX, ARMAX и т.н. ще се използват за означаване на многомерните варианти на тези модели.

ММО модели

Преди да се представят най-често срещаните типови ММО модели са дадени означенията, свързани с полиномните матрици, както и с векторните сигнали, които участват в долните описания.

Полиномна матрица

При прехода от SISO към ММО описания, полиномите в горните модели се заместват с полиномни матрици. Елементите на тези матрици са полиноми. Нека $M(q^{-1})$ е $r \times s$ -мерна реална полиномна матрица (с реални коефициенти), а елементите ѝ са полиноми по степените на аргумента q^{-1} , като максималният ред на полином в $M(q^{-1})$ е n . С други

думи тази полиномна матрици има вида:

$$M(q^{-1}) = \begin{bmatrix} m_{11}(q^{-1}) & m_{12}(q^{-1}) & \dots & m_{1s}(q^{-1}) \\ m_{21}(q^{-1}) & m_{22}(q^{-1}) & \dots & m_{2s}(q^{-1}) \\ \vdots & \vdots & \ddots & \vdots \\ m_{r1}(q^{-1}) & m_{r2}(q^{-1}) & \dots & m_{rs}(q^{-1}) \end{bmatrix}, \quad (2.27)$$

където $m_{ij}(q^{-1}) = m_{ij,0} + m_{ij,1}q^{-1} + \dots + m_{ij,n_{ij}}q^{-n_{ij}}$ е полиномът от i -тия ред и j -тия стълб, а n_{ij} е степента на полинома (тогава $n = \max(n_{11}, \dots, n_{rs})$). Ако полиномът $m_{ij}(q^{-1})$ е от по-нисък ред от n , то $m_{ij,h} = 0$ за $h = \overline{n_{ij} + 1, n}$. За такава матрица може да се използват съкратените записи $M(q^{-1}) \in \mathcal{R}_n^{r \times s}$ или $M \in \mathcal{R}_n^{r \times s}(q^{-1})$.

Освен записа (2.27) друго представяне на $M(q^{-1})$, което също се използва в изложението, е във вид на матричен полином:

$$M(q^{-1}) = M_0 + M_1 q^{-1} + \dots + M_n q^{-n},$$

където матриците $M_h \in \mathcal{R}^{r \times s}$, за $h = \overline{0, n}$ са с реални елементи $m_{ij,h}$. Ако $M(q^{-1})$ участва в регресионния модел, например за отчитане на влиянието на въздействията, то M_h съдържа всички параметри на модела, които отговарят на входните сигнали в $k - h$ -тия такт.

Многомерни сигнали

В монографията са приети следните означения на входно-изходните сигнали и на остатъка:

- $u_k \in \mathcal{R}^m$ е векторният входен сигнал в k -тия такт, т.е. моделът се състои от m входа и

$$u_k = [u_{1,k} \quad u_{2,k} \quad \dots \quad u_{m,k}]^T;$$

- $y_k \in \mathcal{R}^\ell$ е векторният изходен сигнал, т.е. моделът има ℓ изхода, обединени във вектора

$$y_k = [y_{1,k} \quad y_{2,k} \quad \dots \quad y_{\ell,k}]^T;$$

- $e_k \in \mathcal{R}^\ell$ е остатъкът, който в k -тия такт е

$$e_k = [e_{1,k} \quad e_{2,k} \quad \dots \quad e_{\ell,k}]^T.$$

Тъй като остатъчната неопределеност е представена като адитивен сигнал към изхода (виж точка 1.1.4.3), то e_k има същата размерност като тази на y_k .

За по-ясна интерпретация на изразите винаги векторите ще се дефинират като стълбове. Така например векторът x_k е стълб, а x_k^T е вектор ред.

По-долу са разгледани единствено многомерните ARX и ARMAX модели и представянето им в общ вид.

ARX модел

Може би най-често срещаният в практиката регресионен модел за описание на динамични системи е ARX. Неговият многомерен вариант е същият, както в SISO случая, с тази разлика, че размерностите са различни. Видът му е

$$A(q^{-1})y_k = B(q^{-1})u_k + e_k. \quad (2.28)$$

Полиномната матрица на авторегресионната част е $A(q^{-1}) \in \mathcal{R}_{na}^{\ell \times \ell}$, а матрицата, описваща влиянието на входа върху изхода, е $B(q^{-1}) \in \mathcal{R}_{nb}^{\ell \times m}$. В общия случай полиномите (елементи на матриците) са от различен ред. С na и nb са означени максималните степени на полиноми, съответно в $A(q^{-1})$ и $B(q^{-1})$. По аналогия със SISO динамичните модели, матриците:

$$A(q^{-1}) = \begin{bmatrix} a_{11}(q^{-1}) & \dots & a_{1\ell}(q^{-1}) \\ \vdots & \ddots & \vdots \\ a_{\ell 1}(q^{-1}) & \dots & a_{\ell \ell}(q^{-1}) \end{bmatrix}, \quad (2.29)$$

$$B(q^{-1}) = \begin{bmatrix} b_{11}(q^{-1}) & \dots & b_{1m}(q^{-1}) \\ \vdots & \ddots & \vdots \\ b_{\ell 1}(q^{-1}) & \dots & b_{\ell m}(q^{-1}) \end{bmatrix}, \quad (2.30)$$

може да се представят във вида:

$$A(q^{-1}) = I_\ell + A_1 q^{-1} + \dots + A_{na} q^{-na}, \quad (2.31)$$

$$B(q^{-1}) = 0_{\ell \times m} + B_1 q^{-1} + \dots + B_{nb} q^{-nb}, \quad (2.32)$$

където A_i и B_i са матрици, чиито елементи са коефициентите на полиномите, съответно в $A(q^{-1})$ и $B(q^{-1})$, които се умножават по q^{-i} . Смисълът на полиномите в матриците е следният. С текущия ij -ти полином $a_{ij}(q^{-1})$ в матрицата $A(q^{-1})$ се описва влиянието на j -тия изход върху i -тия изход, а полиномът $b_{ij}(q^{-1})$ отразява връзката между i -тия изход и предисторията на j -тия вход. Нека например с ARX модел да се описва поведението на система с три входа и два изхода. Тогава $A(q^{-1})$ трябва да е 2×2 полиномна матрица – така ще се отчете зависимостта на моментните стойности на двата изхода от предисторията на първия и втория изход. От друга страна, $B(q^{-1})$ трябва да е 2×3 полиномна матрица, за да има възможност да се отчете зависимостта на двата изхода от всеки от трите входа. Ако е известно, че някой от входовете не влияе на изход, то съответният полином в $B(q^{-1})$ предварително се приема за нулев. Например, ако u_3 не влияе на y_1 , то полиномът $[B(q^{-1})]_{13} = 0$.

Свободният член в (2.31) е $A_0 = I_\ell$, където I_ℓ е $\ell \times \ell$ мерна единична матрица. МИМО моделът (2.28) може да се разглежда като система от ℓ на брой МISO модела. Причината A_0 да е единична матрица, е, че (по аналогия със SISO моделите) това осигурява в i -тия подмодел ($i = \overline{1, \ell}$) наличието само на i -тия изход в k -тия момент (тъй като $A_0 y_k = y_k$). Това улеснява както получаването на модела, така и неговото използване. Ако по някакви причини $A_0 \neq I_\ell$, то с елементарни преобразувания, приложени към системата уравнения (2.28), тази матрица може да стане единична. Поради вида на A_0 няма други свободни членове в полиномните матрици (отговарящи на q^0), като всички стойности на входно-изходните сигнали в i -тия МISO модел, с изключение на $y_{i,k}$, са в предишни моменти от времето. Това е необходимо и за реализируемостта на модела. Например, ако в първия модел участва и $y_{2,k}$, а във втория $y_{1,k}$ (съответните извъндиагонални елементи на A_0 са ненулеви), МИМО моделът (2.28), в този му вид, няма да има физически смисъл.

Свободният член в (2.32) е $B_0 = 0_{\ell \times m}$ (нулева матрица с размерност $\ell \times m$). Това гарантира, че реакцията в k -тия такт не зависи от стойностите на въздействията в момента k (които все още са неизвестни).

Въпреки че ARX моделът може да се раздели на множество МISO модели, причината той да се разглежда като цяло, а не като множество независими описания, е възможността да се опише връзката на всеки от изходите както със собствената предистория, така и с предишни стойности на останалите входно-изходни сигнали. Също така разглеждането на (2.28) като цяло опростява моделирането, тъй като се построява един,

а не ℓ на брой подмодели. Основните причини за изграждането на ММО модел от гледна точка на приложението му са описани в точка 3.1.2.

В монографията е засегнато приложението на идентификацията и за статични модели, като в този случай матриците $A_1, A_2, \dots, A_{na}$ и $B_1, B_2, \dots, B_{nb}$ в (2.28) са нулеви, т.е. изходите не зависят от предисторията на входно-изходните величини, а свободният член е $B_0 \neq 0_{\ell \times m}$, което значи, че y_k зависи само от k -тите стойности на входните величини. Тогава в статичния случай ARX моделът се свежда до

$$y_k = B_0 u_k + e_k. \quad (2.33)$$

Двете общи форми (2.4) и (2.5) за представяне на линейни многомерни модели са причина за наличието на две реализации на методите за оценяване на параметри. Двата вида оценители имат различни особености (за които ще стане дума в Трета глава). По-долу ще се опишат конкретни варианти на прехода от (2.28) към (2.4) и към (2.5).

Преобразуването на ARX модела в общия вид (2.5) може да се извърши по аналогичен начин както за SISO моделите. Първо, с използване на полиномното представяне на матриците $A(q^{-1})$ и $B(q^{-1})$ за ARX модела се получава:

$$y_k + A_1 y_{k-1} + \dots + A_{na} y_{k-na} = B_1 u_{k-1} + \dots + B_{nb} u_{k-nb} + e_k.$$

След изразяване на y_k и обединяване на параметрите в матрица, а регресорите във вектор, моделът добива вида:

$$y_k = \Theta^T \varphi_k + e_k,$$

където

$$\Theta = [A_1 \ A_2 \ \dots \ A_{na} \ B_1 \ B_2 \ \dots \ B_{nb}]^T \in \mathcal{R}^{z \times \ell},$$

$$\varphi_k = [-y_{k-1}^T \ -y_{k-2}^T \ \dots \ -y_{k-na}^T \ u_{k-1}^T \ u_{k-2}^T \ \dots \ u_{k-nb}^T]^T \in \mathcal{R}^z.$$

Векторът φ_k съдържа na на брой изходни вектора с размерност ℓ и nb входни вектора с размерност m . Така за броя на регресорите се получава:

$$z = \ell \ na + m \ nb.$$

Съответно броят на параметрите е $p = z\ell$.

Възможно е матрицата Θ да се конструира по друг начин. Например, ако полиномите от даден стълб на полиномна матрица се разглеждат независимо един от друг, се получава описание, което има различни свойства.

Нека $A(q^{-1})$ и $B(q^{-1})$ се представят във вида:

$$A(q^{-1}) = [A_{.1}(q^{-1}) \ A_{.2}(q^{-1}) \ \dots \ A_{.\ell}(q^{-1})], \quad (2.34)$$

$$B(q^{-1}) = [B_{.1}(q^{-1}) \ B_{.2}(q^{-1}) \ \dots \ B_{.m}(q^{-1})], \quad (2.35)$$

където $A_{.i}(q^{-1})$ и $B_{.i}(q^{-1})$ са стълбовете на полиномните матрици. Нека също тези полиномни вектори се запишат във вид на векторни полиноми

$$A_{.i}(q^{-1}) = a_{i,0} + a_{i,1}q^{-1} + \dots + a_{i,na_i}q^{-1}, \quad (2.36)$$

$$B_{.i}(q^{-1}) = b_{i,0} + b_{i,1}q^{-1} + b_{i,2}q^{-1} + \dots + b_{i,nb_i}q^{-1}. \quad (2.37)$$

С na_i и nb_i са означени максималните степени на полиномите в i -тите стълбове. Елементите на вектора $a_{i,0}$, съдържащ свободните членове в $A(q^{-1})$, са нули с изключение на i -тия, който е равен на 1, а векторът $b_{i,0}$ от свободните членове на $B(q^{-1})$ е нулев вектор. Тогава ARX моделът може да се запише като

$$a_1(q^{-1})y_{1,k} + \dots + a_\ell(q^{-1})y_{\ell,k} = b_1(q^{-1})u_{1,k} + \dots + b_m(q^{-1})u_{m,k} + e_k.$$

Разписано по-подробно, последното равенство, в което векторът y_k е изразен, е

$$\begin{aligned} y_k = & -a_{1,1}y_{1,k-1} - \dots - a_{1,na_1}y_{1,k-na_1} \\ & -a_{2,1}y_{2,k-1} - \dots - a_{2,na_2}y_{2,k-na_2} - \dots \\ & -a_{\ell,1}y_{\ell,k-1} - \dots - a_{\ell,na_\ell}y_{\ell,k-na_\ell} \\ & +b_{1,1}u_{1,k-1} + \dots + b_{1,nb_1}u_{1,k-nb_1} \\ & +b_{2,1}u_{2,k-1} + \dots + b_{2,nb_2}u_{2,k-nb_2} - \dots \\ & +b_{m,1}u_{m,k-1} + \dots + b_{m,nb_m}u_{m,k-nb_m}. \end{aligned}$$

Векторите $a_{i,j}$ и $b_{i,j}$ съдържат коефициентите на полиномите от i -тия стълб на $A(q^{-1})$ и $B(q^{-1})$, които се умножават по q^{-j} .

След обединяване на параметрите в матрица, а регресорите във вектор, за модела отново се получава:

$$y_k = \Theta^T \varphi_k + e_k,$$

където

$$\Theta = [a_{1,1} \dots a_{1,na_1} \ \dots \ a_{\ell,1} \dots a_{\ell,na_\ell} \ b_{1,1} \dots b_{1,nb_1} \ \dots \ b_{m,1} \dots b_{m,nb_m}]^T,$$

$$\begin{aligned} \varphi_k = & [-y_{1,k-1} \dots -y_{1,k-na_1} \ \dots \ -y_{\ell,k-1} \dots -y_{\ell,k-na_\ell} \\ & u_{1,k-1} \dots u_{1,k-nb_1} \ \dots \ u_{m,k-1} \dots u_{m,k-nb_m}]^T. \end{aligned}$$

Матрицата Θ съдържа $na_1 + na_2 + \dots + na_\ell$ на брой вектора с размерност ℓ , съдържащи параметрите на полиномите в $A(q^{-1})$ и $nb_1 + nb_2 + \dots + nb_m$ вектора, отново с размерност ℓ , които съдържат параметрите на полиномите в $B(q^{-1})$. Така за броя на регресорите се получава:

$$z = \sum_{i=1}^{\ell} na_i + \sum_{i=1}^m nb_i.$$

Двата варианта на представяне, описани по-горе, имат различни свойства, а именно: в първия случай полиномите в дадена полиномна матрица не може да са от различен ред, докато във втория случай такава ограничение се налага единствено върху полиномите, принадлежащи на даден стълб.

Когато параметрите се групират във вектор (виж (2.4)), регресорите се обединяват в матрица. В този случай е по-удобно полиномните матрици да се представят във вида (2.29) и (2.30), а израз (2.28) да се разглежда като система от ℓ уравнения. Тогава i -тият MISO модел (i -тото уравнение) може да се запише по следния начин:

$$A_i(q^{-1})y_k = B_i(q^{-1})u_k + e_{i,k},$$

като $A_i(q^{-1})$ и $B_i(q^{-1})$ са полиномни вектори, отговарящи на i -тите редове на полиномните матрици. След умножението на векторните сигнали по полиномните вектори за горния израз се получава:

$$a_{i1}(q^{-1})y_{1,k} + \dots + a_{i\ell}(q^{-1})y_{\ell,k} = b_{i1}(q^{-1})u_{1,k} + \dots + b_{im}(q^{-1})u_{m,k} + e_{i,k}. \quad (2.38)$$

За опростяване на изложението индексът i , отговарящ на текущия модел, ще се изпусне. В такъв случай полиномът от i -тия ред и j -тия стълб на $A_i(q^{-1})$ е

$$a_j(q^{-1}) = a_{j,0} + a_{j,1}q^{-1} + \dots + a_{j,na_j}q^{-na_j}.$$

Свободният член на полинома е

$$a_{j,0} = \begin{cases} 1, & \text{за } i = j, \\ 0, & \text{за } i \neq j. \end{cases}$$

Полиномът $b_j(q^{-1})$, за $j = \overline{1, m}$ от текущия MISO модел е

$$b_j(q^{-1}) = 0 + b_{j,1}q^{-1} + \dots + b_{j,nb_j}q^{-nb_j}.$$

Тъй като $a_j(q^{-1})$ е полином от j -тия стълб на $A(q^{-1})$, той съдържа параметрите на модела, които се умножават по предишните стойности

на $y_{j,k}$ (за $j = \overline{1, \ell}$). Полиномът $b_j(q^{-1})$ от своя страна съдържа параметрите, които отговарят на предишните стойности на входа $u_{j,k}$ (за $j = \overline{1, m}$). При това положение нека събираемите в (2.38), които са:

$$a_j(q^{-1})y_{j,k} = a_{j,0}y_{j,k} + a_{j,1}y_{j,k-1} + \dots + a_{j,na_j}y_{j,k-na_j},$$

$$b_j(q^{-1})u_{j,k} = b_{j,1}u_{j,k-1} + \dots + b_{j,nb_j}u_{j,k-nb_j},$$

се представят във векторен вид, т.е.

$$a_j(q^{-1})y_{j,k} = \begin{cases} y_{i,k} + y_{j,k}^T a_j, & \text{за } i = j, \\ y_{j,k}^T a_j, & \text{за } i \neq j, \end{cases} \quad (2.39)$$

$$b_j(q^{-1})u_{j,k} = u_{j,k}^T b_j. \quad (2.40)$$

Векторите на параметрите $a_j \in \mathcal{R}^{na_j}$ и $b_j \in \mathcal{R}^{nb_j}$ са:

$$a_j = [a_{j,1} \ a_{j,2} \ \dots \ a_{j,na_j}]^T,$$

$$b_j = [b_{j,1} \ b_{j,2} \ \dots \ b_{j,nb_j}]^T,$$

а векторите на регресорите $y_{j,k} \in \mathcal{R}^{na_j}$ и $u_{j,k} \in \mathcal{R}^{nb_j}$ са:

$$y_{j,k} = [y_{j,k-1} \ y_{j,k-2} \ \dots \ y_{j,k-na_j}]^T,$$

$$u_{j,k} = [u_{j,k-1} \ u_{j,k-2} \ \dots \ u_{j,k-nb_j}]^T.$$

Както се вижда, регресорите в $y_{j,k}$ и $u_{j,k}$ са всъщност предисторията на сигналите до $k-1$ -я такт, а единственият сигнал в k -тия такт (виж (2.39)) в текущия MISO модел е $y_{i,k}$. Това позволява от (2.38) да се изрази $y_{i,k}$ като функция на предисторията на входовете и изходите, т.е.

$$y_{i,k} = -y_{1,k}^T a_1 - \dots - y_{\ell,k}^T a_\ell + u_{1,k}^T b_1 + \dots + u_{m,k}^T b_m + e_{i,k}$$

или накратко (с въвеждане отново на индекса i):

$$y_{i,k} = \Phi_{i,k}^T \theta_i + e_{i,k}. \quad (2.41)$$

Параметрите са групирани във вектора:

$$\theta_i = [\theta_{a_{i1}}^T \ \dots \ \theta_{a_{i\ell}}^T \ \theta_{b_{i1}}^T \ \dots \ \theta_{b_{im}}^T]^T \in \mathcal{R}^{p_i},$$

а регресорите са обединени във вектора:

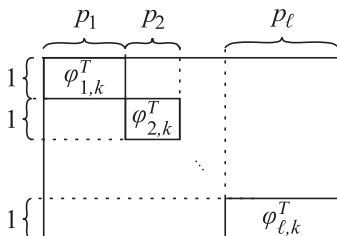
$$\Phi_{i,k} = [-y_{i1,k}^T \ \dots \ -y_{i\ell,k}^T \ u_{i1,k}^T \ \dots \ u_{im,k}^T]^T \in \mathcal{R}^{z_i}.$$

θ_i и $\varphi_{i,k}$ са с еднаква размерност, която е

$$p_i = z_i = \sum_j na_{ij} + \sum_j nb_{ij}.$$

Нека векторът θ съдържа параметрите от всички ММО модели. Той има размерност $p = \sum_i p_i$ и нека е със следната структура:

$$\theta = [\theta_1^T \ \theta_2^T \ \dots \ \theta_\ell^T]^T.$$



Фигура 2.2. Структура на матрицата на регресорите Φ_k на ARX модел в k -тия такт

Тъй като системата е с много изходи, а параметрите са обединени във вектор, необходимо е регресорите да са в матрица. Структурата на матрицата трябва да бъде такава, че при умножението ѝ с вектора θ да се получи вектор с ℓ компоненти, отговарящ на предсказания изход в текущия момент. За целта $\varphi_{i,k}$ се подреждат като редове по диагонала на $\Phi_k \in \mathcal{R}^{\ell \times p}$ (Фигура 2.2), т.е. получава се блокдиагоналната матрица:

$$\Phi_k = \text{diag}(\varphi_{1,k}^T, \varphi_{2,k}^T, \dots, \varphi_{\ell,k}^T).$$

Това е матрицата на регресорите на пълния ММО ARX модел и съдържа всички регресори, участващи във формирането на изхода в k -тия такт. При така дефинирани θ и Φ_k , общият вид на модел (2.28) е

$$y_k = \Phi_k \theta + e_k.$$

Възможни са и други записи на модела във вид с вектор на параметрите. Например, ако първите елементи на θ са всички параметри на матрицата $A(q^{-1})$, следвани от параметрите на $B(q^{-1})$, отново се получава описание с вектор на параметрите. Ако е запазена същата логика при нареждането на параметрите, то Φ_k ще се състои от две блокдиагонални матрици, съдържащи съответно изходните и входните регресори.

ARMAX модел

При изграждането на SISO модел входът и изходът не подлежат на проверка дали трябва да участват в описанието. От друга страна, при избора на структура на MIMO модел, освен уточняването на степените на полиноми и чистите закъснения, се извършва и проверка за значимост на входно-изходните величини, като част от тях може да отпаднат от модела. Освен това има вероятност значими, известни или достъпни за измерване сигнали да не са включени в набора от данни, който се използва за същинското моделиране. Това са предпоставки ARX структурата на модела да се окаже неподходяща за описание на реалната MIMO система.

Немоделираните, но детерминирани аспекти от входно-изходното поведение на системата се причисляват към остатъка и ако те не са пренебрежими, e_k не може да се разглежда като чисто случаен сигнал. Както преди в този случай остатъчната грешка ще се означава с $\tilde{e}_k$.

Ако не е намерен ARX модел с подходяща структура, може да се направи опит за построяване на ARMAX модел, т.е. към ARX структурата да се въведе MIMO филтър от тип MA по отношение на остатъка, т.е.

$$\tilde{e}_k = C(q^{-1})e_k. \quad (2.42)$$

Така за многомерния вид на ARMAX се получава:

$$A(q^{-1})y_k = B(q^{-1})u_k + C(q^{-1})e_k. \quad (2.43)$$

Полиномните матрици $A(q^{-1})$ и $B(q^{-1})$ са като тези в ARX модела, а матрицата на формация филтър е $C(q^{-1}) \in \mathcal{R}_{nc}^{\ell \times \ell}$. Тя може да се запише като матричен полином или в явен вид, съответно като:

$$C(q^{-1}) = I_\ell + C_1 q^{-1} + \dots + C_{nc} q^{-nc} = \begin{bmatrix} c_{11}(q^{-1}) & \dots & c_{1\ell}(q^{-1}) \\ \vdots & \ddots & \vdots \\ c_{\ell 1}(q^{-1}) & \dots & c_{\ell\ell}(q^{-1}) \end{bmatrix}.$$

С nc е означена максималната степен на полиномите в $C(q^{-1})$. Причината за $C_0 = I_\ell$ е същата като при матрицата $A(q^{-1})$, т.е. само в i -тия MISO модел участва $e_{i,k}$. Това осигурява в (2.43) участието на y_k , а не на изчисления изход на модела $\hat{y}_k$ (както е обяснено в точка 2.1.3.1, за целите на идентификацията в описанието е необходимо да присъства y_k , а не $\hat{y}_k$).

Нека с представянето на матриците $A(q^{-1})$, $B(q^{-1})$ и $C(q^{-1})$ като матрични полиноми, ARMAX моделът се запише във вида:

$$y_k + A_1 y_{k-1} + \dots + A_{na} y_{k-na} = B_1 u_{k-1} + \dots + B_{nb} u_{k-nb} + C_1 e_{k-1} + \dots + C_{nc} e_{k-nc} + e_k.$$

Както при ARX модела, за да се достигне до общото описание (2.5), параметрите трябва да се групират в матрица, а регресорите във вектор. След изразяване на y_k и отделяне на параметрите от регресорите се получава:

$$y_k = \Theta^T \varphi_k + e_k.$$

Един вариант е $\varphi_k \in \mathcal{R}^z$ и $\Theta \in \mathcal{R}^{z \times \ell}$ да са:

$$\varphi_k = [-y_{k-1}^T \quad \dots \quad -y_{k-na}^T \quad u_{k-1}^T \quad \dots \quad u_{k-nb}^T \quad e_{k-1}^T \quad \dots \quad e_{k-nc}^T]^T,$$

$$\Theta = [A_1 \quad \dots \quad A_{na} \quad B_1 \quad \dots \quad B_{nb} \quad C_1 \quad \dots \quad C_{nc}]^T,$$

където $z = \ell(na + nc) + m \quad nb$, а $p = z\ell$.

Преобразуването на ARMAX във вида с вектор на параметрите (2.4) може да се извърши по подобен начин на описаното преобразуване на ARX модела. Първо от (2.43) се отделя i -тият ред (MISO модел):

$$A_{i.}(q^{-1})y_k = B_{i.}(q^{-1})u_k + C_{i.}(q^{-1})e_k.$$

Представен по-подробно, моделът е

$$a_{i1}(q^{-1})y_{1,k} + \dots + a_{i\ell}(q^{-1})y_{\ell,k} = b_{i1}(q^{-1})u_{1,k} + \dots + b_{im}(q^{-1})u_{m,k} + c_{i1}(q^{-1})e_{1,k} + \dots + c_{i\ell}(q^{-1})e_{\ell,k}. \quad (2.44)$$

Полиномите $a_{ij}(q^{-1})$ и $b_{ij}(q^{-1})$ са същите като при ARX модела.

По-долу индексът i се изпуска. Във векторен запис, последните ℓ събираеми в (2.44) са:

$$c_j(q^{-1})y_{j,k} = \begin{cases} e_{i,k} + e_{j,k}^T c_j, & \text{за } i = j, \\ e_{j,k}^T c_j, & \text{за } i \neq j, \end{cases}$$

а векторите на регресорите и параметрите са:

$$e_{j,k} = [e_{j,k-1} \quad e_{j,k-2} \quad \dots \quad e_{j,k-nc_j}]^T \in \mathcal{R}^{nc_j},$$

$$c_j = [c_{j,1} \quad c_{j,2} \quad \dots \quad c_{j,nc_j}]^T \in \mathcal{R}^{nc_j}.$$

Това позволява от (2.44) да се изрази $y_{i,k}$ като функция на предистори-

ята на входно-изходните сигнали и на остатъка в k -тия такт, т.е.

$y_{i,k} = -y_{1,k}^T a_1 - \dots - y_{\ell,k}^T a_\ell + u_{1,k}^T b_1 + \dots + u_{m,k}^T b_m + e_{1,k}^T c_1 + \dots + e_{\ell,k}^T c_\ell + e_{i,k}$
или накратко (с въвеждане отново на индекса i):

$$y_{i,k} = \varphi_{\text{ARX}i,k}^T \theta_{\text{ARX}i} + \varphi_{\text{fi},k}^T \theta_{\text{fi}} + e_{i,k}. \quad (2.45)$$

Параметрите на i -тия MISO модел са групирани така:

$$\theta_{\text{ARX}i} = [a_{i1}^T \quad \dots \quad a_{i\ell}^T \quad b_{i1}^T \quad \dots \quad b_{im}^T]^T,$$

$$\theta_{\text{fi}} = [c_{i1}^T \quad \dots \quad c_{i\ell}^T]^T,$$

а регресорите му са обединени в двата вектора:

$$\varphi_{\text{ARX}i,k} = [-y_{i1,k}^T \quad \dots \quad -y_{i\ell,k}^T \quad u_{i1,k}^T \quad \dots \quad u_{im,k}^T]^T,$$

$$\varphi_{\text{fi},k} = [e_{i1,k}^T \quad \dots \quad e_{i\ell,k}^T]^T.$$

Моделите (2.45), записани заедно, са:

$$y_k = \Phi_k \theta + e_k,$$

където векторът θ съдържа параметрите на всички MISO модели и е със следната структура:

$$\theta = [\theta_{\text{ARX}}^T \quad \theta_{\text{f}}^T]^T,$$

като

$$\theta_{\text{ARX}} = [\theta_{\text{ARX}1}^T \quad \theta_{\text{ARX}2}^T \quad \dots \quad \theta_{\text{ARX}\ell}^T]^T$$

е вектор от параметрите на матриците $A(q^{-1})$ и $B(q^{-1})$ (ARX частта на разширения модел). Векторът

$$\theta_{\text{f}} = [\theta_{\text{f}1}^T \quad \theta_{\text{f}2}^T \quad \dots \quad \theta_{\text{f}\ell}^T]$$

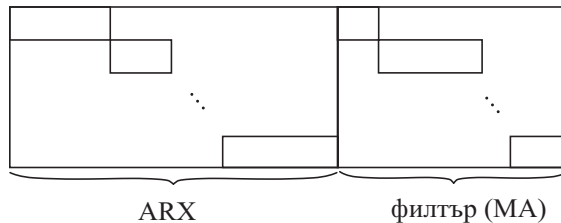
съдържа параметрите на полиномите в матрицата $C(q^{-1})$ (частта на филтъра от тип MA). Тогава матрицата на регресорите на пълния ARMAX модел $\Phi_k \in \mathcal{R}^{\ell \times p}$ е

$$\Phi_k = [\Phi_{\text{ARX},k} \quad \Phi_{\text{f},k}],$$

(структурата ѝ е показана на Фигура 2.3), а $\Phi_{\text{ARX},k}$ и $\Phi_{\text{f},k}$ са:

$$\Phi_{\text{ARX},k} = \text{diag}(\varphi_{\text{ARX}1,k}^T, \varphi_{\text{ARX}2,k}^T, \dots, \varphi_{\text{ARX}\ell,k}^T),$$

$$\Phi_{\text{f},k} = \text{diag}(\varphi_{\text{f}1,k}^T, \varphi_{\text{f}2,k}^T, \dots, \varphi_{\text{f}\ell,k}^T).$$



Фигура 2.3. Структура на матрицата на регресорите Φ_k на ARMAX модел в k -тия такт

Както за ARX модел, така и за ARMAX може да се формират и други представяния с вектор на параметрите. Например възможно е първите елементи на θ да са всички параметри, участващи при описанието на $y_{1,k}$, следвани от всички параметри, свързани с $y_{2,k}$, и т.н. Тогава Φ_k ще е блокдиагонална. Структурата на θ и Φ_k е удачно да се избира така, че да се улесни по-нататъшната работа с модела. Отделянето на параметрите на филтъра от останалите параметри, извършено по-горе, е удобно при извеждането на метода на разширените най-малки квадрати, разгледан в Трета глава.

2.1.3.4 Сравнение на представянето с матрица и с вектор на параметрите

Двете възможни групи представяния не са еквивалентни. Освен това свойствата им може да се променят и в рамките на съответната група, както беше показано за двата варианта за ARX модел в общ вид с матрица на параметрите. За първото представяне важи ограничението – всички полиноми в дадена полиномна матрица да са от един и същи ред, докато при второто представяне ограничението е само върху полиномите от даден стълб. Въпреки че второто представяне води до повече възможности за избор на структурата, конструирането на матрица Θ винаги води до известни ограничения в структурата на регресионния модел.

При представянията на модела с вектор θ промяната на реда на полином, участващ при описанието на даден изход, води единствено до изменение на съответния вектор на параметрите, а също и на свързания с него вектор на регресорите. Тази промяна не оказва влияние на останалата част от модела.

2.1.4 Декомпозиции на MIMO моделите

В точка 2.1.3 са дадени примери за различни регресионни модели, както и възможността някои от тях да се представят в общ линейно параметризиран вид. Освен стремежът описанията да са линейни по параметри, понякога за опростяване на моделирането, на синтеза на регулатор, на анализа на системата и др. многомерните модели се декомпонират на подмодели. Първо е описана декомпозиция на първоначалния MIMO модел на подмодели (в общия случай отново MIMO), а след това и преобразуването на MIMO модел в множество от SISO подмодели.

2.1.4.1 Множество от MIMO или MISO подмодели

Понякога за опростяване задачата на идентификацията, вместо да се търси един модел на многомерната система, тя се разглежда като множество от подсистеми и така, вместо един модел с голяма размерност, се търси множество от модели с по-малка размерност [5, 17].

При тези опростявания на задачата е важно идентификацията на всяка подсистема до голяма степен да може да протече независимо от моделирането на останалите подсистеми. Това е необходима стъпка, особено когато броят на изходите е значителен. В [120, 17] подробно е разглеждана декомпозицията на MISO подсистеми, с което задачата се опростява значително. В [80] е извършена декомпозиция на MIMO система с голяма размерност на MIMO подсистеми.

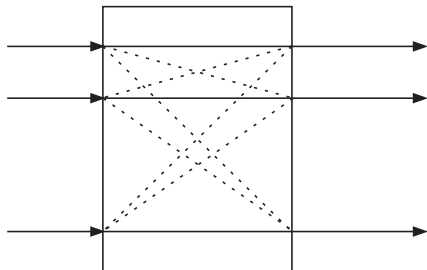
При моделирането на хипермаркет (точка 1.1.3.2), където ℓ е от порядъка на 10^5 , а m е около 10^7 , е невъзможно да се построи директно един MIMO модел, който да е достоверен, тъй като в него биха участвали много лъжливи връзки (дължащи се на случайната съставка в данните). Такива грешки са често срещани, когато твърде много се разчита на статистическите методи, без да се търси логично обяснение на структурата на модела, т.е. без да се използва априорната информация за системата.

В случая с хипермаркета, ако той се разглежда като цяло, то дори и при статичен модел, за всеки от изходите трябва да се проведат (явно автоматизирани) тестове за значимост на всеки един от десетките милиони потенциални фактори. Тези тестове са статистически и зависят от данни, често с малка размерност, като брой наблюдения (тактът на дискретизация е седмица) и с високо ниво на неопределеност. Затова споменатата опасност от грешни изводи за структурата на модела е значителна. Единственото решение в този случай е системата да се декомпозира на подсистеми, като се използват априорни знания за очаквани

зависимости. Така потенциалните връзки в подмоделите се ограничават до подмножество от логични възможности.

2.1.4.2 Множество от SISO подмодели

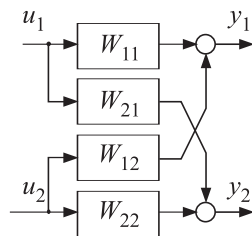
Описаната по-долу декомпозиция се извършва след процеса на моделиране, но е полезна, когато има изискване описанието да се състои от едномерни модели. Такъв е случаят, когато се изгражда система за управление и се разполага с едномерни регулатори. Предимство на този подход е, че многомерните модели се свеждат до едномерни и в този случай директно се прилагат методите за анализ и синтез на SISO системи.



Фигура 2.4. Сепаратни (плътна линия) и кръстосани (пунктирна линия) канали в многомерна система

В днешно време обаче са достъпни многомерни управляващи устройства и често ограниченото им използване се дължи на наследената практика да се използва теорията на едномерните системи въпреки предимствата от използването на МИМО моделите, които са естественият начин за представяне на МИМО системите.

За да се декомпозират многомерните модели на множество от SISO подмодели, се въвеждат компенсиращи звена. Това става, като първо, част от връзките в МИМО структурата се избират за основни (сепаратни), а останалите се приемат за кръстосани, както е показано на Фигура 2.4. Съществуват различни начини за определяне на сепаратните връзки. Например може да се определи чувствителност-



Фигура 2.5. Двусвързан модел с предавателни функции по каналите W_{11} , W_{12} , W_{21} и W_{22} .

та на всеки изход от входовете на модела и за сепаратни канали да се избераат тези, които свързват изходите с входовете, към които са най-чувствителни. В много случаи това разделение е условно.

Добавянето на компенсатори, които да намалят (доколкото е възможно) влиянието на връзките, които са приети за кръстосани (и оттам за нежелани), усложнява цялостното описание на системата. Компенсаторите се определят на базата на параметрите на модела. Тогава, при неточен модел или възникване на нестационарност, качеството на компенсация на кръстосаните влияния се влошава.

Пример. *Компенсирание на кръстосани връзки в двусвързана система*

Нека дадена система се описва с ARX модел:

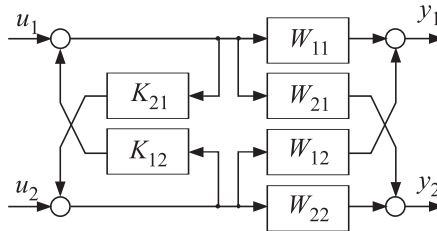
$$A(q^{-1})y_k = B(q^{-1})u_k + e_k.$$

Нека също, за опростяване на примера, $a_{12}(q^{-1}) = a_{21}(q^{-1}) = 0$, т.е. няма кръстосани връзки между изходите, а само от входовете към изходите (Фигура 2.5). Нека също предварително са определени сепаратните и кръстосаните канали, като първият сепаратен канал свързва u_1 и y_1 , а вторият сепаратен канал – u_2 и y_2 . Съответно кръстосаните връзки, които трябва да се неутрализират, са от u_1 към y_2 и от u_2 към y_1 . ARX моделът може да се запише като система от два модела (тъй като толкова са изходите):

$$a_{11}(q^{-1})y_{1,k} = b_{11}(q^{-1})u_{1,k} + b_{12}(q^{-1})u_{2,k} + e_{1,k},$$

$$a_{22}(q^{-1})y_{2,k} = b_{21}(q^{-1})u_{1,k} + b_{22}(q^{-1})u_{2,k} + e_{2,k}.$$

В представената на Фигура 2.6 система са добавени компенсаторите



Фигура 2.6. Разширен двусвързан модел с компенсатори K_{12} и K_{21}

K_{12} и K_{21} . За разширения вид на модела е в сила:

$$a_{11}(q^{-1})y_{1,k} = b_{11}(q^{-1})(u_{1,k} + K_{12}u_{2,k}) + b_{12}(q^{-1})u_{2,k} + e_{1,k},$$

$$a_{22}(q^{-1})y_{2,k} = b_{21}(q^{-1})u_{1,k} + b_{22}(q^{-1})(u_{2,k} + K_{21}u_{1,k}) + e_{2,k}.$$

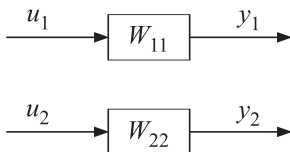
Компенсаторите трябва да удовлетворяват условията:

$$b_{11}(q^{-1})K_{12}u_{2,k} + b_{12}(q^{-1})u_{2,k} = 0, \quad (2.46)$$

$$b_{21}(q^{-1})u_{1,k} + b_{22}(q^{-1})K_{21}u_{1,k} = 0, \quad (2.47)$$

откъдето за K_{12} и K_{21} се получава:

$$K_{21} = -\frac{b_{21}(q^{-1})}{b_{22}(q^{-1})}, \quad K_{12} = -\frac{b_{12}(q^{-1})}{b_{11}(q^{-1})}.$$



Фигура 2.7. При пълно компенсиране на кръстосаните връзки се получават два независими SISO модела.

Ако моделът напълно отразява влиянието на входните сигнали, то кръстосаните връзки се компенсират напълно и разширеният модел се свежда до два независими SISO модела (Фигура 2.7)

$$a_{11}(q^{-1})y_{1,k} = b_{11}(q^{-1})u_{1,k} + e_{1,k},$$

$$a_{22}(q^{-1})y_{2,k} = b_{22}(q^{-1})u_{2,k} + e_{2,k}.$$

(При друга структура на модела K_{12} и K_{21} ще имат друг вид.) ◇

Когато системата е двусвързана, задачата за преход към SISO под-модели не е сложна. Ако системата е с $m = 3$ входа, $\ell = 3$ изхода и всички връзки са значими, в общия случай са нужни 6 компенсатора, при $m = \ell = 4$ нужните компенсатори са 12 и т.н., като с нарастване на броя им се усложнява тяхното извеждане и предавателните им функции.

С изключение на случаите, когато броят на входно-изходните сигнали е много малък, декомпозирането на ММО описанието на SISO модели е практически неприложимо.

Основното предимство на ММО моделите е, че вместо стремежът да се потиснат някои от характерните особености в поведението на многомерната система, тези особености може да се използват за съответните

цели, за които е построен моделът. Например, ако има възможност да се въздейства по различни канали върху важен (от гледна точка на управлението) изход на обекта, тази възможност би трябвало да се използва. Ето защо, ако системата има изразен многомерен характер, тя трябва да се опише и интерпретира по естествения за това начин, а именно с помощта на многомерен модел.

2.2 Експеримент, предварителна обработка и анализ на данни

В първите два етапа на идентификацията се уточнява системата и се избира класът на бъдещия модел. Понякога е възможно да се уточни и неговият тип, евентуално структурата му, а при добре изучени процеси може аналитично да се определят и първоначални стойности на параметрите на модела. Следващият етап е свързан с провеждането на експеримент и подготовката на получените данни за целите на същинското моделиране. Експериментът е необходим за събирането на апостериорна информация (съдържаща се в данните) за обекта. Съществуват два вида експерименти – активни и пасивни. При активните, входните въздействия се формират целенасочено за получаване на информативни данни, а при пасивните данните се събират без активно въздействие върху системата. Понякога е възможно да се планира подходящ експеримент, който да осигури информативни данни, при това – в работната област и средата, в която функционира обектът. Това е най-благоприятният вариант. В много случаи обаче не може (или не е желателно) да се провеждат активни експерименти. Тогава няма гаранция, че данните са достатъчно информативни. Например може да има дълги периоди, в които наблюденията остават постоянни, наличие на нехарактерни стойности, кратки интервали на наблюдение и др. Също така липсата на стойности за наблюдения води до непълни набори от данни, а това поражда проблеми в следващите етапи, особено при многомерните системи, тъй като е необходимо наличието на всички входно-изходни наблюдения за разглеждания интервал. Понякога невъзможността да се проведе подходящ експеримент (или недостатъчното априорно знание) се компенсира с натрупването на голям обем от данни, от който впоследствие се извлича полезна информация за системата.

При некачествени набори от данни обработката и анализът им преди същинското изграждане на модела са от решаващо значение за ре-

зултата от идентификацията. Затова основно внимание е отделено на дейностите след получаването на сурови наблюдения.

Резултатът от този етап е набор/набори от данни, които са подготвени за следващия етап, а именно – построяването на модела. Понякога допълнителен резултат от анализа на експерименталните данни може да е доизясняването на някои особености на системата и в резултат на това моделът да се конкретизира, така че да отразява тези особености.

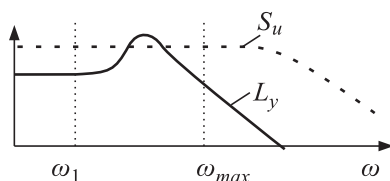
2.2.1 Планиране на експеримент

Планирането на експеримент е обширна тема, която не е цялостно засегната. Например не се обсъждат теми като пълен, дробен, композиционен, ротатабълен, D-оптимален и др. факторен експеримент. Темата за избор на план на експеримента е подробно разгледана в [2, 29]. От друга страна, в монографията са описани някои важни аспекти от планирането на експеримента, които имат отношение към обработката и анализа на получените данни.

2.2.1.1 Входно въздействие

При планирането на експеримента е важно да се избере подходящ входен сигнал, който да осигури информативни данни за изхода на системата. С примера, разгледан по-долу, се показва, че не всеки входен сигнал е подходящ за целите на идентификацията.

Пример. *Избор на входно въздействие*

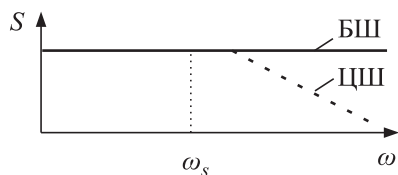


Фигура 2.8. ЛАЧХ на обекта и спектр на u_k (пунктирна линия)

Нека изследваният обект в работната област да може да се апроксимира с линеен модел с ЛАЧХ (логаритмичната амплитудночестотна характеристика) $L_y(\omega)$, дадена на Фигура 2.8. Също така целта на задачата на идентификацията е да се намери модел, описващ обекта в честотния диапазон до ω_{max} – честотата, при която усилването на систе-

мата спада с 3db. (В примера познаването на ω_{max} е всъщност априорна информация за обекта.)

Ако е проведено изследване със сигнал, който в рамките на експеримента е $u(t) = \text{const}$ (с t е означено времето), то входно-изходните данни по никакъв начин няма да отразяват влиянието на входния сигнал върху изхода. В този смисъл за целите на идентификацията е нужно $u(t)$ да се изменя и по този начин да „възбуди“ изхода, пораждайки реакция. Нека входният сигнал е $u(t) = \sin(\omega_1 t)$. Естествено е да се очаква обектът да реагира на това променливо въздействие. Известно е, че при подаване на синусоиден сигнал реакцията на линейна система е синусоида, и то със същата честота като входната. Това означава, че наборът от данни ще съдържа информация за поведението на обекта само за честотата ω_1 . В такъв случай моделът ще има поведение, близко до това на обекта само за тази честотата, а не за целия работен диапазон. С други думи и синусоидалният сигнал не е подходящ за идентификацията, когато се търси модел в определен честотен диапазон. ◇

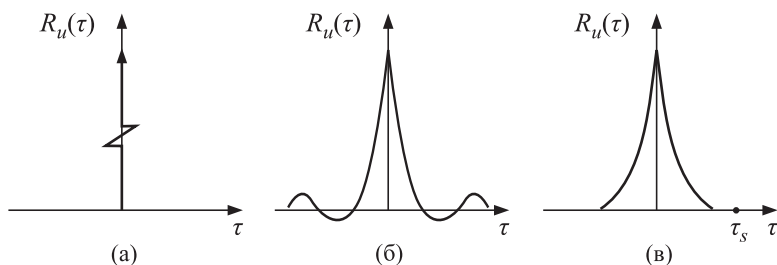


Фигура 2.9. Спектър на непрекъснат бял и цветен шум. Ако ω_s е честотата на дискретизация, то дискретизираният цветен шум е дискретен бял шум.

Условието, което трябва да удовлетворява u_k , за да бъде подходящ за идентификацията, е спектърът му $S_u(\omega)$ да обхваща областта, в която се търси прецизен модел (Фигура 2.8 – пунктирна линия). Това изискване се нарича условие за достатъчна възбудимост на обекта. Често в практиката, за да се изпълни то, u_k се формира като (известен) случаен сигнал с достатъчно широк спектър.

Теоретично идеален сигнал за идентификация е непрекъснатият бял шум (БШ), тъй като спектърът му е константа за всички честоти (Фигура 2.9) [24]. От гледна точка на времевата област БШ е сигнал, стойността на който в даден момент не зависи от предисторията му. Но енергията, необходима за генериране на такъв сигнал, е безкрайно голяма. (Слънцето може да се разглежда като източник на БШ). От друга страна, използването на дискретни данни е свързано с изучаването на сигнали с ограничен спектър. По тези причини често в идентифика-

цията като входно въздействие се използва дискретен бял шум (ДБШ). Неговият спектър (Фигура 2.9) е ограничен и по тази причина той може да се реализира.



Фигура 2.10. Корелационни функции на БШ и ЦШ. Условие за ДБШ.

Условието даден непрекъснат сигнал след дискретизация да е ДБШ, е корелационната му функция да схожда към 0 в рамките на избрания такт на дискретизация. Цветният шум (ЦШ), от друга страна, е случаен сигнал, стойността на който в даден момент зависи, освен от случайна съставка, и от предисторията му. Ако връзката с предисторията му е в рамките на определен период, по-малък от такта на дискретизация, то този (напълно реализируем) ЦШ, може да се разглежда като ДБШ, чийто дискретни стойности са независими. На Фигура 2.10 са дадени корелационните функции на БШ и ЦШ, както и случаят, при който ЦШ удовлетворява условието за ДБШ.

2.2.1.2 Такт на дискретизация

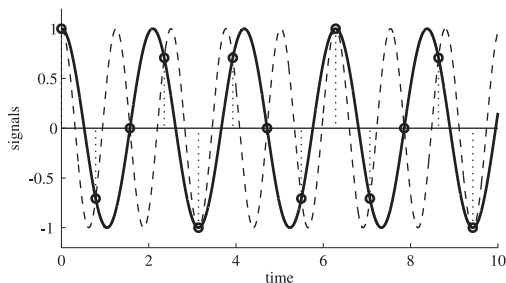
Ако системата е непрекъсната, за да се снимат данни за поведението ѝ, трябва да се избере подходящ период на дискретизация. Понякога има изисквания периодът да е над определена стойност (например поради времето за отчитане на изходите, формирането на управляващи въздействия, оценяването на състояния, параметри и т.н.). От друга страна, ако периодът е твърде голям, при дискретизацията се губи информация за високочестотната съставка на сигналите, която може да е нужна. Затова изборът на период на дискретизация е важна стъпка в идентификацията. От теоремата на Котелников-Шенон е известно, че ако аналогов сигнал се дискретизира с честота, поне два пъти по-висока от най-високата честота в спектъра на сигнала, тогава първоначалният сигнал може да се възстанови напълно от сметите данни. В съвременната формулировка на теоремата честотата на дискретизация трябва

да надвишава два пъти най-високата честота в спектъра на аналоговия сигнал. Така се гарантира, че няма да възникне случай, при който дискретните отчети съвпадат с моментите, когато хармоникът с най-висока честота пресича абсцисата – в този случай той не може да се възстанови. Освен това първоначалният вариант на теоремата е в сила за безкрайна последователност от дискретни наблюдения. На практика при работа с крайни извадки се предпочита честотата на дискретизация да е още по-висока.

Освен загубата на информация, ако се наруши теоремата на Котелников-Шенон, може да възникне и т.нар. шум от квантуване (aliasing). Той се появява, когато честотата на дискретизация е по-ниска от допустимата (според теоремата). Резултатът е погрешно отчитане на високите честоти на оригиналния непрекъснат сигнал като ниски честоти в дискретизирания.

По-долу е разгледан пример, с който се описва шумът от квантуване. В него се използват следните означения: t е времето, измерено в секунди, T_s е периодът на дискретизация (времето между два отчета) и също се измерва в секунди, а $f_s = \frac{1}{T_s}$ е честотата на дискретизация в херци. Понякога е удобно вместо с f_s да се работи с ъгловата честота ω_s , която се измерва в радиани за секунда (тя е подходяща, когато сигналът се разглежда като сума от хармоници). Връзката между двете честоти е $\omega_s = 2\pi f_s$. Нека максималната честота в спектъра на непрекъснат сигнал е ω_{max} . Тогава горната теорема може да се запише като:

$$\omega_s \geq 2\omega_{max}. \quad (2.48)$$



Фигура 2.11. Загуба на информация при дискретизация. Двата сигнала след дискретизация са идентични.

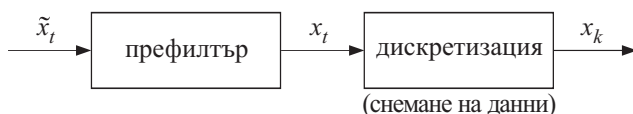
Пример. *Шум от квантуване*

Нека непрекъснатите хармонични сигнали:

$$x_1(t) = \sin(3t) \text{ и } x_2(t) = \sin(5t)$$

, се дискретизират с честота $\omega_s = 8 \text{ rad/s}$. За първия (по-бавно изменящ се) сигнал $x_1(t)$ теоремата на Котелников-Шенон е в сила, защото $\omega_s > 2\omega_1 (= 6)$, но за втория сигнал тя се нарушава, тъй като $\omega_s < 2\omega_2 (= 10)$. На Фигура 2.11 са представени сигналите $x_1(t)$ и $x_2(t)$, както и дискретните стойности, отчетени с честота ω_s . Вижда се, че отчетите за двата сигнала съвпадат. От една страна, това означава, че при дискретизацията на $x_2(t)$ се загубва информация и сигналът не може да се възстанови от измерените данни. От друга страна, при преобразуване на дискретизирания вариант на $x_2(t)$ отново в непрекъснат вид грешно възстановеният сигнал е с по-ниска честота (и съвпада с $x_1(t)$). Така, при неспазване на (2.48), в спектъра на дискретизираните сигнали се наслагват нискочестотни хармоници, дължащи се на наличието на високочестотни хармоници в оригиналните сигнали. $\diamond$

Отстраняването на псевдочестотите в спектъра на дискретизираните сигнали (anti-aliasing) може да се постигне с увеличаване на честотата на дискретизация. Едно практическо правило [148] е честотата на дискретизация да се избере $\omega_s = 10\omega_{max}$.



Фигура 2.12. Филтриране на непрекъснатия сигнал, с цел потискане на шума от квантуване

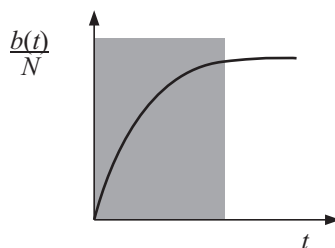
Понякога в сигналите се съдържа високочестотна съставка, която е извън честотната лента, в която се търси модел. Обикновено (тъй като е извън очакваната област на полезната част от сигнала), тази съставка се дължи на шума от измерване, на смущения от околната среда и други фактори, които се приемат за случайни. За да се гарантира, че няма да възникне шум от квантуване, обикновено преди снемането на данните сигналите се филтрират (Фигура 2.12), така че да се потисне тази нежелана съставка. По този начин, вместо да се дискретизира оригиналният сигнал $\tilde{x}_t$, при който би възникнал шум от квантуване, данните се снемат за филтрирания сигнал x_t .

Ако честотният диапазон на наблюдаваните величини не е известен, за да се подsigури достатъчна информативност на набора от данни, е подходящо да се снимат наблюдения с висока честота, а след това, ако е необходимо, ω_s да се намали в процеса на моделиране, като сигналите се филтрират с подходящ цифров филтър (виж точка 2.2.2.1).

2.2.1.3 Интервал на наблюдение

Въпреки че обикновено интервалът на наблюдение, в който е проведен експериментът и броят на наблюденията са директно свързани, има случаи, когато връзката не е толкова явна. Основното при избора на интервала са особеностите на системата, а при определянето на броя на наблюденията акцентът е върху необходимостта от (статистически) достатъчно данни.

Интервалът на наблюдение може да се повлияе от наличието на цикличност в поведението или от необходимостта някои аспекти на системата да се проявят. Сезонността е основна съставка в процесите, протичащи в пазарните системи. Във финансите също се наблюдава сезонност и често това е причина интервалът на наблюдение да включва поне една година. Дори когато се разполага с данни за повече предишни години, поради нестационарният характер в поведението на кредитополучателите и на средата, твърде дълъг интервал не е препоръчителен. Също така в кредитната индустрия се говори за т.нар. „узряване“ на лошите кредитополучатели (Фигура 2.13). В началото на погасителния период болшинството от клиентите на банката изплащат заемите си, но по-рисковите от тях след известно време срещат затруднения и поради просрочия в плащанията попадат в графата на лошите платци. Така, за да се идентифицират лошите клиенти, е нужно време, което на фигурата е означено в сиво.



Фигура 2.13. „Узряване“ на лошите клиенти: $b(t)$ – брой лоши клиенти, N – общ брой клиенти

При изучаването на динамични системи на базата на априорна информация или с формиране на груб модел може да се определи най-голямата времеконстанта T_{max} на обекта, а след това да се използва правилото: продължителността на експеримента да е 10 пъти по-голяма от T_{max} . Естествено, за конкретния случай трябва да се вземе предвид

и степента на неопределеност в данните. Тя има отношение и към броя на наблюденията и затова е обсъдена в следващата подточка.

2.2.1.4 Брой наблюдения

За получаване на достоверен модел е необходимо да се натрупат достатъчно данни. Нека p е броят на търсените параметри. При отсъствие на неопределеност в задачата на идентификацията (всички наблюдения са напълно достоверни), както ще стане ясно в Трета глава, параметрите може да се определят, ако броят на наблюденията N е такъв, че да е възможно да се запише моделът за p различни случая. Ако системата е статична, са необходими $N = p$ наблюдения, а за динамична система, в която е необходима предистория за n предишни стойности при формирането на изхода на модела, са нужни $N = p + n$ наблюдения. В практиката обаче данните съдържат неопределеност и поради това за осигуряване на модел, слабо чувствителен към случайните съставки в данните, е нужно $N \gg p$. Правилото е, колкото по-високо е нивото на неопределеността в данните, толкова повече наблюдения да се използват.

Оказва се обаче, че твърде големи стойности на N не само че не водят до подобрение в качеството на модела, но може да доведат до ненужно забавяне (и следователно до намаляване на ефективността) на процеса на идентификация.

Ако към набор от няколко наблюдения се добави едно допълнително наблюдение, то би довело до чувствително изменение на оценките. От друга страна, ако към няколкостотин наблюдения се добави едно наблюдение, неговото влияние върху оценките е значително по-слабо.

В много случаи броят на наблюденията е ограничен от особеностите на системата. Например натрупването на наблюдения за поведението на ректификационна колона [148] може да продължи седмица. Този случай изглежда оптимистичен на фона на пазарните системи, където данните се отчитат с период една седмица. В кредитната индустрия този период е месец (тъй като кредитите се изплащат месечно).

Проблемът с големият такт на дискретизация в последните две области е решен, като данните се събират от много източници. Има вериги с хиляди магазини. При тях в рамките на седмица се натрупват хиляди реализации на всяка една от входно-изходните величини (например от всеки магазин постъпват данни за продажбите на круши). След подходящо агрегиране на седмичните стойности от всички магазини може да се определи по една обобщена стойност за всяка от наблюдаваните

величини (компонентите на u_k и y_k). Тъй като тези стойности са получени в резултат от осредняване на наблюденията от магазините се очаква неопределеността в тях (например в средните продажби на круши за k -тата седмица) да е значително по-малка от тази, която се съдържа в данните от конкретен магазин.

Във финансите, за да се натрупат достатъчно данни, основно се разчита не на продължителността на пасивния експеримент, а на големия брой клиенти. Данните за кредитоискателите в тази област се получават еднократно при кандидатстването им за конкретна услуга, като тук k е индексът на текущия кандидат. Затова от решаващо значение е броят N на клиентите да е достатъчно голям.

Често в стремежа да се осигурят достатъчно данни, размерността на наборите може да стане твърде голяма. При ограничено време за изпълнение на задачата на идентификацията, ако наборът е твърде голям, моделирането може да не доведе до търсения резултат. Решението в такива случаи е наборът от данни да се редуцира. Ако описанието, което се търси, е статично, тогава най-често се избира поднабор от наблюдения на случаен принцип. Този подход невинаги е удачен. Ако някои особености на системата се наблюдават в значително по-малък брой наблюдения, а други особености са често срещани в данните, желателно е да се премахнат повече наблюдения от втората група, а от първата, наблюденията да се запазят.

Обикновено наборите от данни за кредитополучатели съдържат малко лоши платци и затова е желателно броят им да не се намалява драстично. От друга страна, изключването на част от преобладаващите добри платци не би влошило качеството на модела, но би осигурило достатъчно време за протичането на процеса на идентификация.

Това селективно премахване на наблюдения може да доведе до грешни заключения, тъй като съотношението между добрите и лошите клиенти в данните се променя, а това опорочава някои анализи и изводи, които зависят от първоначалното съотношение на групите клиенти в данните.

За да се запази първоначалната картина, се въвеждат тегла, които се асоциират с всяка от специфичните групи наблюдения. Ако в даден набор броят на добрите клиенти се редуцира 5 пъти, а броят на лошите 2 пъти, то на всяка двойка $\{u_k, y_k\}$, съответстваща на добър клиент, се асоциира тегло $w_k = 5$, а на двойка, отговаряща на лош кандидат, теглото е $w_k = 2$. Впоследствие всеки добър (лош) кандидат се интерпретира като 5 (2) кандидати, които имат еднакви характеристики.

До момента основно е обсъден случаят, когато данните са много, което затруднява процеса на моделиране. Но има много случаи, когато данните от експеримента не са достатъчни и се налага натрупването на още данни. Също е възможно още при планирането да е избран набор от експерименти, които да се проведат. При това положение може да е целесъобразно отделните набори от данни да се обединят. Обикновено директното групиране на наборите не е желателно. Това е в сила особено при динамичните системи, където динамиката в даден момент зависи от предисторията на сигналите. Затова е важно да се отдели специално внимание на началните и крайните условия на границите между наборите.

2.2.2 Предварителна обработка на данни

Предварителната обработка на данните е много важна дейност. В някои случаи правилно обработените данни са ключът към получаването на достоверен модел. От една страна, целта на манипулирането на данните е неопределеността в тях да се потисне, доколкото е възможно, а полезната информация да се запази. Други дейности имат за цел да се опрости същинското оценяване на параметри, например като се измени видът на модела или се избегнат числени проблеми.

Първоначално се започва с дейности по „почистване“ на данните, като: пренебрегване на нехарактерни стойности или на цели сигнали, данните за които са неинформативни или се наблюдава мултиколинearност над допустимо ниво; попълване на липсващи стойности и др. След това се пристъпва към обработка, която включва нормализация, нелинейна трансформация, декомпозиция на сигналите и др.

2.2.2.1 Намаляване честотата на дискретизация

По-горе беше споменато, че ако честотният диапазон, в който се търси модел, не е известен, е подходящо първоначално честотата на дискретизация да е висока. След това, в зависимост от спектъра на изходните сигнали и натрупването на допълнителни познания за процесите, може да се определи честотният диапазон, в който се търси описание на обекта. Тогава наличните данни се обработват с цифров филтър, за да се избегне шумът от квантуване (виж точка 2.2.1.2). Така подготвените данни се обработват с цел намаляване на честотата на дискретизация ω_s . Използването на префилтъра може да доведе до въвеждане на закъснение, което трябва да се отчете при обработката на данните.

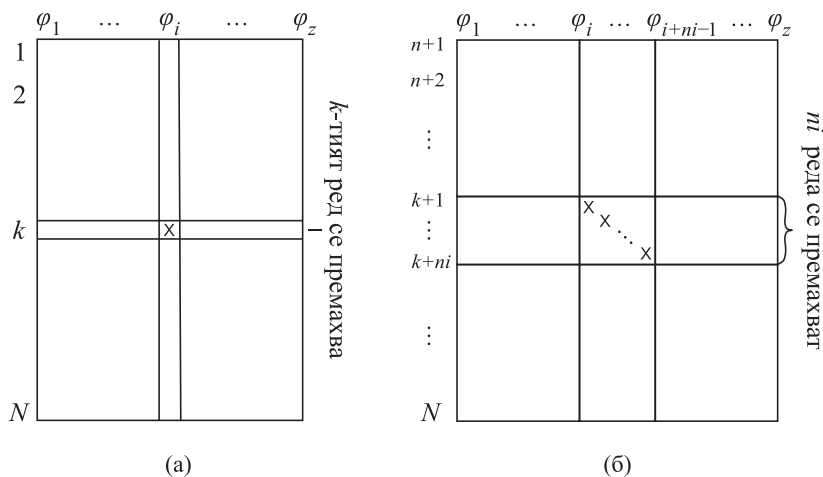
Правилното прилагане на тази дейност води до потискане на високочестотните смущения, като същевременно се запазва полезната съставка в данните.

2.2.2.2 Липсващи стойности

Някои причини за наличието на липсващи (а и на нехарактерни) стойности са грешки при снемане на наблюденията, неправилно агрегиране на данни за едни и същи величини (продажби на определен продукт в магазините от дадена верига), грешки при извличането на данните и др. В [112] са дискутирани някои източници на липсващи и нехарактерни стойности.

Когато наборът от данни не е пълен, т.е. има липсващи стойности, той не може директно да се използва за изграждане на модел. Най-общо съществуват два подхода за справяне с този случай. Това са:

- пренебрегване на наблюдения;
- попълване на липсващите стойности.



Фигура 2.14. Почистване на Φ от липсващи стойности за статичен (а) и динамичен модел (б). Редът на полинома на сигнала с липсваща стойност за случай (б) е ni .

Преди да се разгледат те, първо ще се въведе т.нар. матрица на данните (на факторите), която се използва при оценяването на параметри. Тя е показана на Фигура 2.14, като структурата ѝ във втория случай отговаря на представянето на модела в общ вид с матрица на параметрите.

За статичен модел тя е

$$\Phi = [\varphi_1 \ \varphi_2 \ \dots \ \varphi_N]^T \equiv [u_1 \ u_2 \ \dots \ u_N]^T. \quad (2.49)$$

За динамичен модел, при който най-старата стойност на фактор, влияещ на изхода в k -тия такт, отговаря на момент $k - n$, матрицата на данните е със следната структура:

$$\Phi = [\varphi_{n+1} \ \varphi_{n+2} \ \dots \ \varphi_N]^T. \quad (2.50)$$

Причината първият вектор на регресорите да е с индекс $n + 1$, е, че с първите n налични стойности на регресорите се формира именно φ_{n+1} (а за вектори на регресорите в предишни моменти няма необходимите данни).

При първия подход данните от моментите, в които част от стойностите липсват, не участват по време на моделирането. Когато системата е статична, проблематичните двойки наблюдения $\{u_k, y_k\}$ директно се премахват. От гледна точка на Φ това означава да се премахнат съответните редове, както е показано на Фигура 2.14(а)).

От друга страна, ако се моделира динамиката, то липсата на стойност в даден момент води до необходимостта от пренебрегване на данните от толкова следващи дискретни моменти, в колкото тази стойност участва при формирането на изхода (Фигура 2.14 (б)). Този подход е подходящ при достатъчен брой пълни наблюдения (без липсващи стойности), така че пренебрегването на част от данните статистически да не повлияе на крайния модел. Освен това, за да се определят моментите, в които дадена липсваща стойност участва при формирането на изхода, е необходимо да е известна структурата на модела. Със следващия пример е представен този подход.

Пример. *Липсващи стойности и пренебрегване на наблюдения при динамичен модел*

Нека е избрана следната структура на (MISO) модела:

$$\hat{y}_k = \theta_1 u_{1,k-1} + \theta_2 u_{1,k-2} + \theta_3 u_{1,k-3} + \theta_4 u_{2,k-1},$$

т.е. $\varphi_k = [u_{1,k-1} \ u_{1,k-2} \ u_{1,k-3} \ u_{2,k-1}]^T$. Ако първият входен сигнал в k -тия момент не е известен, то изходът на модела няма как да бъде определен в моментите $k+1$, $k+2$ и $k+3$, тъй като тези предсказани стойности на изхода зависят от $u_{1,k}$. С други думи векторите на регресорите φ_{k+1} , φ_{k+2} и φ_{k+3} , които съдържат $u_{1,k}$, няма как да се формират.

За разгледаните моменти от време изходът на модела е

$$\hat{y}_{k+1} = \theta_1 \underline{u_{1,k}} + \theta_2 u_{1,k-1} + \theta_3 u_{1,k-2} + \theta_4 u_{2,k},$$

$$\hat{y}_{k+2} = \theta_1 u_{1,k+1} + \theta_2 \underline{u_{1,k}} + \theta_3 u_{1,k-1} + \theta_4 u_{2,k+1},$$

$$\hat{y}_{k+3} = \theta_1 u_{1,k+2} + \theta_2 u_{1,k+1} + \theta_3 \underline{u_{1,k}} + \theta_4 u_{2,k+2}.$$

Това означава, че ако се пропускат случаите с липсващи наблюдения, то трябва три стойности на изхода на модела (и свързаните с тях вектори на регресорите) да не участват в моделирането.

Първата стойност, след k -тата, за която има достатъчно предистория (т.е. векторът на регресорите е известен), е $\hat{y}_{k+4}$. Така матрицата на данните (от Фигура 2.14 (б)) ще бъде:

$$\Phi = [\varphi_4 \quad \dots \quad \varphi_k \quad \varphi_{k+4} \quad \dots \quad \varphi_N]^T \in \mathcal{R}^{N-3-3 \times 4}.$$

Ако липсва стойността на втория входен сигнал в k -тия момент, то изходът на модела няма как да се изчисли само в момента $k+1$, тъй като зависи от $u_{2,k}$. Изходът на модела в $k+1$ -я момент е

$$\hat{y}_{k+1} = \theta_1 u_{1,k} + \theta_2 u_{1,k-1} + \theta_3 u_{1,k-2} + \theta_4 \underline{u_{2,k}}.$$

Това означава, че изходът и свързаните с него фактори трябва да се пропуснат. Първата стойност след k -тата, за която има достатъчно предистория, е $\hat{y}_{k+2}$. За матрицата на данните се получава:

$$\Phi = [\varphi_4 \quad \dots \quad \varphi_k \quad \varphi_{k+2} \quad \dots \quad \varphi_N]^T \in \mathcal{R}^{N-3-1 \times 4}.$$

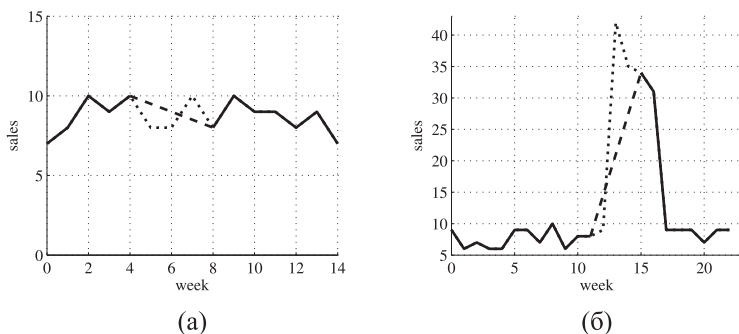
◇

Както се вижда от този пример, при динамичните модели пренебрегването на наблюдения е по-сложно, особено при системи с голяма размерност, тъй като трябва да се отчитат степените на полиномите, отговарящи на броя на предишните стойности във вектора на регресорите по отношение на всеки един канал в модела.

Другият подход за отчитане на случаите с непълен набор от данни е липсващите отчети да се попълнят с подходящи стойности. Най-простият вариант е да се определи по една базова стойност за всеки сигнал, например средната стойност, и тя да се използва вместо липсващите стойности. Друг вариант е данните да се интерполират. Тази техника е подходяща, когато не се наблюдават голям брой съседни липсващи стойности. Също трябва да се внимава със случаите, когато поведението на обекта се изменя значително, както е показано в следния пример.

Пример. *Липсващи стойности и интерполация на сигнали в пазарна система*

На Фигура 2.15 са представени продажбите на продукт, които се характеризират с пикове, дължащи се на прилагането на промоции за кратки периоди от време (примерно до четири седмици). Ако липсват стойности по време на непромоционален период (случай 1), то би могло те да се заместят със средните продажби, характерни за периода, или да се определят на базата на линейна интерполация, с което няма да се въведе голяма грешка. Но ако липсващите стойности са на границата на промоционален период (случай 2), то може интерполираните стойности значително да се отличават от действителните (неизвестни) стойности на продажбите.



Фигура 2.15. Интерполация (пунктир) на продажбите (плътна линия) за попълване на липсващи стойности (линия от точки): непромоционален (а) и на граница с промоционален (б) случай



Интерполацията е приложима както за изходни, така и за входни сигнали. Друг вариант за заместване на липсващите стойности на изхода, особено ако се наблюдават продължителни периоди с последователни липсващи стойности, е да се формират т.нар. „бързи и неточни“ модели, с които да се оценят, макар и грубо, неизвестните стойности. В горния пример с такъв предварителен, макар и не много точен, модел може да се отчете моментът, в който се прилага промоцията, и липсващата стойност да се замени с подходяща, която да отговаря на въздействията (в случая продуктът вече е на промоция и това е известно от цената му и евентуално от флаг за вида на промоцията). Пример за използването на

помощни модели в областта на медицината е отчитането на пациентите с непотвърдена диагноза, както е описано в долния пример.

Пример. *Липсващи стойности за изхода и използване на помощен модел в медицината*

При поставяне на диагноза на базата на данни от изследвания (входове за модела) пациентите се оценяват с помощта на модел. По различни причини този модел може да се окаже неточен и при необходимост той трябва да се подмени с нов. Диагнозата (изход на модела) на част от пациентите не е потвърдена, тъй като те не са продължили с изследвания и/или лечение. Като цяло има зависимост между вида на диагнозата и споменатата реакция на индивидите. Затова тези пациенти не може да се разглеждат като случайно подбрани. Тогава, ако заради липсата на изходни данни съответните двойки наблюдения $\{u_k, y_k\}$ се пренебрегнат, то новият модел не би описвал коректно бъдещи пациенти със симптоми и резултати от изследвания, сходни с тези с непотвърдени диагнози в миналото (а такива пациенти се очаква и за в бъдеще да има).

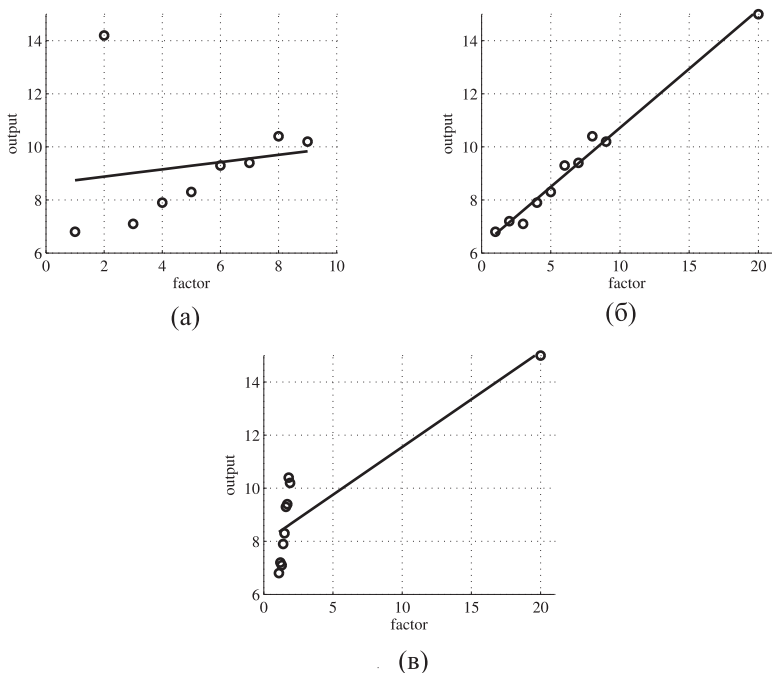
Възможно решение е да се моделират пациентите, за чието поведение има данни. Полученият (неточен) модел отразява особеностите само на тази част от пациентите. Поради различното поведение на двете групи индивиди моделът допълнително трябва да се измени, така че да бъде логичен от гледна точка на предварителните знания за поведението на пациентите. Тази вторична настройка може да се извърши ръчно, като изследователят използва натрупания практически опит, или автоматично, при предварително зададени правила. Пригоденият към особеностите на дискутираната група помощен модел се прилага към известните входни данни, с цел да се оценят липсващите стойности на изхода. Всъщност този помощен модел дава приблизителен отговор на въпроса каква е вероятността пациент с непотвърдена диагноза да е болен. $\diamond$

2.2.2.3 Нехарактерни стойности

Наблюдения, които не са типични, може да се дължат, освен на споменатите причини в началото на предната тема, и на рядкото проявяване на някои особености на системата може би поради неподходящите условия, при които е проведен експериментът. Например, ако твърде рядко даден продукт е на промоция, големият брой продажби може да се разглежда като нехарактерно поведение. Друг пример е решението на човек с много висок доход да изтегли малък потребителски кредит.

Тогава нехарактерна стойност би могла да бъде размерът на дохода му (отличаващ се значително от доходите на хората, прибягващи до тази услуга на банката).

Смущенията от околната среда също могат да породят нехарактерни пикове в данните. Примери за това са ефектите от някои специфични кампании на конкурентни търговски вериги върху продажбите.



Фигура 2.16. Данни с нехарактерна стойност и модел. Нехарактерен изход (а), нехарактерен фактор и изход, но в направлението (б) и извън направлението (в) на регресионната линия.

Типично за нехарактерните стойности е, че се срещат рядко в данните. Това означава, че дори да се дължат на особеност на системата, няма натрупана статистика за това поведение, т.е. не може със статистически методи да се изгради надеждно описание на тази особеност. От друга страна, наличието на пикови стойности може да измести модела от средната тенденция в поведението на системата.

На Фигура 2.16 са дадени три варианта на данни с нехарактерна стойност и линеен SISO модел, построен с тези данни. В случай (а) стойността на изхода е нетипична. Това наблюдение представлява отдалечена от общата тенденция точка (outlier) в пространството на данните и за конкретния пример води до повдигане на регресионната линия.

Във втория случай особената точка е отдалечена от останалите стойности и на фактора, и на изхода, но лежи в оптималното направление на линията и затова не я променя. Нежеланият ефект от това наблюдение е, че то изкуствено намалява стандартната грешка на параметрите (виж точка 2.4.4.1), което се отразява на изводите от анализа на качеството на модела. Казано по друг начин, данните не са достатъчни, за да се приеме, че моделът е достоверен и за такава нехарактерна стойност на фактора.

Случай (в) отразява влиянието на точка, която е отдалечена едновременно спрямо оптималната тенденция в данните и спрямо типичните стойности на фактора. Този тип наблюдения имат значителен ефект върху качеството на модела. По тази причина наблюденията с нехарактерни стойности на факторите (варианти (б) и (в)) се наричат още лостови точки (leverage points) [110].

Правилното действие при наличие на нехарактерни наблюдения е те да се премахнат. Съществуват различни начини за намирането и коририрането им. В някои случаи се налагат ограничения върху стойностите на сигналите в зависимост от априорната информация за областта на реалистичните данни. При известни ограничения на скоростта на изменение на даден сигнал може да се следи разликата между неговите съседни стойности и т.н. Друг вариант е свързан с дефиниране на областта на реалистичните стойности с използване на наличните данни. Тъй като тази дейност може да се приложи, както към входните, така и към изходните данни, по-долу се използва общото означение x_k за сигнала, който се обработва. С помощта на филтър може да се определи нискочестотната съставка $x_{t,k}$ (трендът) на x_k . За тази цел например в [147] се използва филтър на Бътърфърд [20]. Полученият нискочестотен сигнал, заедно със стандартното отклонение на сигнала, от който трендът е премахнат ($\tilde{x}_k = x_k - x_{t,k}$) и евентуалната информация за разпределението му, може да се използва за определяне на долната и горна граници на изменение на x_k :

$$x_{l,k} = x_{t,k} - n\sigma_{\tilde{x}} \quad \text{и} \quad x_{u,k} = x_{t,k} + n\sigma_{\tilde{x}}. \quad (2.51)$$

Тези граници не са постоянни, както е показано на Фигура 2.17. Стойностите на сигнала извън границите се приемат за нехарактерни, а заместването им с подходящи стойности може да се извърши например с интерполация. Когато се използва линейна интерполация, се гарантира, че всички коригирани стойности попадат в характерната област на изменение.

Параметърът n в (2.51) зависи от доверителния интервал, в който стойностите на x_k се приемат за нормални. Информация, която може да подпомогне избора на n , е разпределението на $\tilde{x}_k$. Например, ако то е нормално, стойностите на x_k попадат в интервала $[x_{l,k}, x_{u,k}]$ с вероятност 0.997 [4], когато $n = 3$. Ако има стойности извън тази област, те се заместват с подходящи – в примера на Фигура 2.17, като за тази цел се използва линейна интерполация. По този начин всички коригирани стойности принадлежат на приетата за нормална област.

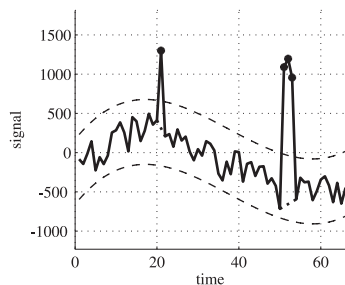
Опростен вариант на този метод е вместо нискочестотен филтър, да се използва средната стойност на x_k , която вместо $x_{t,k}$ участва при формирането на $x_{l,k}$ и $x_{u,k}$.

Дейностите, разгледани до момента, са свързани с първоначалното почистване на данните. В следващите точки са дадени техники, ориентирани към допълнително изменение на сигналите, с цел опростяване на моделирането и подобряване на резултата от него.

2.2.2.4 Нормализация

Характерно за много системи е наблюдаваните величини да се изменят в различни граници. Това, особено при многомерните системи, може да доведе до значителни грешки, както в оценяването на параметрите, така и при анализа на получения модел. В резултат на това, ако не са предприети специални мерки за числено устойчиви изчисления, има вероятност въобще да не се стигне до задоволителен модел.

Затова, след провеждане на експеримент, е удачно данните да се нормализират. Целта на нормализацията е входно-изходните сигнали



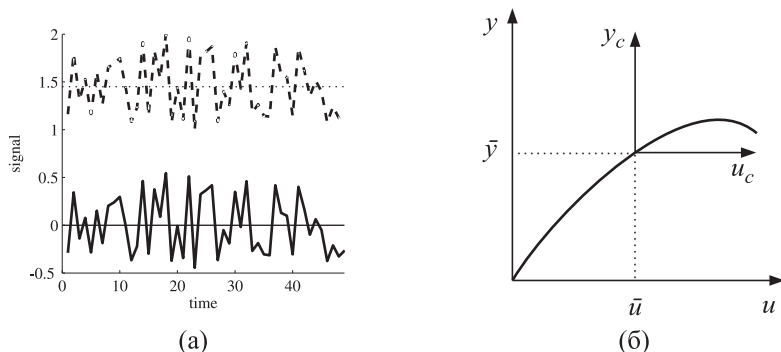
Фигура 2.17. Отчитане на нехарактерни стойности с нискочестотен филтър и попълване с линейна интерполация

да се уеднаквят в някакъв смисъл, така че да се избегне този нежелан ефект. По-долу са представени двете най-често срещани нормализации и са обсъдени ползите от тях.

Стандартизация

Една от най-разпространените нормализации е т.нар. стандартизация [147]. При нея сигналите се центрират и мащабират, както е описано по-долу. Тъй като тази дейност е приложима към входовете и към изходите, по-долу отново се използва общото означение x_k за обработвания сигнал.

Центриране



Фигура 2.18. Центриране на сигнал (а); центрирането на u и y е придвижване на началото на координатната система в работната точка (б)

Нека x_k е скаларен сигнал. Центрирането му се свежда до определяне на средната стойност

$$\bar{x} = \frac{1}{N} \sum_{k=1}^N x_k$$

за интервала на наблюдение, а след това до изместването на x_k , така че новата последователност $x_{c,k}$ да се изменя около нулата (Фигура 2.18 (а)), т.е.

$$x_{c,k} = x_k - \bar{x}.$$

Ако u_k и y_k се центрират, това означава, че координатната система, в която се описва статичната характеристика на системата, се придвижва така, че центърът на новата ѝ позиция да съвпадне с работната точка (Фигура 2.18 (б)), в околността на която е проведен експериментът.

Тази дейност е подходяща, когато средните стойности на сигналите се отличават значително една от друга.

Пример. Необходимост от центриране

Нека входовете и изходите се изменят в интервалите

$$u_k \in [2 - 10^{-2}, 2 + 10^{-2}], \quad y_k \in [10^4 - 10^2, 10^4 + 10^2]$$

и нека бъде избран оценител на параметри по метода на най-малките квадрати. Тогава матрицата на Фишер (наричана още информационна матрица) $F = \Phi^T \Phi$, която се използва при изчисляване на оценките (виж (2.49) и (2.50)), се състои от много малки и много големи стойности, а това може да доведе до влошаване на обусловеността ѝ. Важно е F да е добре обусловена, тъй като при работата на оценителя тя се обръща. (Числените проблеми са обсъдени подробно в точка 3.4 от Трета глава.) След центриране се получават сигнали в интервалите:

$$u_{c,k} \in [-10^{-2}, 10^{-2}], \quad y_{c,k} \in [-10^2, 10^2].$$

Ако те се използват за определяне на F , вероятността от числени проблеми намалява. Въпреки това дисперсията на изхода е значително по-голяма от тази на входа, което впоследствие може да доведе до големи разлики в стойностите на оценките. $\diamond$

За допълнително подобрене на данните от гледна точка на оценяването е желателно те да се мащабират.

Мащабиране

При мащабирането се определя стандартното отклонение на сигнала:

$$\sigma_x = \sqrt{\frac{1}{N-1} \sum_{k=1}^N x_{c,k}^2}.$$

То се използва за мащабиране на центрираните стойности, като

$$x_{s,k} = \frac{x_{c,k}}{\sigma_x}.$$

Така се получава стандартизиран сигнал $x_{s,k}$, който е центриран и има стандартно отклонение 1. Ако разпределението на x_k е нормално, след стандартизацията (разпределението не се променя) се получава т.нар. стандартно нормално разпределение, наричано още разпределение на Лаплас [6].

Ако u_k и y_k се мащабират, това означава, че мащабите в координатната система на статичната характеристика се уеднаквяват. На Фигура

2.19 (б) е показан този ефект, като за удобство е прието, че входът и изходът са центрирани.

Ако сигналът е векторен (нека $x_k \in \mathcal{R}^n$), то стандартизирането му може да се запише по следния начин:

$$x_{s,k} = S_x^{-1}(x_k - \bar{x}). \quad (2.52)$$

Векторът $\bar{x} \in \mathcal{R}^n$ и матрицата $S_x \in \mathcal{R}^{n \times n}$ са:

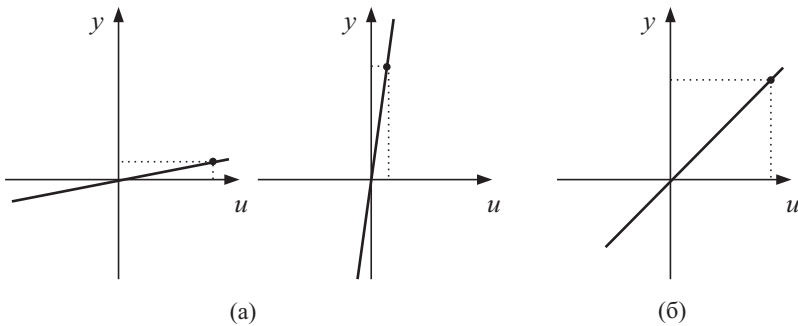
$$\bar{x} = [\bar{x}_1 \quad \bar{x}_2 \quad \dots \quad \bar{x}_n]^T, \quad (2.53)$$

$$S_x = \text{diag}[\sigma_{x,1} \quad \sigma_{x,2} \quad \dots \quad \sigma_{x,n}]. \quad (2.54)$$

Ефектът от стандартизацията на данните от гледна точка на оценяването на параметри е представен със следния пример.

Пример. Необходимост от стандартизация

Нека регресорите в даден MISO модел $\varphi_{1,k}$ и $\varphi_{2,k}$ са еднакво информативни и независими. В такъв случай, ако за средните им стойности е в сила $\bar{\varphi}_1 \gg \bar{\varphi}_2$, това би довело (както стана дума в предишния пример) до много разнородни елементи на матрица F и до евентуално влошаване на нейната обусловеност, а оттам и до нарастване на числените грешки при определянето на оценките на параметрите. Ако в модела няма свободен член, различните средни стойности на двата фактора водят също и до по-малка стойност на оценката $\hat{\theta}_1$ на параметъра, свързан с $\varphi_{1,k}$, в сравнение със стойността на оценката $\hat{\theta}_2$ на параметъра, отговарящ на $\varphi_{2,k}$. Причината за $\hat{\theta}_1 \ll \hat{\theta}_2$ е, че за да се изравни влиянието на



Фигура 2.19. Ефект от мащабирането. Статични характеристики на немащабирани (а) и мащабирани (б) сигнали u_k и y_k .

регресорите, така че големите стойности на $\varphi_{1,k}$ да не доминират при формирането на изхода, те се умножават с по-малката стойност $\hat{\theta}_1$. Когато в модела има постоянен член θ_0 , стойностите на параметрите не са толкова разнородни (това е показано по-долу при дестандартизацията). Въпреки това с θ_0 не се гарантира добра обусловеност на матрицата F .

От друга страна, ако стандартните отклонения на $\varphi_{1,k}$ и $\varphi_{2,k}$ са $\sigma_{\varphi,1} \gg \sigma_{\varphi,2}$, това също води до влошаване на обусловеността на F и до нуждата от значително по-малка стойност на параметъра, свързан с първия фактор. Причината отново е в различната област на изменение на регресорите.

Но след стандартизирането, поради това че първоначалните фактори допринасят с еднакво количество информация за описание на изхода, информационната матрица е добре обусловена и оценките на параметрите приемат стойности в много по-близки граници. $\diamond$

Освен численото подобряване на резултата от оценяването, когато данните са стандартизирани, съществува директна връзка между големината на параметрите и значимостта на съответните величини (за повече подробности виж точка 2.4.4.1). Заради тази връзка, при анализа на получения модел много изследователи използват не оригиналния, а стандартизирания модел (получен със стандартизирани данни).

Дестандартизация

Нека моделът е линеен и се използва описание (2.5) с матрица на параметрите (тогава факторите са подредени във вектор) и нека регресорите са стандартизирани. В резултат на тази обработка и след оценяване на параметрите се получава стандартизираният модел:

$$y_{s,k} = \hat{\Theta}_s^T \varphi_{s,k} + e_{s,k}. \quad (2.55)$$

Този модел не може директно да се използва с оригиналните данни, тъй като за неговото получаване се използват изменени сигнали. Затова, след оценяването на параметрите, моделът трябва да се дестандартизира, или с други думи да се преобразува във вид, подходящ за работа с първоначалните величини.

Нека със z' се означаи броят на факторите в стандартизирания модел, с векторите $\bar{\varphi} \in \mathcal{R}^{z'}$ и $\bar{y} \in \mathcal{R}^{\ell}$ (по подобие на (2.53)) – съответно средните стойности на регресорите и изходите, а с $S_{\varphi} \in \mathcal{R}^{z' \times z'}$ и $S_y \in \mathcal{R}^{\ell \times \ell}$ – диагоналните матрици, съдържащи (по подобие на (2.54)), съответните стандартни отклонения. Тогава стандартизираният вектор на регресо-

рите и този на изхода може да се запишат като:

$$\varphi_{s,k} = S_{\varphi}^{-1}(\varphi_k - \bar{\varphi}), \quad (2.56)$$

$$y_{s,k} = S_y^{-1}(y_k - \bar{y}). \quad (2.57)$$

След оценяване на параметрите, е получена оценката $\hat{\Theta}_s$ в (2.55). За формиране на дестандартизирания модел се извършват следните преобразувания. С използване на горните означения и зависимости моделът (2.55) може да се така:

$$\begin{aligned} S_y^{-1}(y_k - \bar{y}) &= \hat{\Theta}_s^T S_{\varphi}^{-1}(\varphi_k - \bar{\varphi}) + e_{s,k}, \\ y_k &= \bar{y} - S_y \hat{\Theta}_s^T S_{\varphi}^{-1} \bar{\varphi} + S_y \hat{\Theta}_s^T S_{\varphi}^{-1} \varphi_k + S_y e_{s,k}, \\ y_k &= \hat{\Theta}^T \varphi_k + e_k. \end{aligned} \quad (2.58)$$

Последният модел е дестандартизираният. Параметрите му са обединени в матрицата $\hat{\Theta} \in \mathcal{R}^{z \times \ell}$, която е

$$\hat{\Theta} = \begin{bmatrix} \hat{\theta}_0^T \\ \hat{\Theta}' \end{bmatrix}, \quad (2.59)$$

като $z = z' + 1$. Тъй като S_{φ}^{-1} и S_y са диагонални, от (2.58) и (2.59), за дестандартизираните оценки може да се запише:

$$\begin{aligned} \hat{\Theta}' &= S_{\varphi}^{-1} \hat{\Theta}_s S_y \in \mathcal{R}^{z' \times \ell}, \\ \hat{\theta}_0 &= \bar{y}^T - \hat{\Theta}'^T \bar{\varphi} \in \mathcal{R}^{\ell}. \end{aligned}$$

Добавеният първи ред е с индекс 0, тъй като той съдържа свободните членове в модела. Това са първите две събираеми в (2.58), които не зависят от регресорите. За да се изравни броят на стълбовете на $\hat{\Theta}$ и на елементите на φ_k , към вектора на регресорите се добавя 1 като първи елемент, т.е.

$$\varphi_k \leftarrow \begin{bmatrix} 1 \\ \varphi_k \end{bmatrix}.$$

Наличието на свободните членове $\hat{\Theta}_0$ отразява факта, че работната точка не съвпада с началото на координатната система (Фигура 2.18).

Както се вижда, за получаване на дестандартизираните оценки се изчисляват матриците S_{φ}^{-1} и S_y^{-1} . Тъй като те са диагонални, формирането им се свежда до изчисляване на $\sigma_{\varphi,i}^{-1}$ и $\sigma_{y,i}^{-1}$. Проблем би възникнал, ако има стандартни отклонения, които са близки до нула. Но ако сигнали с такива стандартни отклонения се премахват още при почистването на данните, не възникват числени проблеми на този етап.

Въпреки че стандартизацията е често срещана обработка на данните, има случаи, когато нейното прилагане не е необходимо. В точка 2.2.2.8 е описан начин за използване на символни (нечислови) данни, като се въвеждат т.нар. фиктивни променливи, приемащи стойности 0 или 1. Това води до естествено нормализирани регресори в диапазона $[0, 1]$. Такава нормализация е разгледана по-долу.

Нормализация в зададени граници

Друга нормализация, която намира приложение в идентификацията, е мащабирането на данните в предварително зададени граници (range normalization). Нека желаният интервал на изменение на нормализирания сигнал е $[0, 1]$. Тогава, ако x_k е скаларен, то нормализираният $x_{r,k}$ се формира, като

$$x_{r,k} = \frac{x_k - x_{min}}{x_{max} - x_{min}}.$$

С x_{min} и x_{max} са означени минималната и максималната стойност на x_k , изчислени от наличните данни.

Ако са зададени желаните долна x_l и горна x_u граници на обработения сигнал $x_{r,k}$, тогава нормализиращата формула е

$$x_{r,k} = \frac{x_k - x_{min}}{x_{max} - x_{min}}(x_u - x_l) + x_l. \quad (2.60)$$

Тази трансформация изисква по-малко изчисления от стандартизацията и за разлика от нея, където интервалът на изменение на $x_{s,k}$ зависи от данните, тук диапазоните на сигналите се задават в явен вид. Въпреки това стандартизацията намира най-широко приложение поради това, че трансформиранията величини са с удобни статистически характеристики ($m_{x,s} = 0$ и $\sigma_{x,s} = 1$), което улеснява анализа на получения модел. От тази гледна точка е удачно стандартизацията да се използва, когато величините имат унимодални разпределения. Ако е известно, че някое разпределение е с повече максимуми, то анализът няма да е коректен. В този случай е подходяща нормализацията (2.60).

2.2.2.5 Трансформация на данни

Когато в системата се наблюдават нелинейни зависимости на изходите от факторите, в някои случаи е възможно, чрез нелинейни трансформации на сигналите, те да се изменят така, че връзката между трансформираните фактори и изходи да се доближи до линейна [90]. Тази обработка е свързана и с допълнителен анализ, като целта е да се получат линейно параметризирани модели. Стремешът моделите да са линейни по отношение на параметрите, е обсъден в точка 2.1.2.5.

Трансформацията на данните може да се раздели на две групи: трансформация на отделни фактори и на група от фактори.

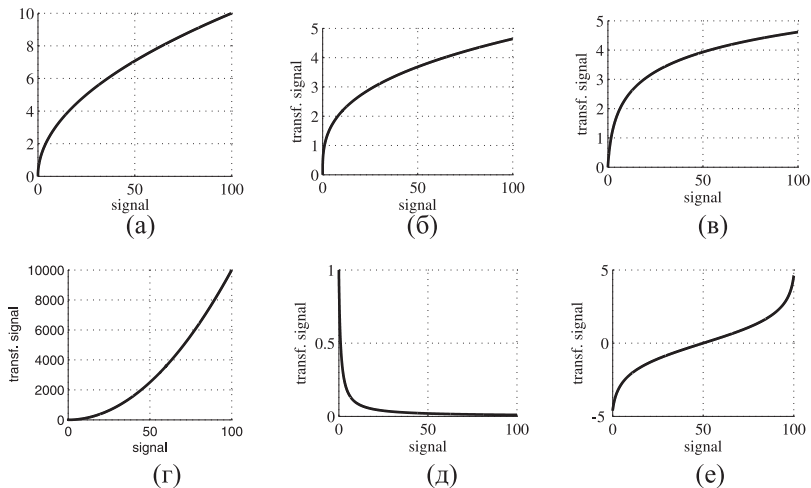
Трансформация на отделни фактори

Софтуерите за статистическа обработка на данни и моделиране, като SAS (Statistical Analysis System), SPSS (Statistical Product and Service Solutions) и R (от първата буква в имената на създателите Robert Gentleman и Ross Ihaka) предлагат набор от вградени нелинейни трансформации, които може да се използват при предварителната обработка на данните и са особено удобни, когато идентификацията е автоматизирана. Някои от тях, съобразени за целочислени стойности на x_k , са дадени в Таблица 2.1 и са представени графично на Фигура 2.20.

Таблица 2.1. Често използвани трансформации, съобразени за $x_k \in \mathcal{Z}$

трансформация	израз	сложност	забележка
квадратен корен	$\tilde{x}_k = \sqrt{x_k}$	2	ако $x_k < 0$, $\tilde{x}_k$ - липсваща
логаритмична	$\tilde{x}_k = \ln(x_k + 1)$	3	ако $x_k < 0$, $\tilde{x}_k$ - липсваща
квадратична	$\tilde{x}_k = x_k^2$	4	
кубичен корен	$\tilde{x}_k = \sqrt[3]{x_k}$	5	
реципрочна	$\tilde{x}_k = \frac{1}{x_k + 1}$	6	ако $x_k = -1$, $\tilde{x}_k$ - липсваща
сгъваща	$\tilde{x}_k = \ln \frac{x_k + 1}{x_{max} - x_k + 1}$	7	ако $x_k < 0$, $\tilde{x}_k$ - липсваща

Възможно е да се използва и полиномна интерполация от по-висок ред (от 2). Този вариант е удачен, ако областта на изменение на фактора е предварително известна и са налични данни в целия диапазон. Причината е, че извън диапазона на изменение на сигнала няма гаранция, че трансформираният полином е адекватен.



Фигура 2.20. Функции на трансформациите: квадратен корен (а), кубичен корен (б), логаритмична (в), квадратична (г), реципрочна (д) и сгъваща (е)

При значителен брой потенциални фактори, изборът на подходяща трансформация трябва да се автоматизира. Задачата за трансформация на факторите с цел получаване на линейно параметризирани модели може да се реши, като се приложи предварително избран набор от трансформации към всеки фактор и се изчисли корелационният коефициент на Пиърсън между получените нови фактори $\tilde{\varphi}_{i,k}$ и изходите на обекта. Този коефициент отразява силата на линейна зависимост между величините.

За скалярни сигнали $\tilde{\varphi}$ и y коефициентът на Пиърсън е

$$\rho_{\tilde{\varphi}y} = \frac{\sum_{k=1}^N (\tilde{\varphi}_k - \bar{\tilde{\varphi}})(y_k - \bar{y})}{\sqrt{\sum_{k=1}^N (\tilde{\varphi}_k - \bar{\tilde{\varphi}})^2} \sqrt{\sum_{k=1}^N (y_k - \bar{y})^2}} = \frac{\sigma_{\tilde{\varphi}y}^2}{\sigma_{\tilde{\varphi}} \sigma_y}.$$

Ако $\tilde{\varphi}_k$ и y_k са векторни, като първият елемент на $\tilde{\varphi}_k$ е оригиналният фактор, а останалите $h - 1$ елемента са неговите трансформации, стандартните им отклонения са обединени във векторите:

$$\sigma_{\tilde{\varphi}} = [\sigma_{\tilde{\varphi}_1} \quad \sigma_{\tilde{\varphi}_2} \quad \dots \quad \sigma_{\tilde{\varphi}_h}]^T,$$

$$\sigma_y = [\sigma_{y_1} \quad \sigma_{y_2} \quad \dots \quad \sigma_{y_\epsilon}]^T,$$

а $\Sigma_{\tilde{\varphi}y}$ е ковариационната матрица:

$$\Sigma_{\tilde{\varphi}y} = \begin{bmatrix} \text{cov}(\tilde{\varphi}_1, y_1) & \text{cov}(\tilde{\varphi}_1, y_2) & \dots & \text{cov}(\tilde{\varphi}_1, y_\ell) \\ \text{cov}(\tilde{\varphi}_2, y_1) & \text{cov}(\tilde{\varphi}_2, y_2) & \dots & \text{cov}(\tilde{\varphi}_2, y_\ell) \\ \vdots & \vdots & \ddots & \vdots \\ \text{cov}(\tilde{\varphi}_h, y_1) & \text{cov}(\tilde{\varphi}_h, y_2) & \dots & \text{cov}(\tilde{\varphi}_h, y_\ell) \end{bmatrix},$$

то коефициентът на Пийърсън е матрица с размерност $h \times \ell$, която се формира, като

$$\rho_{\tilde{\varphi}y} = \Sigma_{\tilde{\varphi}y} // (\sigma_{\tilde{\varphi}} \sigma_y^T).$$

Със символа ‘//’ е означено действието поелементно деление. Стойността на елемента $[\rho_{\tilde{\varphi}y}]_{ij}$ показва доколко зависимостта между i -тия елемент на $\tilde{\varphi}_k$ и j -тия изход е линейна.

Коефициентът на Пийърсън е подходящ за анализа, тъй като е нормиран (изменя се между -1 и 1) и освен това е удобен при автоматизиране на описаната дейност. Той е разгледан по-подробно в точка 2.2.3.3.

Ако изчислителната сложност е от значение, то е подходящо изборът на трансформацията да се извърши с отчитане на степента ѝ на сложност (виж трета колона в Таблица 2.1). Например, ако за най-подходящата трансформация се получава $\rho_{\tilde{\varphi}y_best} = 0.6$ и тя има сложност 5, то може да се въведе допустим праг за влошаване на корелираността, в рамките на който, ако се намери трансформация с по-малка сложност, тя да бъде избрана. Така, ако избраният праг е 3%, то се избира трансформацията, за която $\rho_{\tilde{\varphi}y} \geq 0.582 = 0.6 \times 97\%$, а сложността е по-малка от 5.

Трансформация на група от фактори

Понякога е подходящо вместо връзката между трансформациите на даден фактор и изход да се изследва влиянието на група от фактори върху изхода. Причина за търсенето на такива зависимости може да е заключение, базирано на знанието за процесите в системата.

Ако е известно, че връзката между даден фактор $\varphi_{i,k}$ и изход зависи от друг фактор $\varphi_{j,k}$, то е подходящо да се провери доколко би подобрило бъдещия модел въвеждането на нов (вторичен) фактор, например

получен, като

$$\varphi_{h,k} = \varphi_{i,k} \varphi_{j,k}. \quad (2.61)$$

В социологията например, когато се търси връзката на дохода (изход) от образованието (първичен фактор), може да възникне въпросът доколко полът на индивида влияе на тази връзка. Тогава на етапа на предварителната обработка е желателно да се въведе допълнителният фактор, който е произведение на споменатите два.

Възможни са и други подходящи зависимости между факторите, като най-вече предварителното знание за системата подпомага конкретния избор. В кредитната индустрия например кандидат с нисък доход може да е рисков за даден заем, а друг с висок доход да не е рисков кандидат за същия заем. Затова често се използва величината ‘заем като % от дохода’, т.е.

$$\varphi_{h,k} = \frac{\varphi_{i,k}}{\varphi_{j,k}} \times 100\%,$$

като факторът ‘заем’ е $\varphi_{i,k}$, а ‘доход’ е $\varphi_{j,k}$.

При автоматизирания подход, където човешката намеса е ограничена, отново се формират множество вторични фактори, които впоследствие участват при избора на подходяща структура на модела. Поради тази причина характерно за автоматизирането на идентификацията е, че въпреки изключването на неинформативни фактори крайният набор от предварителната обработка може да нарасне значително.

Трансформация на изходи

В примера на точка 2.1.3.2 е показан случай, при който за получаване на линейно параметризиран модел се трансформира изходът. Това е типична дейност при моделирането на пазарни системи. В случая се търси такова нелинейно изменение на изхода, при което се постига максимално достоверен линеен модел.

Нека първите ℓ елемента на вектора $\tilde{y}_k$ са компонентите на y_k , а останалите елементи са трансформирания изход. Ако се формира линеен модел

$$\tilde{y}_k = \tilde{\Theta}^T \varphi_k + e_k,$$

е възможно да се определят търсените трансформации. Отново тази дейност е характерна за автоматизираната идентификация.

2.2.2.6 Декомпозиция на сигнали

Понякога подобряването на резултата от оценяването на параметрите се постига с декомпозиране на наблюдаваните величини [121]. Например, ако даден сигнал съдържа полезни съставки във високите и ниските честоти, той може да се декомпозира и двете съставки да се отчетат поотделно в моделирането. Често бавната компонента (дрейф, тренд) се разглежда като адитивна съставка.

Пример за декомпозиция на първоначален сигнал x_k е представянето му като сума от неговата средна стойност $\bar{x}$ (константен тренд) и центрирания сигнал $x_{c,k}$. Това е част от стандартизацията, описана в точка 2.2.2.4. При тази декомпозиция съставката $\bar{x}$ се определя преди оценяването на параметрите и се отчита в модела, когато той се дестандартизира. От друга страна, втората съставка $x_{c,k}$ участва в явен вид, като фактор в оценявания модел.

В икономическите системи величините често се декомпонират на тренд, сезонност и циклична съставка. Тези компоненти са разгледани по-долу.

Тренд

Някои сигнали е подходящо да се декомпонират на линеен тренд $x_{t,k}$ и остатъчна съставка $\tilde{x}_k$, съдържаща по-бързите вариации на първоначалния сигнал x_k . При адитивно наслагване на двете съставки декомпозицията е

$$x_k = x_{t,k} + \tilde{x}_k.$$

Линейният тренд, както и константният, са частни случаи на нелинейния тренд. Един начин за отделянето му е да се оценят параметрите на полином по степените на времето, т.е. моделът на тренда е

$$x_{t,k} = \varphi_{t,k}^T \theta_t + e_{t,k}.$$

Векторът на регресорите е $\varphi_{t,k} = [1 \ k \ \dots \ k^n]^T$, а $\theta_t \in \mathcal{R}^{n+1}$. Съществуват процедури, на базата на които определянето на реда на полинома може да се автоматизира.

За някои величини в икономиката са характерни цикличните и сезонните съставки. Има различни методи за тяхното оценяване, както и начини за отчитането им в модела. Периодът на сезонността е постоянен, а амплитудата обикновено се изменя в по-малки граници в сравнение с цикличния тренд. Освен това в общия случай периодът на цикличността е променлив. Поради нерегулярния си характер цикличността често

се обединява с тренда и тяхното оценяване се извършва едновременно.

Сезонност

Величина, за която е налице сезонност, може да се декомпозира на компонента, от която е премахната сезонността, и набор от коефициенти за всеки дискретен момент от периода на сезонността, които отразяват наличието на тази съставка. Един вариант е коефициентите да са пропорционални на средните стойности, съответстващи на първия, втория и т.н. момент, в рамките на един период на сезонността. Например, ако величината се отчита месечно и първият месец е януари, то коефициентът $\hat{\theta}_{s,1}$ зависи от средната стойност на отчетите в различните години за януари, $\hat{\theta}_{s,2}$ е пропорционален на средната стойност на отчетите за февруари през годините и т.н. Общата формула за получаване на сезонните коефициенти е

$$\hat{\theta}_{s,j} = \frac{\bar{x}_j}{\bar{x}},$$

където $\bar{x}$ е средната стойност на x_k за интервала на наблюдение, а индексът $j = \overline{1, h}$, като h е броят на дискретните моменти, с които се описва един период на сезонността (в примера $h = 12$). Ако отчетите не са цяло число години, средните стойности $\bar{x}_j$ в дискретните моменти, в рамките на периода на сезонността се изчисляват по следния начин:

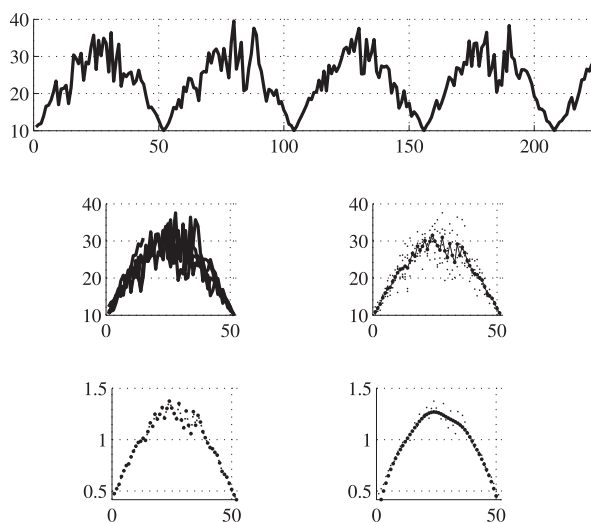
$$\bar{x}_j = \begin{cases} \frac{1}{n} \sum_{i=0}^{n-1} x_{j+ih}, & \text{за } N \leq j + nh, \\ \frac{1}{n+1} \sum_{i=0}^N x_{j+ih}, & \text{за } N > j + nh, \end{cases}$$

като $n = \lfloor N/h \rfloor$ е броят на пълните цикли на сезонността в набора от данни, а N е броят на наблюденията. Ако данните са за цели години (N е кратно на h), $\bar{x}_j$ се изчислява от първата формула. На Фигура 2.21, графично е представен начинът на формиране на сезонните компоненти. Обикновено броят n на стойностите, които се усредняват, е малък, а освен това сезонността зависи от много фактори, за които няма информация. Затова, с цел да се намали остатъчната неопределеност в елементите на вектора $\hat{\theta}_s$, е желателно те допълнително да се изгладят.

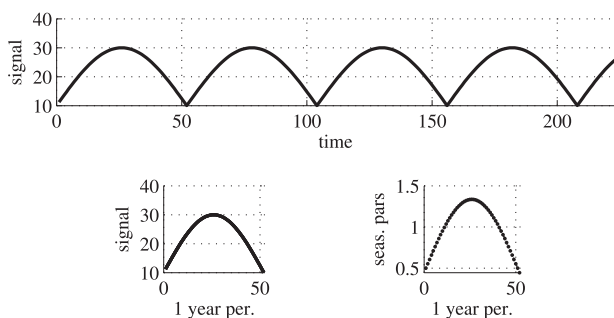
След определянето на $\hat{\theta}_s$ компонентата на x_k , от която е премахната сезонността, е

$$\tilde{x}_k = \frac{x_k}{\hat{\theta}_{s,j}}, \text{ за } j = k - \left\lfloor \frac{k-1}{h} \right\rfloor h.$$

Индексът j на параметъра е определен така, че да отговаря на съответ-



Фигура 2.21. Формиране на сезонните компоненти



Фигура 2.22. Формиране на сезонните коефициенти в идеализиран случай, при който x_k съдържа само сезонна компонента

ния момент от сезонността, в който е отчетена стойността x_k . Алгоритъмът лесно може да се провери, ако се приложи към идеализирания сигнал x_k на Фигура 2.22, в който, освен сезонна компонента, няма други съставки (тренд, случайна компонента и т.н.). След премахване на сезонността би трябвало $\tilde{x}_k$ да е константа. В този случай в дискретните моменти от периода на сезонността, независимо от конкретната година,

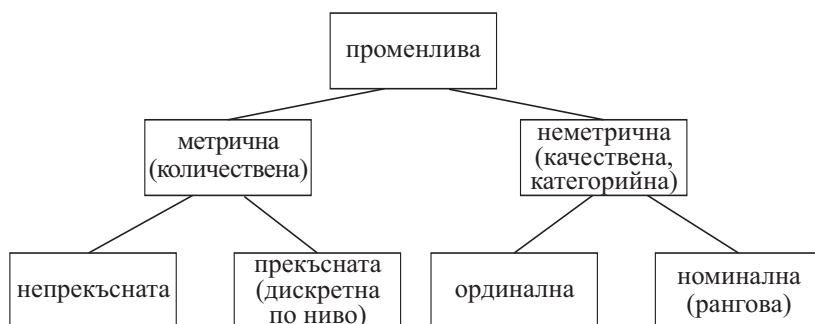
x_k приема една и съща стойност и тогава $\hat{\theta}_{s,j} = \frac{x_j}{\bar{x}}$. Наистина за този идеален сигнал, след премахване на сезонността, за $\tilde{x}_k$ се получава

$$\tilde{x}_k = \frac{x_k}{\hat{\theta}_{s,j}} = \bar{x}.$$

2.2.2.7 Видове променливи

Предварителната обработка до момента е описана най-вече от гледна точка на динамичните системи. Затова в болшинството от случаите данните се разглеждат като наблюдавани стойности на сигнали, изменящи се във времето. При това е прието, че съответните величини са количествени, а не качествени. Ако изследваната система е статична и особено ако някои данни не са числови, по-удачно е величините да се разглеждат най-общо като променливи (например променливата ‘кръвна група’ на пациент в медицинско изследване). По-долу се описва обработка, характерна за количествени и качествени величини, зависещи и независещи от времето.

Въпреки че понятието сигнал се използва в монографията в по-общ смисъл, т.е. функция, която не е задължително да зависи от времето, в болшинството от случаите това понятие се свързва с времето. Затова вместо сигнал ще се използва понятието „променлива“. Освен това, преди да се представят дейностите кодиране и категоризация, ще се разгледат по-подробно понятията количествена (метрична) и качествена (неметрична) променлива, както и техните разновидности.



Фигура 2.23. Класификация на променливите

На Фигура 2.23 е представена една класификация на променливите. Най-общо те се разделят на метрични и неметрични. Метричните (коли-

чествени) променливи може да са непрекъснати, ако приемат стойности от неизброимо множество (например реални числа), и прекъснати, ако приемат дискретни по ниво стойности (например цели числа). Между стойностите на метричните променливи съществува подредба, като разликите между отделните стойности имат определен количествен смисъл. Например температурата в пещ x_k е непрекъснатата променлива. Тя е метрична, тъй като две различни температури x_1 и x_2 може да се сравнят, например $x_1 < x_2$, както и разликата $\Delta x = x_2 - x_1$ има ясен количествен смисъл (x_2 е с Δx по-висока от x_1). Също така x_k е непрекъснатата променлива, защото може да се намери валидна стойност между кои да е две различни температури. От друга страна, броят на продадените бири на ден е прекъснатата (дискретна по ниво) променлива. Тя е метрична по същата причина като x_k , но е прекъснатата, защото между две съседни стойности, например 42 и 43, не може да се намери валиден брой продадени бири.

Неметричните (качествени) променливи се делят на ординални (между стойностите, на които също има подредба) и номинални (при които няма подредба). Характерно за неметричните променливи е, че разликата между стойностите им няма смисъл. Пример за ординална променлива е ‘образование’ – стойностите ѝ се подреждат възходящо така:

‘без образование’, ‘основно’, ‘средно’, ‘висше’,

но разликата

‘висше’ – ‘средно’

няма смисъл. Пример за номинална променлива е ‘кръвна група’. Възможните ѝ стойности са:

‘А’, ‘В’, ‘АВ’ и ‘0’,

като между тях няма естествена подредба. Неметричните (качествени) променливи приемат символни стойности (‘А’, ‘В’, ‘С’ или ‘1’, ‘2’, ‘3’ – в смисъл на символи, а не на числа) или стрингови (‘без образование’, ‘основно’ и т.н.). С помощта на такива величини може да се дефинират категории (наричани още групи или класове) от наблюдения (примерно категорията на висшистите в набора от данни).

2.2.2.8 Кодиране на променливи

Когато данните описват неметрични променливи, за да може да участват в идентификацията, те трябва да се преобразуват в метричен вид.

Важно е да се прави разлика между метричните променливи и неметричните, чиито стойности в набора от данни са дефинирани с цифри. Нека стойностите на номиналния фактор φ_k – ‘вид на промоцията’, да са зададени с цифри (примерно се разграничават 6 вида промоции и съответно $\varphi_k = \{‘1’, ‘2’, ‘3’, ‘4’, ‘5’, ‘6’\}$). Това обаче не означава, че в един линеен модел промоцията, кодирана с 1, трябва да има два пъти по-слабо влияние върху продажбите от промоцията, кодирана с 2. Тъй като стойностите на този фактор са условни и характеризират категории, те трябва да се интерпретират като символи, а не като числа.

От друга страна, поради факта, че ординалните променливи не са количествени (въпреки наличието на подредба между категориите), те също трябва да се интерпретират по подходящ начин.

По-долу са описани двата подхода за преобразуване на неметрични променливи в метрични. Първият подход се използва и за преобразуване на метрични променливи (както е описано в точка 2.2.2.9) и с това се постигат някои желани свойства на данните, а вторият е характерен за променливи без естествена подредба между стойностите.

Кодиране с фиктивни променливи

Може би най-често срещаният подход за преобразуване на неметрични данни в метрични е въвеждането на фиктивни променливи (dummy variables или dummies). В следващия пример се описва заместването на неметрична величина с набор от такива променливи.

Пример. *Кодиране на неметричен фактор с фиктивни променливи*

Нека има налични данни за технологичен обект, който преобразува един материал в друг, като съставът на първоначалния материал се е изменял по време на експеримента в зависимост от това, от кое хранилище е доставян.

Също така нека възможните източници на входен материал са четири. Очаква се изменението на състава да оказва влияние на поведението на обекта. Затова в процеса на моделиране трябва да се отчете текущият състав или с други думи да се въведе нова променлива ‘вид на състава’. Тъй като със стойностите ѝ не може да се извършват аритметични действия (тя е категорийна), то информацията в тази променлива трябва да се представи в числов вид. За целта нека се въведат четири фиктивни променливи $d_{1,k}$ (отговаряща на състав 1), $d_{2,k}$ (на състав 2) и т.н.

За удобство се въвежда векторът d_k , елементите на който са стойностите на четирите нови променливи в k -тия такт. Всяка от тях може

Таблица 2.2. Кодирание на променливата 'вид на състава' с фиктивни променливи

такт	d_1	d_2	d_3	d_4
1	1	0	0	0
2	1	0	0	0
3	0	1	0	0
4	0	1	0	0
5	0	0	1	0
6	0	0	1	0
7	0	1	0	0
8	0	0	1	0
9	0	0	0	1
10	0	0	0	1

да приема следните стойности:

$$d_{i,k} = \begin{cases} 1, & \text{ако материалът в момент } k \text{ е от източник } i, \\ 0, & \text{ако материалът не е от източник } i. \end{cases}$$

Въпреки че променливите приемат една от две възможни стойности и имат смисъл на флагове, те са метрични (дискретни по ниво) величини, които директно може да участват в моделирането.

В Таблица 2.2 са показани примерни стойности на променливите. Тъй като постановката е такава, че в даден момент материалът е от едно хранилище, то във всеки момент стойността само на една променлива е единица, а останалите са нула. $\diamond$

Фиктивните променливи описват пълен набор от взаимно изключващи се събития (затова в горния пример, ако $d_{i,k} = 1$, то всички останали са 0). В резултат на това една от фиктивните променливи ще е излишна, защото не допринася с нова информация към съдържащата се в останалите променливи. За горния пример е очевидно, че ако съставът не отговаря на източник 1, 2 и 3, той със сигурност отговаря на състава от източник 4. С други думи винаги една от фиктивните променливи е линейно зависима от останалите, а това може да доведе на следващия етап от идентификацията до числени проблеми (при оценяването на параметрите). Ето защо след кодирането една от променливите трябва да се изключи. Това може да стане ръчно, ако е необходимо да се отчете спецификата в поведението на обекта. Например, ако в основната част от наблюденията материалът е доставян от източник 3, фиктивната променлива, с която тази характеристика е кодирана, може

да се изключи от данните. По този начин моделът ще отчита ефекта от по-рядко срещаните случаи спрямо най-често срещания (а тъй като съответната фиктивна променлива е премахната от набора, тя няма да влияе на изхода на модела). Другият вариант е изключването да се извърши автоматично, както е споменато в точка 2.2.3.3, като съображенията за избора са свързани с числените особености на задачата, а не с априорната информация.

Независимо дали съществува подредба между стойностите на променливата, логиката за кодирането ѝ с фиктивни променливи не се променя. Когато се прилага за метрични величини, обикновено тази техника се комбинира с дискретизацията по ниво, описана в точка 2.2.2.9.

Едно удобство, свързано с употребата на фиктивни променливи, е, че полученият модел може лесно да се интерпретира. Тъй като тези променливи приемат стойности 0 и 1, то параметрите на линеен модел отразяват директно степента на влияние на кодираните от тях случаи върху изхода. Например, ако моделът е

$$y_k = \theta_1 d_{1,k} + \theta_2 d_{1,k}$$

и $\theta_1 = 2$, то когато $d_{1,k} = 1$, изходът нараства с 2.

Кодиране с една метрична променлива

Има и други подходи за преобразуване на променливи в метричен вид. Описаните по-долу техники се използват най-вече когато величините са номинални. Понякога е желателно, вместо да се въвежда множество от променливи, да се формира една метрична величина, която да обобщи информацията в първоначалната променлива, имаща отношение към изхода. Ако възможните стойности на номиналната величина са твърде много, тогава не е удачно да се използва подходът с фиктивните променливи, тъй като техният брой би бил голям, а това може излишно да усложни процеса на моделиране. По-подходящо е в този случай да се формира една нова променлива. Преобразуването в метричен вид означава, че трябва да се дефинира подредба между стойностите на първоначалната променлива и те да имат ясен количествен смисъл. Със следващите два примера са описани такива преобразувания за MISO система.

Пример. *Кодиране на номинална променлива с използване на средната стойност на изхода*

При изграждането на модел, описващ влиянието на различни видове реклама върху продажбите на продукт, е възможно да се подходи

по следния начин. Тъй като променливата ‘реклама’ (означена с φ_k) е номинална, то за въвеждане на подредба между стойностите ѝ за всяка от прилаганите реклами може да се определи средната стойност на продажбите:

$$\bar{y}_i = \left\{ \frac{1}{N_i} (\sum_{k=1}^N y_k | \varphi_k = 'i\text{-ти вид реклама}') \right\},$$

като N_i е броят на наблюденията, в които е използвана i -тата реклама. Така вместо номиналния фактор φ_k с възможни стойности:

{‘1-ви вид реклама’, ‘2-ри вид реклама’, ..., ‘ n -ти вид реклама’},

за построяване на модел може да се използва метричният фактор:

$$\tilde{\varphi}_k = \{\bar{y}_1, \bar{y}_2, \dots, \bar{y}_n\},$$

между стойностите на който съществува подредба. Освен това стойностите на $\tilde{\varphi}_k$ имат количествен смисъл. Ако ефектът на два вида (i -ти и j -ти вид) реклами е сходен върху продажбите, то и съответните стойности на $\tilde{\varphi}_k$ са сходни (т.е. $\bar{y}_i \approx \bar{y}_j$). $\diamond$

Използването на средните стойности на изхода (свързани с отделните групи наблюдения) е подходящо, когато изходът е непрекъсната метрична променлива. Описаният подход в следващия пример е по-удачен, когато изходът е дискретна метрична или неметрична променлива.

В долния пример y_k е биномен (приема две възможни състояния). Такива изходи се използват в техниката при изграждането на адаптивни системи, където се оценява вероятността за възникване на повреда; в медицината – за прогнозиране на вероятността пациент да страда от заболяване и т.н. В примера въвеждането на подредба между стойностите на номинален фактор е онагледено с приложение във финансите.

Пример. *Кодиране на номинална променлива с т.нар. „тегло на достоверност“ (WoE – Weight of Evidence)*

При формиране на модел за оценка на кредитния риск y_k е вероятността индивид да принадлежи към една от две възможни групи (‘добър’ и ‘лош’), т.е. изходът на системата е биномен. Ако даден фактор е неметричен, то всяка негова стойност може да се преобразува така, че новополучената променлива да е метрична. Например нека факторът φ_k е ‘местоживеене’ и приема следните нечислови стойности: ‘собственик на жилище’, ‘наемател на жилище’, ‘наемател с хазяин’ и ‘живее с родители’. Също така, за удобство стойностите на y_k са кодирани с двете числови стойности: 1 – ‘добър’ и 0 – ‘лош’. Тогава измежду наличните

собственици (в набора от данни) броят на добрите g_1 и на лошите b_1 кредитополучатели е

$$g_1 = \{\sum_{k=1}^N y_k | \varphi_k = \text{'собственик на жилище'}\},$$

$$b_1 = \{\sum_{k=1}^N (1 - y_k) | \varphi_k = \text{'собственик на жилище'}\}.$$

Съответно g_2 и b_2 са броят на добрите и на лошите наематели на жилище в набора, g_3 и b_3 – броят на добрите и лошите наематели с хазяин в същото жилище, а g_4 и b_4 – броят на добрите и на лошите в категорията на живеещите с родители. Освен това сумата на всички добри е $G = \sum g_i$, а на лошите е $B = \sum b_i$. Тогава възможните стойности WoE_i на новия, вече числов фактор $\tilde{\varphi}_k$ се дефинират за всяка категория, като

$$\text{WoE}_i = \ln \left(\frac{g_i/G}{b_i/B} \right), \text{ за } i = \overline{1,4}.$$

Отношението на общия брой на лошите към общия брой на добрите в набора от данни G/B е константа, която не зависи от конкретната категория, следователно подредбата между стойностите WoE_i зависи само от отношението g_i/b_i . Нека за конкретния набор е в сила:

$$\frac{g_3}{b_3} < \frac{g_2}{b_2} < \frac{g_4}{b_4} < \frac{g_1}{b_1}.$$

Тъй като натуралният логаритъм е монотонно нарастваща функция, то

$$\text{WoE}_3 < \text{WoE}_2 < \text{WoE}_4 < \text{WoE}_1.$$

Освен въведената подредба WoE_i имат и количествен смисъл. Например, ако групите на живеещите с родители и с хазяин са със сходно поведение (тогава и $g_2/b_2 \approx g_3/b_3$), а тези със собствено жилище се отличават в по-голяма степен от споменатите две групи, то разликата между стойностите WoE_2 и WoE_3 ще е малка, а между WoE_1 и WoE_2 (или между WoE_1 и WoE_3) ще е по-голяма.

Така, вместо номиналния фактор:

$$\varphi_k = \{\text{'собственик на жилище'}, \text{'наемател на жилище'}, \\ \text{'наемател с хазяин'}, \text{'живее с родители'}\},$$

за построяването на модел може да се използва количественият фактор:

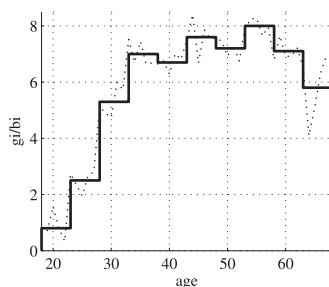
$$\tilde{\varphi}_k = \{\text{WoE}_1, \text{WoE}_2, \text{WoE}_3, \text{WoE}_4\}. \quad \diamond$$

2.2.2.9 Категоризация на променливи

Дейността в тази подточка се нарича категоризация, защото резултатът от прилагането ѝ е нова променлива, която винаги е категорийна и по-точно ординална. При категоризацията стойностите на променливата се обединяват в статистически различни групи (категории). Тя се използва най-вече при статични системи и е основна дейност в някои, установени в практиката, методологии за изграждане на статичен модел. Затова е описана, въпреки, че не е удачна за динамични системи.

Тази обработка може да се приложи към метрични и към неметрични променливи и може да се разглежда като дискретизация по ниво (всяка категория отговаря на определено дискретно ниво на сигнал). В резултат на категоризирането на метрична променлива се загубва нейният количествен смисъл (резултатната ординална променлива е качествена, а не количествена). Ако първоначалната променлива е ординална, то категоризацията е всъщност окрупняване на първоначалните категории. Тъй като между стойностите на номиналните променливи няма подредба, а както ще стане ясно, категоризирането изисква наличието на такава, то е необходимо първо номиналните променливи да се преобразуват в метрични и след това да се категоризират. За да се запази последователността и след категоризацията, сливането на групите се извършва винаги между съседни стойности (групи).

Тъй като стойностите на резултатната променлива са агрегирани, на пръв поглед категоризацията води до загуба на информация. Но от статистическа гледна точка правилно извършената категоризация води до потискане на смущенията. Ако вариацията между стойностите на y_k в рамките на дадена група се дължи на случайни фактори, то е по-добре да се работи не с тези стойности, а със средната $\bar{y}_i$ за групата. Причината е, че при усредняването влиянието на случайната съставка се посикта. Например (виж Фигура 2.24), ако хората на възраст от 18 до 22 години имат сходно поведение (в някакъв смисъл), то обединяването на стойностите 18, 19, 20, 21 и 22 на променливата 'възраст' в новата категория '[18, 23)' не би довело до значима загуба на информа-



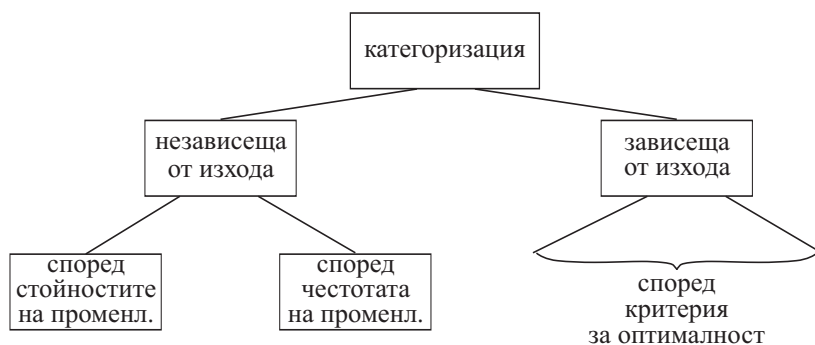
Фигура 2.24. Отношението g_i/b_i за групите от кандидати на еднаква възраст и g_i/b_i за групите в интервали през 5 години

ция, а същевременно до потискане на неопределеността във вариацията на изхода. Това е показано на фигурата и тъй като величината y_k е биномна, по ординатата е дадено отношението g_i/b_i .

С категоризацията може да се отчете и априорна информация за системата. Групата '[18, 23]' например може да съдържа най-вече студенти и да се очаква, че тяхното поведение като цяло е сходно, но същевременно се различава от средното за групата '[23, 27]' и т.н.

Така с категоризацията може да се потисне неопределеността в данните с цената на неголяма загуба на полезна информация. Неслучайно в области като оценка на кредитния риск категоризацията (зависеща както от чисто статистически съображения, от бизнес логиката, така и от опита на изследователя) отнема значителна част от времето на цялостния процес на идентификация.

Друга възможност, която категоризацията дава, е моделирането на нелинейни зависимости с помощта на линеен модел. На Фигура 2.24 се вижда, че до около 55 години с нарастване на възрастта нараства и тенденцията кредитоискателите да са добри, след което се наблюдава спад. Също така първоначалното нарастване не е линейно. По тази причина не е удачно да се търси линейна връзка между възрастта и вероятността за добро поведение на клиентите на банката. От друга страна, ако факторът възраст се категоризира и след това се кодира с фиктивни променливи, всяка категория би участвала в линейния модел независимо от останалите, като на всяка фиктивна променлива би отговарял съответен параметър на модела.



Фигура 2.25. Методи за категоризация

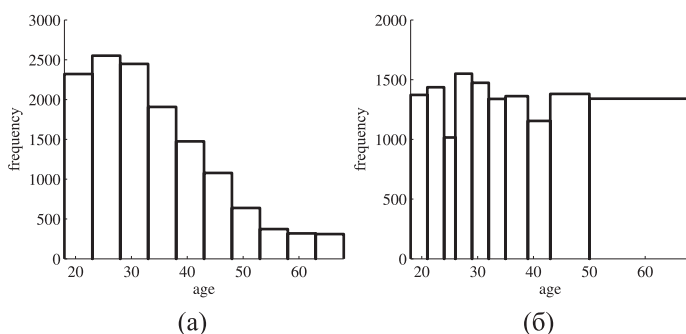
Видове категоризация

В долните разглеждания за опростяване на изложението отново се използват MISO модели. Ако системата е с повече изходи, тогава методите се прилагат поотделно за всеки изход.

Една класификация на методите за категоризация е дадена на Фигура 2.25. Според нея има две основни групи методи. Едната е без отчитане, а другата е с отчитане стойностите на изхода.

Категоризация, независеща от изхода

Методите, които не зависят от изхода, се прилагат, когато променливата е метрична. Те се делят на такива, които използват директно стойностите на променливата, подлежаща на категоризация, и такива, които използват честотата на формираните групи (брой или процентно съотношение на появите на стойности от всяка група в данните).



Фигура 2.26. Категоризация, независеща от изхода, с критерии: равни интервали (а) и равен брой наблюдения в групите (б)

Стойностите на променливата може да се използват, за да се формират категории, например през равни интервали, както е показано на Фигура 2.26 (а) – индивидите са разделени на групи през 5 години, но за разлика от Фигура 2.24 тук по ординатата са нанесени честотите. При методите с отчитане на честотата на групите най-често стремежът е да се формират групи с равен (доколкото е възможно) брой наблюдения във всяка група. Така, ако към данните от Фигура 2.24 се приложи категоризация с отчитане на честотата, тъй като 70% от кандидатите са на възраст между 18 и 35 години, то се очаква интервалите, отговарящи на групите в този диапазон, да са значително по-малки от тези, извън диапазона. На Фигура 2.26 (б) е показан резултатът от тази категори-

зация. От статистическа гледна точка стремежът във всяка група да има достатъчно данни, е с цел информацията, съдържаща се в групите, да се отчете коректно от модела. Ако малко наблюдения попадат в определена група, то е възможно оценката на параметъра, свързан със съответната фиктивна променлива, да е неточна, защото зависи твърде много от неопределеността. Също така, ако в една група попадат почти всички наблюдения, то фиктивната променлива (нека е $d_{i,k}$), която я описва, почти винаги би приемала стойност 1. Това също не е желателно от гледна точка на моделирането, тъй като се губи информация от сливането, а поради слабата вариация на $d_{i,k}$ оценката на параметъра отново би била неточна.

Методите за категоризация, описани по-долу, са предпочитани, тъй като категориите се формират така, че да са статистически различни от гледна точка на зависимата променлива. Също така, въпреки по-сложната логика на следващите методи, всички методи за категоризация може да се автоматизират.

Категоризация, зависеща от изхода

Методите от тази група може да се приложат, когато изходите приемат стойности от неизброимо множество (непрекъснати по ниво променливи), като приемат дискретни по ниво стойности (например биномен изход: ‘здравословно състояние’ = {‘болен’, ‘здрав’}), а също и ако са неметрични (‘вид повреда’ = {‘прекъсната статорна намотка’, ‘повреда в редуктора’, ‘повреда в тахометъра’ и др.}).

Обикновено тази категоризация се извършва, когато моделът ще се използва за класифициране на поведението на обекта според описващите го характеристики. За разлика от първата група методи, които може да се приложат и върху изхода, тук категоризацията се прилага само върху факторите. Целта е новосформираните групи да са колкото е възможно по-разнородни от гледна точка на стойностите на изхода. Както беше описано в предишната подточка, ако има стойности на фактор, за които изходът приема сходни стойности (в рамките на нивото на неопределеност), те е желателно да се обединят в обща група. Категоризирането на факторите с цел максимизиране на разнородността на групите е оптимизационна задача. Методите се различават най-вече по критерия, който се оптимизира. Често използвани критерии са χ^2 , VoI (Value of Information), Gini и др. [41], описани в примера на точка 2.2.3.2.

Оптимизацията се извършва при наличието на ограничения като:

- максимален брой нива на категоризираната променлива;
- монотонно изменение на средните стойности на изхода за групите

($\bar{y}_i$ е монотонна спрямо i);

- минимален брой наблюдения в група над определен праг;
- максимален брой наблюдения в група под определен праг и др.

С условието за монотонен тренд може да се отчитат някои априорни знания. Обикновено това ограничение в оптимизационната процедура произлиза от знанията за приложната област. Например известно е, че с нарастване на дохода трендът в стойностите на g_i/b_i се очаква да е положителен, т.е. очаква се кандидатите да стават по-малко рискови, когато доходът им се увеличава. Съответно банката не би приела модел, ако има нелогични трендове и няма сериозно бизнес обяснение за това. В такива случаи е желателно монотонността на очаквания тренд да се гарантира още на етапа на предварителната обработка на данните.

Когато броят на възможните стойности на първоначалната променлива е голям и особено ако тя е непрекъсната, първо се извършва категоризация без отчитане на стойностите на изхода. Целта е да се ограничат нивата на променливата, преди да се стартира значително по-сложният оптимизационен метод за категоризация с отчитане на изхода.

2.2.3 Анализ на данни

Обикновено получаването на набор от данни, който да доведе до формирането на достоверен модел, още на този етап от идентификацията е свързано с провеждането на серия от анализи. Повечето от описаните до момента стъпки също изискват анализ, например за коректно отчитане на нехарактерни и липсващи стойности, за качествена категоризация, зависеща от изхода и т.н.

На този етап от моделирането се разграничават анализи на:

- отделна променлива;
- групи от променливи.

От гледна точка на това, доколко изследователят предварително е запознат със системата, анализът може да е незаменим за първоначалната ориентация в данните.

Част от анализирането на данните и вземането на решения може да се автоматизира. По-трудна за автоматизация е тази част, която е свързана с отчитането на априорната информация. Въпреки това, когато идентификацията в определена приложна област се изпълнява регулярно, дори и тази част (до известна степен) може и е целесъобразно да се автоматизира.

2.2.3.1 Първоначална ориентация в данните

Качественото протичане на идентификацията често е свързано с изучаване на особеностите на обекта по време на всеки неин етап. Понякога изследователят има нужда да изучи внимателно данните не само за да избегне грешни стойности или да подпомогне обработката, но и за да вникне в скритата логика, като намира потвърждение на априорните знания в данните или (в обратна посока) търси обяснение на наблюдаваните свойства в данните, събирайки допълнителна информация за системата. Благодарение на тази дейност се натрупат нови знания и е възможно резултатът от идентификацията качествено да се подобри.

Особено когато изследователят не е участвал в провеждането на експеримента, той има нужда от първоначална ориентация – какви данни се съдържат в предоставения му набор. В това отношение са описани два анализа, наречени статистически и честотен. При единия се определят статистически характеристики, описващи като цяло данните за отделните променливи, а при втория анализ се изследват експерименталните разпределения на величините.

Статистически анализ

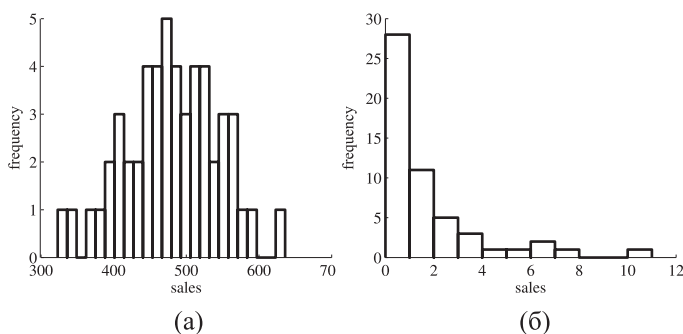
Първият анализ след получаването на данните, който е желателно да се проведе, е статистическият. При него изследователят се запознава общо с данните за всяка величина и евентуално прави връзка с предварителното знание за системата. Това включва изолирането на липсващи стойности, определянето на типа на променливата (от който зависи как ще се интерпретират и обработват съответните стойности), определянето на диапазона на изменение, сумата от стойностите (кое-то е полезно при променливи, имащи смисъл на бройки или флагове), средната стойност, стандартното отклонение, медианата, модата, проценти (примерно 5-и, 95-и) и др.

Тези и други статистически характеристики са описани подробно в [112]. С тяхна помощ, определена информация за разпределението на променливите се представя със скаларни стойности, което улеснява автоматизираното вземане на решения за следващи стъпки, които трябва да се предприемат по отношение на данните. Например автоматично може стойностите под 5-ия и/или над 95-ия процентил да се заместят със стойността на съответния процентил. Това е стандартна дейност в някои автоматизирани методологии, с която се намалява влиянието на нехарактерните стойности на променливите, водещи понякога до т.нар.

тежки опашки в разпределенията на данните (виж точка 3.2.7). Също при голям брой липсващи стойности променливата може да се изключи от набора и т.н.

Честотен анализ

Този анализ се прилага за изследване на разпределението на категориална (неметрична или дискретизирана по ниво метрична) променлива. При него за всяка от срещаните в набора стойности на променливата се определя честотата, в смисъл на брой (или по-често в проценти) появи на стойността. Също често се изследва кумулативната честота, средната стойност на всяка група, ако преди това променливата е категоризирана, и др. Визуализацията на честотите улеснява анализа.



Фигура 2.27. Честота на продажбите на хляб (а) и хладилници (б)

На Фигура 2.27 са представени честотите (брой продажби) в различни подинтервали на променливата ‘продажби’. Така изследователят може да направи заключение за това, дали определени продукти се търсят регулярно (като хляб, плодове и др., при които продажбите имат разпределение, близко до нормалното), или търсенето е нерегулярно (например на бяла техника, телевизори и др., продажбите на които имат поасоново разпределение). В зависимост от вида на разпределението, се избира и подходящ метод за същинското построяване на модела [83, 121]. Оценката на разпределението на продукта също може да се автоматизира, като се определи разликата между разпределението на данните и на теоретичното нормално и поасоново разпределение.

2.2.3.2 Информативност на фактор

Понякога този анализ се нарича характеристичен, тъй като се изследва връзката между даден фактор и изход на системата. Невинаги голямата вариация на фактор означава, че обектът е качествено възбуден. Може изходът да не е чувствителен към фактора и това означава, че тази променлива има смисъл по-скоро на шум, отколкото на полезен сигнал. В долния пример са споменати някои статистики, с които се оценява количеството информация за изхода, съдържащо се в даден фактор.

Пример. *Информативност на фактор при биномен изход*

Ако y_k е биномен, то при характеристичния анализ се оценява дискриминативността на фактора, в смисъл колко качествено той разграничава стойностите на изхода. За да се получи количествена представа за това, се изчисляват статистики като степен на информативност VoI (Value of information, наричана още Kullback–Leibler divergence):

$$\text{VoI} = \sum_{i=1}^N \left(\frac{g_i}{G} - \frac{b_i}{B} \right) \ln \frac{g_i/G}{b_i/B},$$

индекс на различие (Diversity Index):

$$D = 1 - \sum_{i=1}^N \frac{g_i b_i}{g_i + b_i} \frac{G + b}{GB},$$

който се изменя от 0 до 1, като при $D = 1$ факторът дискриминира напълно стойностите на изхода, както и коефициент на Джини (Gini), статистика на Колмогоров-Смирнов KS (Kolmogorov-Smirnov) и др. Те са описани подробно в [41]. $\diamond$

Важно условие за получаване на информативни данни е всеки от входните сигнали да възбужда изходите на системата. При провеждането на активен експеримент това може да се постигне, но когато данните са събрани пасивно, няма гаранция за изпълнение на това условие. В резултат от анализите до момента може да се прецени кои входове не възбуждат обекта и кои изходи са невъзбудени. Следващата стъпка е фокусът на изследователя да се концентрира върху тази част от системата, за която има подходящи за моделиране данни. Въпреки това (както е описано в следващата точка) не трябва да се бърза с премахането на фактори, които (поотделно) не влияят силно на изхода.

2.2.3.3 Информативност на група фактори

Освен изучаването на факторите поотделно е важно те да се разглеждат и на групи. Възможно е някои фактори, които не са свързани силно с изхода, когато са в група, да допринесат чувствително за неговото описание. Такава ситуация се наблюдава, когато факторите не са корелирани помежду си и сумарната информация в тях осезателно подобрява модела. Също така може фактори, които са информативни, в група да не допринасят чувствително повече, отколкото поотделно. Обяснение за това е, че те са силно корелирани помежду си, което значи, че голяма част от информацията в тях се припокрива и участието на един от факторите в модела обезсмисля използването на останалите. Ето защо не трябва да се бърза с изводите, преди да се провери информативността на факторите в група.

Корелационен анализ

При анализа на връзката между сигнали обикновено се определя силата на зависимостта, нейната посока и форма. Първата задача е обект на корелационния анализ, а втората и третата задача – на регресионния. Регресионният анализ няма да се разглежда отделно, тъй като той е свързан с определянето на модел (в случая за целите на анализа), а тази тема е подробно разгледана в Трета глава.

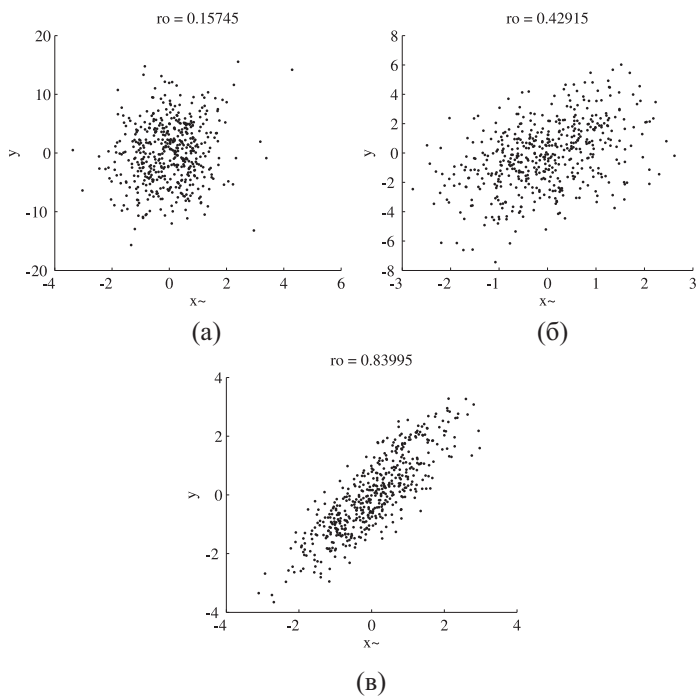
Според това, дали се изучава единична връзка или множество от връзки с даден изход, корелационният анализ е единичен или множествен.

Единичен корелационен анализ

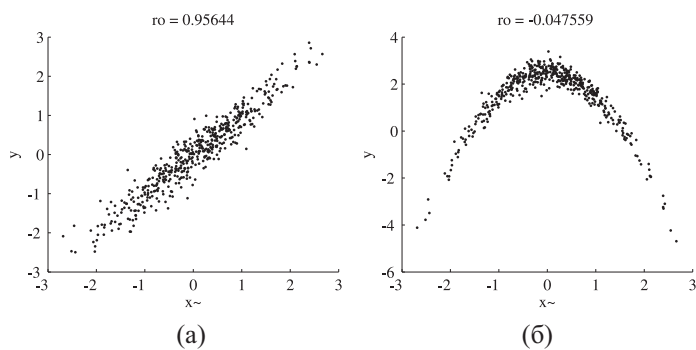
Ако скаларните величини x и y са количествени (не е нужно едната да е изходна, а другата да е фактор – може да се проверява зависимостта между фактори или между изходи), подходящ измерител на силата на връзката между тях е корелационният коефициент на Пийърсън:

$$\rho_{xy} = \frac{\sigma_{xy}}{\sigma_y \sigma_x} = \frac{\sum_{k=1}^N (x_k - \bar{x})(y_k - \bar{y})}{\sqrt{\sum_{k=1}^N (x_k - \bar{x})^2 \sum_{k=1}^N (y_k - \bar{y})^2}}.$$

Той е описан в 2.2.2.5 от гледна точка на трансформирането на факторите. С него може да се установи типът на зависимостта между x и y . Например, ако стойностите на x се коренуват и ρ_{xy} нарасне чувствително спрямо стойността му за първоначалните данни, то може да се предположи, че функционалната зависимост $y = f(x)$ е близка до квадратичната.



Фигура 2.28. Корелационен коефициент при различна степен на линейна зависимост: слаба (а), умерена (б) и силна (в) зависимост между $\tilde{x}$ и y



Фигура 2.29. Корелационен коефициент при много силна линейна (а) и не-линейна (б) зависимост между $\tilde{x}$ и y

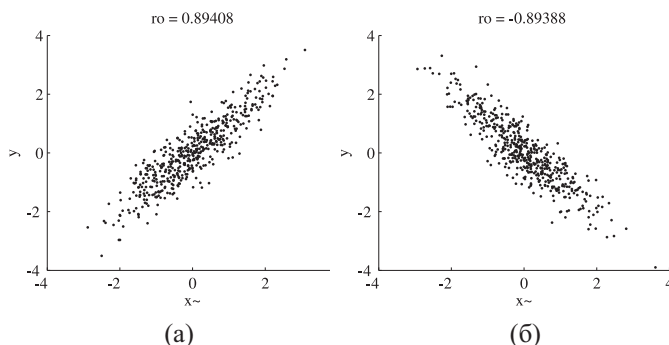
От последната формула се вижда, че коефициентът на Пиърсън е комутативен, в смисъл че $\rho_{xy} = \rho_{yx}$. Стойността му е в интервала $[-1, 1]$. Колкото ρ е по-голям по абсолютна стойност, толкова връзката между сигналите е по-близка до линейната (което се използва в трансформационния анализ).

На Фигура 2.28 са дадени диаграмите на разсейване (наричани още точкови диаграми) на x и y при различна сила на линейната зависимост. На Фигура 2.29 е представена силна линейна и нелинейна зависимост. Вижда се, че във втория случай корелационният коефициент е близък до нула.

Таблица 2.3. Степен на корелационна зависимост и стойности на $|\rho|$

степен на корелация	диапазон
слаба	$0 \leq \rho \leq 0.3$
умерена	$0.3 < \rho \leq 0.5$
значителна	$0.5 < \rho \leq 0.7$
силна	$0.7 < \rho \leq 0.9$
много силна	$0.9 < \rho \leq 1$

В Таблица 2.3 са дадени приетите степени на сила на връзката и съответните диапазони на изменение на $|\rho|$.



Фигура 2.30. Нарастваща (а) и намаляваща (б) зависимост и знак на ρ

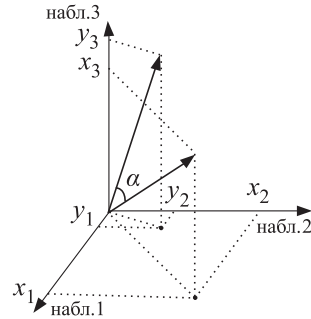
Когато $\rho > 0$, корелацията е права (Фигура 2.30) – с нарастване на едната величина тенденцията е да нараства и другата. В противен случай (при $\rho < 0$) е обратна, т.е. с нарастване на едната величина, тенденцията е другата да намалява).

Тъй като стойността на коефициента на корелация се определя на базата на данни, съдържащи неопределеност, стойността му (както и всички статистически характеристики, определени на базата на експериментални данни) има вероятностен характер. Поради това корелационният анализ може да завърши с проверка на хипотезата за статистическа значимост на стойността му. Нулевата хипотеза е, че не съществува връзка между величините, т.е. теоретичната стойност е $\rho_T = 0$, което означава, че ако емпирично определената стойност е различна от 0, това се дължи на случайни фактори. В [22] подробно е засегната темата за проверката на значимостта на експериментално получената стойност за коефициента на Пиърсън.

Графично коефициентът на Пиърсън се интерпретира като косинуса на ъгъла между векторите, съдържащи стойностите $x_{c,k}$ и $y_{c,k}$ на центрираните сигнали за интервала на наблюдение, т.е.

$$\rho_{xy} = \frac{x_c^T y_c}{\|x_c\|_2 \|y_c\|_2} = \cos \alpha,$$

където $x_c = [x_{c,1} \ x_{c,2} \ \dots \ x_{c,N}]^T$ и $y_c = [y_{c,1} \ y_{c,2} \ \dots \ y_{c,N}]^T$. Например, ако броят на наблюденията е $N = 3$, то ρ_{xy} е косинусът на пространствения ъгъл между векторите $x, y \in \mathcal{R}^3$ (Фигура 2.31).



Фигура 2.31. Интерпретация на коефициента на Пиърсън

До момента е разгледан случаят, при който се търси връзка между скаларни величини. Ако те са векторни, т.е. в k -тото наблюдение $x_k \in \mathcal{R}^m$ и $y_k \in \mathcal{R}^n$, коефициентът на Пиърсън е матрицата $\rho_{xy} \in \mathcal{R}^{m \times n}$, която е

$$\rho_{xy} = \Sigma_{xy} // (\sigma_x \sigma_y^T). \quad (2.62)$$

Както е описано в 2.2.2.5, действието ‘//’ е поелементно деление. Тук също важи връзката $\rho_{xy} = \rho_{yx}^T$.

Корелационният анализ може да се приложи както за определяне на това, доколко силен (информативен) е даден фактор (т.е. да се изследва връзката между фактор и изход), така и колко зависими са два фактора. Въпреки че в последното разглеждане сигналите са векторни, а

ρ_{xy} е матрица, изследваната корелация все още се нарича единична, тъй като елементите на матрицата $[\rho_{xy}]_{ij}$ характеризират единични връзки (между x_i и y_j).

Множествен корелационен анализ

Когато се изследва до каква степен множество от фактори влияе на определен изход, корелационният анализ се нарича множествен или множествена корелация. При изучаването на тази връзка се изчислява множественият коефициент на корелация $\rho_{y|\varphi} \in \mathcal{R}$ по следния начин:

$$\rho_{y|\varphi} = \sqrt{\rho_{\varphi y}^T \rho_{\varphi}^{-1} \rho_{\varphi y}},$$

където $\rho_{\varphi y} \in \mathcal{R}^z$ е векторът с елементи – коефициентите на Пиърсън между изхода и всеки от факторите, а ρ_{φ} е матрицата (2.62), но по отношение на взаимната връзка между факторите (тук $x \leftarrow \varphi$ и $y \leftarrow \varphi$). Диагоналните елементи на ρ_{φ} са 1 поради пълната корелация между всеки фактор, сравнен със себе си. За разлика от коефициента на Пиърсън $\rho_{y|\varphi}$ се изменя между 0 и 1 (а не в интервала $[-1, 1]$). Освен това, множественият корелационен коефициент не е комутативен. Например, нека са дадени три скаларни сигнала x_1 , x_2 и x_3 , като фактори са x_1 и x_2 , а зависима от тях променлива (изход) е x_3 . Тогава $\rho_{x_3|x_1, x_2}$ е множественият коефициент на корелация. От друга страна, ако за фактори се приемат x_1 и x_3 , а за зависима променлива – x_2 , то $\rho_{x_2|x_1, x_3}$ е множественият коефициент на корелация, отразяващ силата на влияние на x_1 и x_3 върху променливата x_2 . Съвсем естествено е връзката на x_3 с x_1 и x_2 да не е същата като тази между x_2 и двойката $\{x_1, x_3\}$ и съответно $\rho_{x_3|x_1, x_2} \neq \rho_{x_2|x_1, x_3}$.

Смисълът на множествения корелационен коефициент може да се обясни по следния начин. Колкото по-силна е изследваната зависимост – толкова по-голяма част от дисперсията на описваната променлива се дължи на влиянието на описващия/описващите фактори.

Съществуват и други показатели на силата на корелационните зависимости. Например коефициентът на определеност ρ^2 , който е равен на квадрата на коефициента на Пиърсън, отразява доколко факторите обясняват вариациите на изхода. Често в литературата този показател се записва като R^2 и затова в монографията се използва това означение. Използва се и коригираният коефициент на определеност R_a^2 , и показателят VAF (Variance Accounted For), който дава същата информация, но в проценти. Тези показатели са представени подробно в 2.4.

В зависимост от вида на променливите се използват различни коефициенти на корелация. Например, ако променливите са категорийни, се използва коефициент на контингенция. Различните варианти са подробно разгледани в [22].

Мултиколинеарност и премахване на линейно зависими величини

Ако факторите в набор от данни са много силно корелирани, това означава че те са близки до линейно зависими. Но обратното – слаба корелация между факторите (недиагоналните елементи на ρ_φ са близки до 0) не гарантира, че в набора няма линейна зависимост. Ако една променлива е линейна функция например на 9 други променливи, то е напълно възможно корелацията между нея и всяка от 9-те променливи да е слаба ($\rho < 0.3$). Ето защо е необходимо да се проведе специален тест, с който да се открият линейно зависимите променливи. Често поради наличието на неопределености връзките може да не са чисто линейно зависими. Затова се въвежда праг на значимост, над който се приема, че променливите са линейно зависими.

Ако връзката, дефинирана в статистически смисъл, между два (или повече) сигнала клони към линейно зависима, те се наричат колинеарни (мултиколинеарни). От друга страна, ако връзката е линейно зависима (напълно детерминирана), те се наричат напълно колинеарни (напълно мултиколинеарни).

Един метод за откриване на линейно зависими променливи е свързан с факторизация на матрицата на Фишер:

$$F = \Phi^T \Phi,$$

дефинирана в темата за стандартизацията, в точка 2.2.2.4, като стълбовете на Φ (z на брой) са стойностите на факторите за интервала на наблюдение (виж (2.49) и (2.50) в точка 2.2.2.2). Идеята на метода е, че ако факторите не са линейно зависими, то Φ и респективно F е от пълен ранг, но ако $\text{rank } \Phi = r < z$, то и $\text{rank } F = r < z$ (тъй като за произволна матрица A е в сила $\text{rank } A = \text{rank}(A^T A)$). Ако с елементарни преобразувания матрицата F се преобразува в триъгълна, то след това, анализирайки структурата на резултантната матрица, може да се определи рангът на F , равен на този на Φ и на броя на линейно независимите фактори. В случай на открита линейна зависимост от извършените преобразувания може да се изолират фактор/фактори, които са линейно зависими с останалите. Те може да се премахнат от набора, без това да влоши информативността на данните като цяло.

Този анализ напълно може да се автоматизира и затова е имплементиран в някои софтуери за извличане на информация от данни като стандартна процедура, изпълняваща се точно преди оценяването на параметри.

2.2.3.4 Сегментация на данни

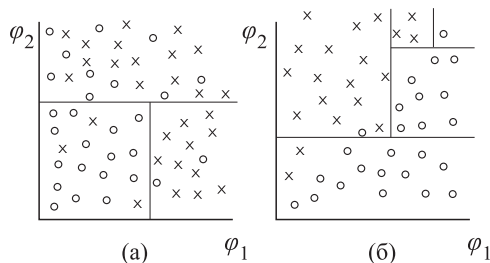
Понякога данните отразяват поведение на системата в различни работни области. Един подход за извличане на информацията от такива данни е представянето на реалното поведение с множество модели (например почастично линейното описание на система с изразено нелинейно поведение). Важна стъпка при този подход е определянето на множеството от модели. Те не трябва да имат твърде сходно поведение, а да са достатъчно различни. В противен случай може да възникнат числени проблеми. От друга страна, ако моделите описват твърде различно поведение, точността на многомоделното описание [40, 142] като цяло намалява.

В кредитната индустрия, когато наборите съдържат данни за разнородни групи от хора (на различна възраст, с различен доход, образование и т.н.), често се прибегва до сегментиране на популацията на подпопулации, като за целта се използват дървета на решенията [72, 138]. Те декомпозират факторното пространство на сегменти, които са различни от гледна точка на поведението на системата. Например хората с висок и с нисък доход се очаква да имат (статистически) различно поведение от гледна точка на погасяването на кредити. Ако е нужно, сегментацията може да се извърши на няколко нива – първо, цялата популация се разделя на подпопулации според възрастта. След това, най-младите кандидати може допълнително да се разделят по образование, тези на средна възраст – на такива с висок, среден или нисък доход, и т.н. Така се получават различни групи от наблюдения, които в рамките на групата отразяват сходно поведение на системата.

Процесът на сегментиране (изграждането на дърво на решенията) може да се автоматизира. Ако има предварителна информация за това – кои характеристики потенциално са подходящи за сегментиране, те могат да се отделят и от тях автоматично да се изберат най-подходящите, според наличните данни.

Предимство на дърветата на решенията са, че те лесно се тълкуват, което улеснява проверката доколко отразяват априорните знания. Също така те не са сложни за употреба (извършват се сравнения на стойности). Но този тип модели има ограничено приложение най-вече поради

това, че сегментирането на текуща група е по определен фактор, както е показано на Фигура 2.32.



Фигура 2.32. Сегментиране при независими (а) и зависими (б) фактори

Основният недостатък на дърветата на решенията е, че те не са подходящи, когато факторите са взаимосвързани. Например, ако два фактора са свързани по следния начин:

$$\varphi_{1,k} = a\varphi_{2,k} + b,$$

както е показано на Фигура 2.32 (б), сегментацията с този подход не би довела до добър резултат. Темата за сегментирането е подробно разглеждана в [72].

Съществуват и други дейности, които се извършват при подготовката на данните за оценяване на параметри и избор на структура на модела. Понякога се използват специфични техники, свързани с конкретното приложение и избраната методология, най-вече от гледна точка на отчитането на априорната информация. Като цяло описаните методи може да се автоматизират, а това е много важно условие при идентификацията на многомерни системи, особено при работата с големи набори от данни.

2.3 Оценяване на параметри и избор на структура на модела

В четвъртия етап на идентификацията, описан в тази точка, се извършва същинското изграждане на модел на системата. Целта е да се определят оптималните, в смисъл на избран критерий, структура и параметри на модела. Обикновено изборът на структурата е свързан с определянето на множество от модели. По тази причина първо се раз-

глежда задачата за намиране на оптималните параметри на модел с фиксирана структура, а след това вниманието се насочва към различни методи за избор на структура му.

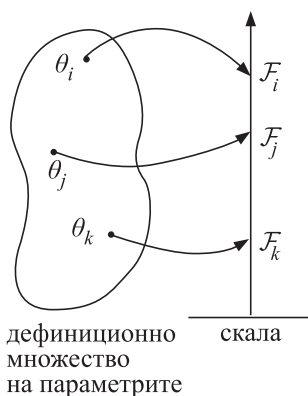
2.3.1 Оценяване на параметри

Определянето на „най-добрите“ параметри на модела е оптимизационна задача. Най-общо оптимизацията е процесът на получаване на най-добър резултат в определен смисъл и при определени условия. Смыслът се изразява с помощта на критерий, който се дефинира като минимум или максимум на някакъв показател. Показателят може да бъде функция или функционал и зависи от определен брой параметри, чиито оптимални стойности се търсят.

Обект на оптимизацията в задачата за оценяване на параметри е моделът на системата, в който се съдържа свобода по отношение на стойностите на определен брой параметри. (Част от тях са структурните параметри, които за момента се приемат за известни.) В долните разглеждания, когато моделът се приема за линейно параметризиран, е използвано представянето в общ вид с вектор на параметрите.

2.3.2 Целева функция

Целевата функция $\mathcal{F}(\theta)$ е математическа зависимост, дефинираща целта на задачата за оптимизация. Тя зависи от търсените параметри и в идентификацията има смисъл на показател на качеството на модела. Обикновено $\mathcal{F}(\theta)$ е скалярна функция, която може да се интерпретира като правило, което на всяка точка от дефиниционното множество на параметрите θ съпоставя определена стойност от скала (Фигура 2.33). От тази гледна точка целта на оптимизацията е намирането на тази точка от дефиниционното множество, на която отговаря минимална стойност от съответната скала.



Фигура 2.33. Целева функция

Най-разпространеният случай е $\mathcal{F}(\theta)$ да е сума от квадратите на остатъка $e_k = y_k - \hat{y}_k$, т.е.

$$\mathcal{F}(\theta) = \sum_k e_k^T e_k = \sum_k (e_{1,k}^2 + e_{2,k}^2 + \dots + e_{\ell,k}^2). \quad (2.63)$$

Когато зависимостта $\hat{y}(\theta)$ (и съответно e_k) е линейна, оптимизацията се свежда до т.нар. задача на най-малките квадрати (Least Squares – LS), а когато $\hat{y}$ е нелинейна функция на θ , се решава задачата на нелинейните най-малки квадрати (Nonlinear LS – NLS).

Ако моделът е линеен по отношение на θ , тогава (2.63) е квадратична функция на параметрите и техните оптимални стойности може да се определят аналитично. От друга страна, при NLS, въпреки че $\mathcal{F}(\theta)$ не е квадратична, нейният градиент и Хесиан (аналозите на първата и втората производна) по отношение на параметрите имат специален вид, което води до възможности за опростяване на оптимизацията. Също така, стойността на $\mathcal{F}(\theta)$ е по-чувствителна към големите (абсолютни) стойности на остатъка, в сравнение с малките. Често това свойство е желано в задачата за оценяване на параметри. Това са основните причини за популярността на тази целева функция.

От (2.63) се вижда, че остатъците по всеки изход участват с еднакво тегло при формирането на $\mathcal{F}(\theta)$. Понякога има нужда елементите на сумата да участват с различни тегла. Например, ако системата е нестационарна, то последните остатъци носят най-актуална информация за несъответствието между системата и модела и затова влиянието им върху $\mathcal{F}(\theta)$ трябва да е по-голямо от това на предишните остатъци.

Нека елементите на вектора w_k са тегла, с които компонентите на e_k участват при формирането на целевата функция, а матрицата $W_k \in \mathcal{R}^{\ell \times \ell}$ е

$$W_k = \text{diag}[w_{1,k} \quad w_{2,k} \quad \dots \quad w_{\ell,k}]^T.$$

Тогава по-общият вид на (2.63) е претеглената сума от квадратите на остатъка:

$$\mathcal{F}(\theta) = \sum_k e_k^T W_k e_k = \sum_k (w_{1,k} e_{1,k}^2 + w_{2,k} e_{2,k}^2 + \dots + w_{\ell,k} e_{\ell,k}^2). \quad (2.64)$$

С въвеждането на теглата в целевата функция при линейно параметризирани модели се получава задачата на претеглените най-малките квадрати (Weighted Least Squares – WLS).

Освен за нестационарни системи, друго приложение на (2.64) е, когато се формира множество от прости модели, описващи различни аспекти

от поведението на сложни системи. Тук, с подходящи стойности на w_k , може да се формира множество от модели, като се използва един и същи набор от данни. При формирането на конкретен модел с теглата се акцентира на съответен аспект на обекта. Пример за такова приложение на (2.64) във финансите е, когато един модел отразява поведението на младите кредитоискатели, а друг – на хората на средна възраст. В случая w_k е функция на фактора ‘възраст’. За модела, описващ младите, теглата намаляват с нарастване на възрастта, а при оценяване на модела за средна възраст, теглата имат максимум в околност на избраната възраст.

Друг вариант е теглата да зависят от качеството на данните. Например в k -тия момент може да е известно, че 20% от регресорите са били липсващи стойности, преди данните да се обработят, или неопределеността в стойностите им е по-голяма от обичайното поради условията на проведения експеримент и т.н. В тези случаи елементите на e_k трябва да влияят по-слабо на $\mathcal{F}(\theta)$, а оттам и на оценките на параметрите. Така задачата за оценяване става по-слабо чувствителна към неопределеността в данните.

В долните разглеждания се използва следното представяне на наличните стойности на изхода и на остатъка. Нека n е броят тактове, отговарящ на максималното закъснение на сигнал, участващ при формирането на изхода на модела. В такъв случай данните от първите n такта са необходими за формиране на началните условия на модела, с които се предсказва първата стойност на изхода му (това е $\hat{y}_{n+1}$). Ако моделът е статичен, то $n = 0$. Нека стойностите на изхода и на остатъка се подредят във векторите y , $e \in \mathcal{R}^{\ell(N-n)}$ по следния начин:

$$y = [y_{n+1}^T \ y_{n+2}^T \ \dots \ y_N^T]^T \text{ и } e = [e_{n+1}^T \ e_{n+2}^T \ \dots \ e_N^T]^T. \quad (2.65)$$

С това представяне, (2.63) и (2.64) може да се запишат съответно като

$$\mathcal{F}(\theta) = e^T e \text{ и } \mathcal{F}(\theta) = e^T W e. \quad (2.66)$$

За целта е въведена тегловната матрица W , която е

$$W = \text{diag}[w_{n+1}^T \ w_{n+2}^T \ \dots \ w_N^T].$$

Когато матрицата W не е диагонална, а $\hat{y}_k$ е линейна функция на параметрите, се получава методът на обобщените най-малки квадрати (General Least Squares - GLS). Разликата между WLS и GLS е описана в точка 2.3.4.4.

Съществуват и други форми на целевата функция (които не са квадратични по отношение на остатъка). Конкретният вид на $\mathcal{F}(\theta)$ е свързан,

както с особеностите на системата, така и с избрания метод за оценяване на параметрите (който зависи най-вече от това, каква информация ще се използва при оценяването на параметрите (виж точка 2.3.4).

2.3.3 Критерий за оптималност

С цел да се унифицира задачата за оптимизация, обикновено се приема, че се търси минимум на целевата функция. Тогава критерият за оптималност ще се записва като:

$$\min_{\theta \in \Omega} \mathcal{F}(\theta).$$

Ω е областта на допустимите стойности на вектора θ , измежду които се търси оптималното решение. Когато параметрите са подредени във вектор, Ω е подмножество или съвпада с множеството на p -мерните реални вектори ($\Omega \subseteq R^p$), а когато параметрите са обединени в $z \times \ell$ мерната матрица Θ , то $\Omega \subseteq R^{z \times \ell}$.

Ако не се налагат ограничения на стойностите на параметрите, то за модел с вектор на параметрите критерият ще се записва така:

$$\min_{\theta} \mathcal{F}(\theta).$$

Ако се търси максимум, например при метода на максималното правдоподобие, разгледан в следващата точка, за унифициране на задачата целевата функция се изменя така, че да се търси минимум. В такъв случай максимизирането на функцията $\tilde{\mathcal{F}}(\theta)$ е еквивалентно на минимизиране на $\mathcal{F}(\theta) = -\tilde{\mathcal{F}}(\theta)$, т.е.

$$\max_{\theta} \tilde{\mathcal{F}}(\theta) \rightarrow \min_{\theta} \mathcal{F}(\theta).$$

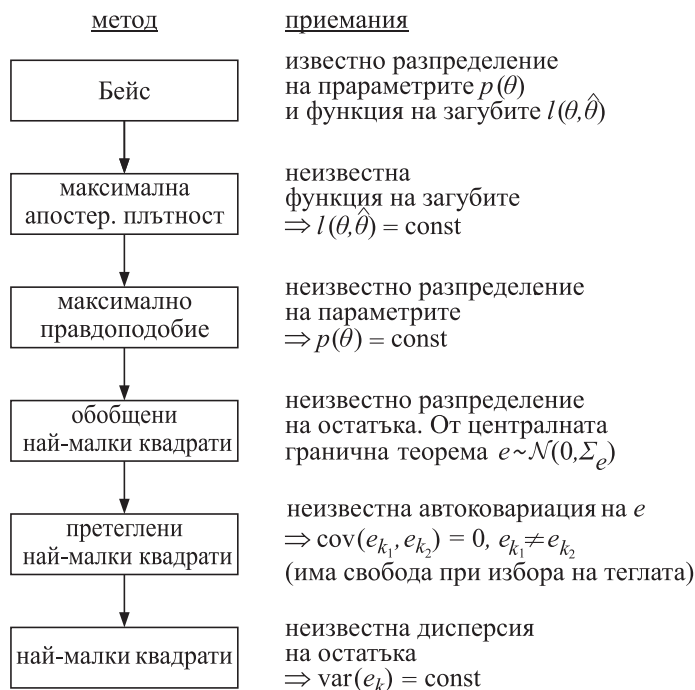
2.3.4 Методи за оценяване на параметри

Често има стремеж видът на целевата функция да е такъв, че оптималните параметри да се изразят като явна функция на данните и оттам директно да се изчислят. Но също така има случаи, когато този опит за опростяване не е удачен или е невъзможен. Тогава решението на проблема за оценяване не е аналитично (еднократно) и поради това, параметрите на модела се търсят числено, с помощта на метод за оптимизация.

В тази точка най-общо са разгледани методите за оценяване на θ в зависимост от наличната априорна информация за изследвания обект,

като не се набляга на конкретни методи за числена оптимизация, с които се реализират оценителите.

На Фигура 2.34 са показани разгледаните по-долу методи и приеманията при формулирането им. Опростяването, породено най-вече от непълната априорна информация за обекта, е за сметка на отдалечаване от оптималното решение. Въпреки това, както ще стане ясно, много често оценителите, които използват малко предварителни знания, напълно задоволяват практическите изисквания към моделите.



Фигура 2.34. Бейсовият метод е най-общ, но изисква най-много априорна информация. Колкото по-малко е тя, толкова повече предположения се правят и така се стига до съответните частни случаи.

2.3.4.1 Метод на Бейс

Най-общият метод за оценяване на параметрите е този на Бейс. Той се базира на следната постановка. От гледна точка на изследваната система търсените параметри може да са напълно детерминирани величини, зависещи от спецификата на процесите, но поради непълните

априорни знания и неопределеностите, свързани с даден експеримент, стойностите им трябва да се приемат с определена сигурност. В основата на Бейсовия подход (от гледна точка на идентификацията) е приемането, че както данните съдържат неопределеност и следователно имат вероятностен характер, така и търсените параметри са случайни величини. В този смисъл вероятностната плътност на разпределение на θ може да се интерпретира като мярка за сигурността по отношение на възможните стойности на параметрите [39]. Това позволява връзката между неизвестните параметри и данните да се изрази с помощта на теорията на вероятностите. Затова се въвежда условната плътност [16] на разпределение на параметрите – при условие че данните за реакцията на системата са налични. Тя ще се бележи с $p(\theta|y)$ и се дефинира по отношение на y , тъй като към него се стреми зависещият от параметрите изход на модела.

Данните съдържат апостериорна информация за системата. От формулата на Бейс:

$$p(\theta|y)p(y) = p(y, \theta) = p(y|\theta)p(\theta),$$

условната апостериорна плътност се изразява по следния начин:

$$p(\theta|y) = \frac{p(y|\theta)p(\theta)}{p(y)}. \quad (2.67)$$

Тъй като разпределението $p(y)$ не зависи от θ , то не влияе на търсената оценка $\hat{\theta}$. От друга страна, от съществено значение за процеса на оценяване е връзката

$$p(\theta|y) \propto p(y|\theta)p(\theta) \quad (2.68)$$

(символът ' $\propto$ ' означава, че зависимостта е пропорционална), която се използва при конструирането на целевата функция. Величината $p(\theta)$ е априорната плътност на разпределение на параметрите (преди получаването на данните), която в метода на Бейс се приема, че е известна. Тя се задава от изследователя и зависи от предварителната информация за разпределението на θ .

Ако $p(\theta)$ не е известна, то с определени приемания оценяването на параметрите се свежда до прилагането на някои от методите (частни случаи на Бейсовия), описани в следващите точки. Освен $p(\theta)$ в (2.68) участва условната плътност $p(y|\theta)$, която също трябва да е известна. Тя отразява близостта на поведението на модел с параметри θ до това на системата. На практика конкретна стойност на $p(y|\theta)$ е вероятността за

наблюдаване на данните при избраните параметри на модела. Така задачата се довежда до варианта, при който еднократно или итеративно (с метод за оптимизация) параметрите на модела се оценяват на базата на фиксиран набор от данни. Зависимостта $p(y|\theta)$ се нарича още функция на правдоподобие. От тази гледна точка $p(y|\theta)$ описва доколко правдоподобни са наблюдаваните данни при конкретни стойности на параметрите.

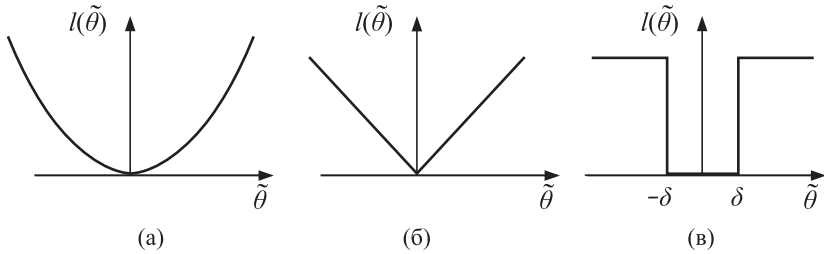
Изразите (2.67) или (2.68) може да се интерпретират като правило за преобразуване на априорното знание за параметрите (плътността $p(\theta)$) в апостериорно (плътността $p(\theta|y)$), като се извлича информация от достъпните измервания. Колкото по-ниско е нивото на неопределеността в данните или колкото повече измервания са достъпни, толкова по-точно може да се определи връзката на данните от параметрите, т.е. функцията на правдоподобие $p(y|\theta)$ да е с по-остра форма. Това означава (съгласно (2.68)), че и $p(\theta|y)$ е по-тясна и стръмна, и съответно по-точно се определят параметрите на модела. Подобен извод може да се направи, ако предварителната информация за параметрите е по-точна и тогава $p(\theta)$ е тясна и стръмна. От друга страна, при по-неинформативни данни $p(y|\theta)$ (или малко априорни знания, $p(\theta)$) и $p(\theta|y)$ са по-разлети и оценката $\hat{\theta}$ става по-неточна.

Постановката на задачата за оценяване е най-обща при метода на Бейс, тъй като зависи от най-малко приемания, които да ограничат общността, но цената за това е нуждата от най-много априорни знания.

Целевата функция по метода на Бейс зависи от $p(\theta|y)$ (или съгласно (2.68)) от $p(\theta)$ и $p(y|\theta)$. За пълното дефиниране на $\mathcal{F}(\theta)$ първоначално в разглежданията се въвежда т.нар. функция на загубите $l(\theta, \hat{\theta})$. Тя зависи от параметрите на текущия модел, отразява загубите от разликата между тях и оптималните им стойности и практически се избира от изследователя. $p(\theta|y)$ допълнително се претегля с $l(\theta, \hat{\theta})$. Наличието на функцията на загубите в $\mathcal{F}(\theta)$ е всъщност това, което отличава Бейсовия метод от метода на максималната апостериорна плътност (разгледан в следващата точка). Тук $\mathcal{F}(\theta)$ има вида:

$$\mathcal{F}(\theta) = \int_{-\infty}^{\infty} l(\theta, \hat{\theta}) p(\theta|y) d\theta = \int_{-\infty}^{\infty} l(\theta, \hat{\theta}) p(y|\theta) p(\theta) d\theta. \quad (2.69)$$

За да се изясни ролята на $l(\theta, \hat{\theta})$, ще се определи оптималното решение $\hat{\theta}$ при различни нейни варианти.



Фигура 2.35. Функция на загубите: квадратична (а), абсолютна (б) и режекторна (в), където $\tilde{\theta} = \theta - \hat{\theta}$

Квадратична функция на загубите

Може би най-често използваната функция на загубите е квадратичната (виж Фигура 2.35 (а)):

$$l(\theta, \hat{\theta}) = \|\theta - \hat{\theta}\|_2^2 = (\theta - \hat{\theta})^T (\theta - \hat{\theta}). \quad (2.70)$$

При този вариант целевата функция е

$$\mathcal{F}(\tilde{\theta}) = \int_{-\infty}^{\infty} \|\theta - \hat{\theta}\|_2^2 p(\theta|y) d\theta.$$

С използване на правилото на Лайбниц за диференциране на интеграли:

$$\nabla_x \int_{h_1(x)}^{h_2(x)} f(x, y) dy = \int_{h_1(x)}^{h_2(x)} \nabla_x f(x, y) dy + f(h_2, y) \frac{\partial h_2}{\partial y} - f(h_1, y) \frac{\partial h_1}{\partial y},$$

за градиента ѝ се получава

$$\begin{aligned} \nabla \mathcal{F}(\theta) &= -2 \int_{-\infty}^{\infty} (\theta - \hat{\theta}) p(\theta|y) d\theta \\ &= -2 \int_{-\infty}^{\infty} \theta p(\theta|y) d\theta + 2\hat{\theta} \int_{-\infty}^{\infty} p(\theta|y) d\theta \\ &= -2 \int_{-\infty}^{\infty} \theta p(\theta|y) d\theta + 2\hat{\theta}. \end{aligned}$$

В оптимума той се нулира и за търсените параметри $\hat{\theta}$ се получава:

$$\hat{\theta} = \int_{-\infty}^{\infty} \theta p(\theta|y) d\theta = \text{mean } p(\theta|y), \quad (2.71)$$

т.е. при квадратична функция на загубите оптималната оценка на параметрите съвпада с математическото очакване на разпределението. Казано по друг начин, когато целевата функция е (2.69), $\hat{\theta}$ е условното

математическо очакване на θ , при условие че са наблюдавани данните от експеримента.

Абсолютна функция на загубите

Друг вариант за функцията на загубите е

$$l(\theta, \hat{\theta}) = \|\theta - \hat{\theta}\|_1 = \sum_{i=1}^p |\theta_i - \hat{\theta}_i|.$$

(Фигура 2.35(б)). Целевата функция може да се развие така:

$$\begin{aligned} \mathcal{F}(\theta) &= \int_{-\infty}^{\infty} l(\theta, \hat{\theta}) p(\theta|y) d\theta \\ &= \int_{-\infty}^{\infty} \|\theta - \hat{\theta}\|_1 p(\theta|y) d\theta \\ &= \sum_{i=1}^p \int_{-\infty}^{\hat{\theta}_i} (\hat{\theta}_i - \theta_i) p(\theta_i|y) d\theta_i + \sum_{i=1}^p \int_{\hat{\theta}_i}^{\infty} (\theta_i - \hat{\theta}_i) p(\theta_i|y) d\theta_i. \end{aligned}$$

След диференциране и приравняване $\nabla \mathcal{F}(\theta)$ на нула се получава:

$$-\sum_{i=1}^p \int_{-\infty}^{\hat{\theta}_i} p(\theta_i|y) d\theta_i + \sum_{i=1}^p \int_{\hat{\theta}_i}^{\infty} p(\theta_i|y) d\theta_i = 0.$$

Равенството се удовлетворява, когато площите на $p(\theta_i|y)$ отляво и отясно на $\hat{\theta}_i$ са равни, т.е. всяка оценка отговаря на медианата на съответното едномерно разпределение $p(\theta_i|y)$. За оптималните оценки $\hat{\theta}$ се получава:

$$\hat{\theta} = \text{median } p(\theta|y).$$

Режекторна функция на загубите

Трети случай за $l(\theta, \hat{\theta})$ е

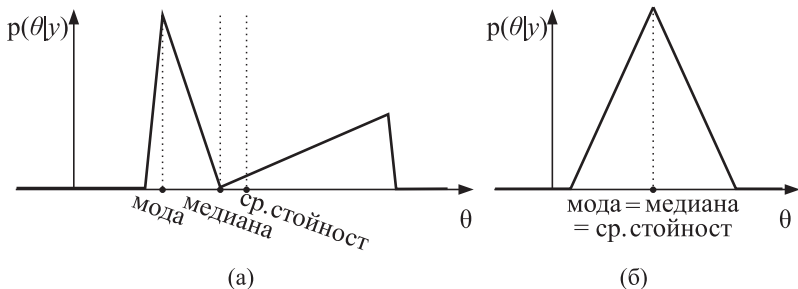
$$l(\theta, \hat{\theta}) = \sum_{i=1}^p l_i, \text{ където } l_i = \begin{cases} 0, & \text{за } |\theta_i - \hat{\theta}_i| < \delta, \\ 1, & \text{за } |\theta_i - \hat{\theta}_i| \geq \delta, \end{cases}$$

(дадена на Фигура 2.35(в)). При този вариант

$$\begin{aligned} \mathcal{F}(\theta) &= \sum_{i=1}^p \int_{-\infty}^{\hat{\theta}_i - \delta} 1 p(\theta_i|y) d\theta_i + \sum_{i=1}^p \int_{\hat{\theta}_i + \delta}^{\infty} 1 p(\theta_i|y) d\theta_i \\ &= 1 - \sum_{i=1}^p \int_{\hat{\theta}_i - \delta}^{\hat{\theta}_i + \delta} p(\theta_i|y) d\theta_i, \end{aligned}$$

т.е. минимизирането на $\mathcal{F}(\theta)$ е равносилно на максимизиране на интегралите $\int_{\hat{\theta}_i - \delta}^{\hat{\theta}_i + \delta} p(\theta_i|y) d\theta_i$ за $i = \overline{1, p}$. Когато $\delta \rightarrow 0$, оценката $\hat{\theta}_i$ отговаря на максималната стойност (модата) на $p(\theta_i|y)$ и в такъв случай за оценката на вектора на параметрите се получава:

$$\hat{\theta} = \text{mode } p(\theta|y).$$



Фигура 2.36. Несиметрични разпределения с повече от един екстремум: средна стойност, медиана и мода не съвпадат (а), а при симетрични – съвпадат (б).

Когато разпределението е симетрично и унимодално, статистическите характеристики: средна стойност, медиана и мода съвпадат (Фигура 2.36 (б)). Но в общия случай, както е показано на Фигура 2.36 (а), такова съвпадение не се наблюдава.

От представените варианти за функцията на загубите се вижда, че с подходящ избор на $l(\theta, \hat{\theta})$, се постигат желани свойства на оценките.

В следния пример се онагледява приложението на метода на Бейс.

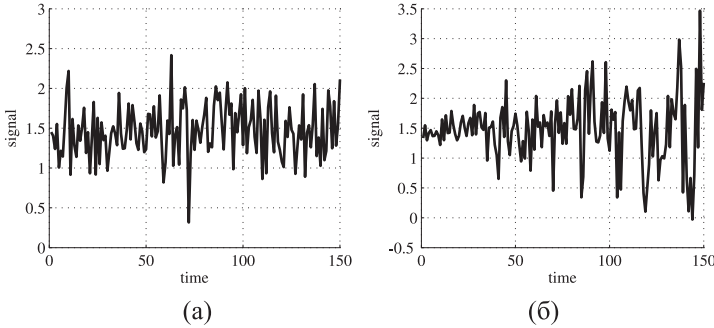
Пример. Метод на Бейс

Нека за описание на поведението на ММО система е избран модел, представен в линейно параметризиран вид с вектор на параметрите, т.е. $\hat{y}_k = \Phi_k \theta$ и предварително е известно, че разпределението на параметрите е нормално $\theta \sim \mathcal{N}(\theta_a, \Sigma_{\theta_a})$. Величините $\theta_a \in \mathcal{R}^p$ и $\Sigma_{\theta_a} = \text{cov } \theta \in \mathcal{R}^{p \times p}$ може да са изведени или определени на базата на данни от предишен експеримент. Тази априорна информация се изразява с плътността:

$$p(\theta) = \frac{1}{\sqrt{2\pi \det \Sigma_{\theta_a}}} e^{-\frac{1}{2}(\theta - \theta_a)^T \Sigma_{\theta_a}^{-1}(\theta - \theta_a)}. \quad (2.72)$$

Плътността $p(y|\theta)$ отразява близостта на поведението на модел с параметри θ до това на системата, или казано по друг начин: опис-

ва вероятността изходът $\hat{y}_k$ на модела да повтори изхода y_k на системата. Несъвпадението между тези два сигнала се изразява с остатъка $e_k = y_k - \hat{y}_k$. Целта на оценяването на параметри е да се построи модел,



Фигура 2.37. Ергодичен (хомоскедастичен) (а) и неергодичен (хетероскедастичен) (б) сигнал

който да опише детерминираната съставка в y_k , и така остатъкът да съдържа само случайната съставка. За по-голяма общност в долните разглеждания ще се приеме, че e_k е неергодичен, т.е. е с променливи статистически характеристики, както е показано на Фигура 2.37 където $\Sigma_{e,k} = \text{cov } e_k$ зависи от времето. Нека също така структурата на модела да осигурява намирането на описание, което отразява детерминираната съставка в y_k . Тогава ковариационната матрица на остатъка за различни моменти от времето k_1 и k_2 е $\text{cov}(e_{k_1}, e_{k_2}) = 0$.

При фиксиран набор от N независими наблюдения апостериорната плътност (функцията на правдоподобие), зависеща от стойностите на параметрите, е

$$p(y|\theta) = \prod_{k=n+1}^N \frac{1}{(2\pi)^{\ell/2} \det \Sigma_{e,k}^{1/2}} e^{-\frac{1}{2}(y_k - \hat{y}_k)^T \Sigma_{e,k}^{-1} (y_k - \hat{y}_k)}. \quad (2.73)$$

Нека по аналогия с (2.65) стойностите на изхода на модела се подредят във вектора $\hat{y} \in \mathcal{R}^{\ell(N-n)}$ по следния начин:

$$\hat{y} = [\hat{y}_{n+1}^T \quad \hat{y}_{n+2}^T \quad \dots \quad \hat{y}_N^T]^T,$$

а във $\Phi \in \mathcal{R}^{\ell(N-n) \times p}$ се подредят матриците на регресорите, т.е.

$$\Phi = [\Phi_{n+1}^T \quad \Phi_{n+2}^T \quad \dots \quad \Phi_N^T]^T.$$

Тогава за интервала на наблюдение е в сила $\hat{y} = \Phi\theta$ и $p(y|\theta)$ може да се запише във вида:

$$p(y|\theta) = \frac{1}{(2\pi)^{\frac{\ell N}{2}} \det \Sigma_e^{\frac{1}{2}}} e^{-\frac{1}{2}(y-\Phi\theta)^T \Sigma_e^{-1}(y-\Phi\theta)}, \quad (2.74)$$

като $\Sigma_e \in \mathcal{R}^{p\ell(N-n) \times \ell(N-n)}$ е блокдиагоналната матрица:

$$\Sigma_e = \text{diag}(\Sigma_{e,n+1}, \Sigma_{e,n+2}, \dots, \Sigma_{e,N}).$$

Тя е диагонална, ако елементите на остатъка в един и същи момент от времето не са корелирани помежду си, т.е.

$$\Sigma_{e,k} = \text{diag}(\sigma_{e,1,k}^2, \sigma_{e,2,k}^2, \dots, \sigma_{e,\ell,k}^2).$$

Ако съществува корелация между стойностите на остатъка в различни моменти от времето, тогава матрицата Σ_e^{-2} ще има ненулеви елементи извън блоковете по диагонала, отговарящи на ковариациите $\text{cov}(e_{k_1}, e_{k_2})$. Тези случаи са разгледани по-подробно в точка 2.3.4.4.

Тъй като в (2.72) и (2.74) параметрите участват единствено в степените на експонентите, то дробните членове (константи) не влияят на оптималното решение. Затова от гледна точка на оценяването са важни зависимостите:

$$p(\theta) \propto e^{-\frac{1}{2}(\theta-\theta_a)^T \Sigma_{\theta_a}^{-1}(\theta-\theta_a)} \text{ и } p(y|\theta) \propto e^{-\frac{1}{2}(y-\Phi\theta)^T \Sigma_e^{-1}(y-\Phi\theta)}.$$

От (2.72) и (2.74) лесно може да се види, че тъй като двете плътности са на нормално разпределени величини, то и $p(\theta|y)$ също отговаря на нормално разпределение, като

$$p(\theta|y) \propto p(y|\theta)p(\theta) \propto e^{-\frac{1}{2}(\theta-\theta_a)^T \Sigma_{\theta_a}^{-1}(\theta-\theta_a) - \frac{1}{2}(y-\Phi\theta)^T \Sigma_e^{-1}(y-\Phi\theta)}. \quad (2.75)$$

За дефинирането на $\mathcal{F}(\theta)$ по метода на Бейс е нужно да се конкретизира и функцията на загубите. Нека е избрана квадратичната функция:

$$l(\theta, \hat{\theta}) = \|\theta - \hat{\theta}\|_2^2.$$

По-горе беше показано, че при нея оптимумът на $\mathcal{F}(\theta)$ (виж (2.71)) се достига при

$$\hat{\theta} = \text{mean } p(\theta|y).$$

Всъщност, тъй като θ е с нормално разпределение, то $p(\theta|y)$ е симетрична и унимодална. Тогава математическото очакване, медианата и модата съвпадат, т.е. $\hat{\theta}$ съвпада за разгледаните варианти на $l(\theta, \hat{\theta})$.

Максимумът на $p(\theta|y)$ и респективно на $p(y|\theta)p(\theta)$ се достига, когато и степента на експонентата е максимална. След разкриване на скобите и известни преобразувания в степента на (2.75) се получава полиномът

$$f(\theta) = \theta^T A \theta + b^T \theta + c.$$

A е отрицателно определена, симетрична матрица и ако постановката на задачата е коректна, нейната обратна матрица съществува. Тогава математическото очакване на $p(\theta|y)$, което е и оптималната оценка, отговаря на максимума на $f(\theta)$. За квадратичната функция стойността на θ в максимума ѝ е $\hat{\theta} = -\frac{1}{2}A^{-1}b$, а ковариационната матрица на получената оценка е $\Sigma_{\hat{\theta}} = -\frac{1}{2}A^{-1}$. $\diamond$

2.3.4.2 Метод на максималната апостериорна плътност

Този метод е частен случай на метода на Бейс. При него отново се задава априорното знание за параметрите с помощта на разпределението $p(\theta)$ и след това тази информация се актуализира (според формулата на Бейс) с използване на апостериорните данни. За разлика от предишния случай се приема, че $l(\theta, \hat{\theta}) = \text{const}$, и така целевата функция при метода на максималната апостериорна плътност (Maximum A-Posteriori – MAP) е

$$\mathcal{F}(\theta) = p(\theta|y),$$

а критерият съответно е

$$\max_{\theta} p(\theta|y) = \max_{\theta} p(y|\theta)p(\theta).$$

Нека априорната информация за параметрите са първоначалните стойности θ_a , както и точността, с която всеки параметър е определен (т.е. има информация за дисперсиите на компонентите $\theta_{a,i}$). Това може да се изрази със задаването на $p(\theta)$ със средна стойност θ_a и различна острота по отношение на всеки параметър (на по-неточните параметри отговарят по-разлети плътности). По този начин MAP, а и Бейсовият метод отчитат априорната информация – параметрите, за които е характерна по-голяма априорна неопределеност, са по-чувствителни към данните, а тези, които са известни с по-голяма точност, са по-чувствителни към предварителната информация. Тази особеност може да се интерпретира като т.нар. регуляризация на решението, разгледана по-подробно в точка 3.4.3.

В следващия пример се решава задачата за определяне на апостериорните оценки на параметрите на модел с идеята, заложена в MAP.

Пример. Метод на максималната апостериорна плътност

Нека за графичното представяне на ефекта от комбинираното използване на априорната и апостериорната информация параметърът на модела от предишния пример е един, като предварителната информация за него е $\theta_a = 5$, а неопределеността при определянето на тази стойност (например от предишен експеримент) е изразена с дисперсията $\sigma_{\theta_a}^2 = 1$. Тогава

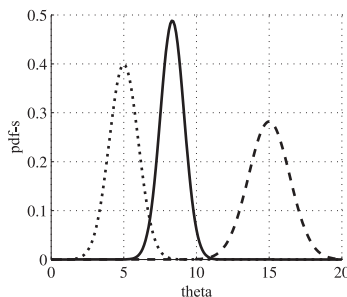
$$p(\theta) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(\theta-5)^2}. \quad (2.76)$$

Нека е проведено едно наблюдение ($N = 1$), при което $\varphi_1 = 1$ и $y_1 = 15$, а дисперсията на неопределеността в системата, приведена към изхода, е $\sigma_e^2 = 2$. За функцията на правдоподобие се получава:

$$p(y|\theta) = \frac{1}{\sqrt{4\pi}} e^{-\frac{1}{4}(15-\theta)^2}. \quad (2.77)$$

Тогава, с пренебрегване на членовете, които не влияят на оптималното решение, за целевата функция се получава:

$$\mathcal{F}(\theta) = e^{-\frac{1}{2}(\theta-5)^2 - \frac{1}{4}(15-\theta)^2} = e^{-\frac{1}{4}(3\theta^2 - 50\theta + 275)} = e^{-0.75\theta^2 + 12.5\theta - 68.75}.$$



Фигура 2.38. Априорна плътност на θ (линия от точки), функция на правдоподобие $p(y|\theta)$ (прекъсната линия) и актуализирана апостериорна плътност на θ (плътна линия)

На Фигура 2.38 са представени разпределенията (2.76), (2.77) и разпределението на апостериорната оценка $\hat{\theta}$, която максимизира $\mathcal{F}(\theta)$. Стойността на $\hat{\theta}$, както и нейната дисперсия $\sigma_{\hat{\theta}}^2$ отново може да се определят аналитично. За конкретния случай се получава:

$$\hat{\theta} = 8.3(3) \text{ и } \sigma_{\hat{\theta}}^2 = 0.6(6).$$

◇

В примера е представена актуализацията на априорната информация $(\theta_a \text{ и } \sigma_{\theta_a}^2)$ с апостериорните данни, съдържащи се в $\{\varphi_1, y_1\}$. Тази актуализация може да се повтори за повече двойки $\{\varphi_k, y_k\}$, като всеки път преди прилагането на Бейсовото правило (2.68) апостериорната информация се приема за априорна $(\theta_a \leftarrow \hat{\theta} \text{ и } \sigma_{\theta_a}^2 \leftarrow \sigma_{\hat{\theta}}^2)$. Ако $N > 1$, крайната апостериорна плътност не зависи от реда, в който се отчитат $\{\varphi_k, y_k\}$, както и дали те се обединяват на групи.

Моделът в примера е статичен. В общия случай, ако векторът на регресорите изисква измервания от предишни (най-много n) дискретни моменти и y_k е вектор, то за формиране на двойката $\{\Phi_1, y_1\}$ са нужни $N = n + 1$ наблюдения.

2.3.4.3 Метод на максималното правдоподобие

В много практически задачи липсва предварителна информация за плътността на разпределение на оценките $p(\theta)$. В такъв случай е удачно да се приеме, че всяка стойност на параметър преди отчитането на информацията в данните е еднакво вероятна, т.е. $p(\theta_i) = \text{const}$, за $i = \overline{1, p}$. Тогава априорната плътност $p(\theta)$ престава да влияе на оптимума на $\mathcal{F}(\theta)$ и MAP се свежда до метода на максималното правдоподобие (Maximum Likelihood – ML). С отпадането на нуждата от априорно знание за параметрите основният източник на информация остават данните. Така методът ML води до получаване на оценки, които максимизират вероятността да се наблюдават наличните данни. Казано по друг начин, $\hat{\theta}$, получена по метода ML, е оценката на параметрите, при която y е „най-правдоподобната“ реакция на системата. Целевата функция се свежда до $\mathcal{F}(\theta) = p(y|\theta)$, а критерият на ML е

$$\max_{\theta} p(y|\theta).$$

Нека $p(y_k|\theta)$ е условната плътност на вероятността – изходът на модела да приеме наблюдаваната стойност y_k , при условие че параметрите на модела са θ . Логично е, колкото по-неподходящи са оценките на параметрите, толкова по-малка да е вероятността да се наблюдава измереният изход y_k . Освен това, поради наличието на неопределеността e_k , максималната вероятност е по-малка от 1. Но колкото по-подходяща е структурата на модела или влиянието на околната среда и шума от измерване са по-незначителни (т.е. колкото нивото на e_k е по-ниско), толкова максимумът на $p(y_k|\theta)$ ще се доближава до 1.

Величината $p(y_k|\theta)$ е приносът на k -тото наблюдение към функцията на правдоподобие $p(y|\theta)$. По-долу, за да се наблегне на това, че функци-

ята на правдоподобие зависи от стойностите на параметрите, а данните са фиксирани, вместо $p(y_k|\theta)$ и $p(y|\theta)$ ще се въведат означенията $L_{\theta,k}$ и L_θ , които често се използват в литературата.

Нека от входно-изходните данни се формират $N - n$ на брой ($k = 1, N - n$), независими една от друга двойки $\{\varphi_k, y_k\}$. Тогава връзката между целевата функция L_θ и приносите $L_{\theta,k}$ е

$$L_\theta = \prod_{k=n+1}^N L_{\theta,k}. \quad (2.78)$$

В [38, 37] са разгледани показатели, които са оптимални за различни разпределения на остатъка. Целевата функция във вида (2.78) не е удобна за оптимизация особено когато се изчисляват производни на L_θ по отношение на параметрите. Например, ако се изчислява градиентът на L_θ при два приноса, се получава:

$$\nabla L_\theta = \nabla(L_{\theta,1}L_{\theta,2}) = \nabla L_{\theta,1}L_{\theta,2} + L_{\theta,1}\nabla L_{\theta,2},$$

което за произведение от $N - n$ на брой функции води до значително усложняване на изчислителния процес. От друга страна, логаритмуването (а то е монотонна трансформация и не влияе на оптимума) на (2.78) води до преобразуването на произведението от приносите $L_{\theta,k}$ в сумата

$$\ln L_\theta = \sum_{k=1}^{N-n} \ln L_{\theta,k}.$$

Диференцирането на този израз при два приноса води до:

$$\nabla \ln L_\theta = \nabla(\ln L_{\theta,1} + \ln L_{\theta,2}) = \nabla \ln L_{\theta,1} + \nabla \ln L_{\theta,2},$$

като при голям брой наблюдения нараства единствено броят на събираемите в горната сума, без те да се усложняват. Освен това, за да се преобразува задачата за оптимизация в унифициран вид (търсене на минимум), за целева функция в ML се използва $\mathcal{F}(\theta) = -\ln L_\theta$. Повечето статистически софтуери, в които е реализиран ML, визуализират стойностите на функцията:

$$\mathcal{F}(\theta) = -2 \ln L_\theta. \quad (2.79)$$

В следния пример се онагледява приложението на ML.

Пример. *Метод на максималното правдоподобие*

Нека отново се разгледа предишният пример, като се приеме, че липсва информация за $p(\theta)$. С други думи е проведено едно наблюдение,

при което $\varphi = 1$ и $y = 15$, а $\sigma_e^2 = 2$. За функцията на правдоподобие се получава:

$$L_\theta = \frac{1}{\sqrt{4\pi}} e^{-\frac{1}{4}(15-\theta)^2}.$$

Тогава изменената целева функция (2.79) е

$$\mathcal{F}(\theta) = -2 \ln L_\theta = \ln(4\pi) + \frac{1}{2}(15 - \theta)^2 \propto (15 - \theta)^2.$$

За ML оценката, която минимизира $\mathcal{F}(\theta)$, и за нейната дисперсия се получава:

$$\hat{\theta}_{\text{ML}} = 15, \quad \sigma_{\hat{\theta}, \text{ML}}^2 = 2.$$

За сравнение с предишния случай, където априорната информация е $\theta_a \sim \mathcal{N}(5, 1)$, т.е.

$$\theta_a = 5, \sigma_{\theta_a}^2 = 1,$$

по метода MAP е определено:

$$\hat{\theta}_{\text{MAP}} \approx 8.333, \quad \sigma_{\hat{\theta}, \text{MAP}}^2 \approx 0.667.$$

Вижда се, че ако има предварителна информация за параметъра на модела, нейното използване води до по-точна оценка ($\sigma_{\hat{\theta}, \text{MAP}}^2 < \sigma_{\hat{\theta}, \text{ML}}^2$), а оценката $\hat{\theta}_{\text{MAP}}$, отчитаща двата типа информация, е приела компромисна стойност между априорната и апостериорната оценка. $\diamond$

2.3.4.4 Метод на обобщените и на претеглените най-малки квадрати

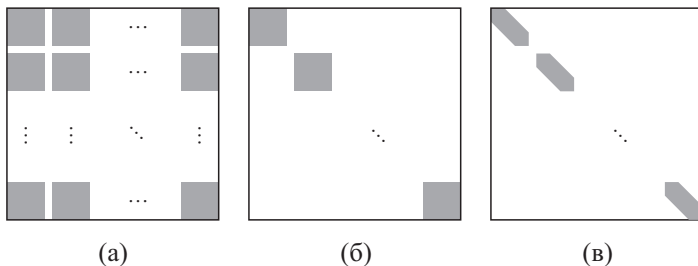
С ML се получава оценка, която максимизира правдоподобие на наблюдаваните входно-изходни стойности за интервала на наблюдение. Неудобството при този метод е, че функцията на правдоподобие трябва да е известна, а в много приложения разпределението на изхода (с въведеното означение това е L_θ) е неизвестно. По-долу се приема, че разпределението е нормално. В точка 2.4.1 се обсъжда по-подробно това приемане. Тогава функцията на правдоподобие (2.73) може да се запише така:

$$L_\theta = \frac{1}{(2\pi)^{\frac{\ell N}{2}} \det \Sigma_e^{\frac{1}{2}}} e^{-\frac{1}{2} e^T \Sigma_e^{-1} e}. \quad (2.80)$$

В общия случай, когато моделът е нелинеен,

$$\hat{y} = f(\Phi, \theta), \quad (2.81)$$

при нормално разпределение на остатъка L_θ може да се запише във вида (2.80). Нека e_k е хетероскедастична величина (ковариационната матрица $\Sigma_{e,k}$ се изменя във времето). Също така нека първо се разгледа общият случай, когато съществува корелация между стойностите на остатъка. Единият вариант е тя да е между компонентите на e_k в един и същи момент т.е. $\text{cov}(e_{i,k}, e_{j,k}) \neq 0$. Другият вариант е корелацията да



Фигура 2.39. Структура на Σ_e^2 при хетероскедастичен остатък и корелация в различни моменти от времето – пълна матрица (а), корелация в текущия момент – блокдиагонална (б) и липса на корелация – диагонална (в). Сивите области съдържат ненулеви стойности.

е между стойностите на остатъка в различни моменти от времето т.е. $\text{cov}(e_{k_1}, e_{k_2}) \neq 0$, като това означава, че наличната входно-изходна информация не е достатъчна за описанието на изследваната система или структурата на модела не е подходяща.

Независимо от корелацията, щом тя съществува, ковариационната матрица за интервала на наблюдение Σ_e^2 не е диагонална и има структура, показана на Фигура 2.39 (а) или (б). Тъй като моделът (2.81) е нелинеен, оценяването на параметрите се извършва по метода на нелинейните обобщени най-малки квадрати (Nonlinear Generalized LS – NGLS). От друга страна, ако между елементите на остатъка не съществуват споменатите корелации (или са пренебрежими), то Σ_e е диагонална (Фигура 2.39 (в)) и задачата за намиране на $\hat{\theta}$ се свежда до метода на нелинейните претеглени най-малки квадрати (Nonlinear Weighted LS – NWLS).

По същите съображения като в предната точка при NGLS и NWLS, вместо да се максимизира L_θ , се използва целевата функция $\mathcal{F}(\theta) = -2\ln L_\theta$, на която се търси минимум (което не променя оптималното решение). За нормално разпределение, с използване на (2.80), $\mathcal{F}(\theta)$ може

да се развие по следния начин:

$$\mathcal{F}(\theta) = -2 \ln L_\theta = \ln((2\pi)^{\ell N} \det \Sigma_e) + e^T \Sigma_e e \propto e^T \Sigma_e e.$$

Нека се положи $W = \Sigma_e^{-1}$. Така, вместо по-общата целева функция в ML, в методите NGLS и NWLS се използва (2.66):

$$\mathcal{F}(\theta) = e^T W e. \quad (2.82)$$

При NGLS, поради наличието на недиагонални елементи в тегловната матрица, $\mathcal{F}(\theta)$ е претеглена сума от произведенията $e_{i,k_1} e_{j,k_2}$ (e_{i,k_1} е i -тата компонента на остатъка (ℓ -мерен вектор) в момент k_1). При NWLS, поради диагоналната форма на тегловната матрица, претеглената сума е само от квадратите на остатъците $e_{i,k}^2$.

До момента се приема, че тегловната матрица е инверсната на ковариационната матрица на остатъка. Има и други варианти за избор на W , част от които бяха обсъдени в точка 2.3.2.

В точка 2.3.4.1 е показано, че ако моделът е линеен и представен с вектор на параметрите, т.е. за интервала на наблюдение $\hat{y} = \Phi\theta$, L_θ може да се запише във вида:

$$L_\theta = \frac{1}{(2\pi)^{\frac{\ell N}{2}} \det \Sigma_e^{\frac{1}{2}}} e^{-\frac{1}{2}(y - \Phi\theta)^T \Sigma_e^{-1} (y - \Phi\theta)}. \quad (2.83)$$

Ако тегловната матрица е диагонална и няма смисъл на обратна на ковариационната матрица на остатъка, вместо Σ_e в (2.82) се използва общият запис W . При извършване на същите разглеждания като тези за нелинеен модел, когато тегловната матрица е недиагонална, се стига съответно до метода на обобщените най-малки квадрати (GLS), а при диагонална матрица W – задачата за намиране на $\hat{\theta}$ се свежда до метода на претеглените най-малки квадрати (WLS). Пропускането на думата „нелинеен“ в името на метода означава, че моделът, за който се прилага съответният метод, е линеен по параметри. От друга страна, дори моделът да е нелинеен като връзка между първоначалните входно-изходни величини, ако е приведен в линейно параметризиран вид, методите за оценяване на параметри отново са GLS и WLS.

За целевата функция във вида $\mathcal{F}(\theta) = -2 \ln L_\theta$ (с използване на означението W) е в сила:

$$-2 \ln L_\theta \propto (y - \Phi\theta)^T W (y - \Phi\theta). \quad (2.84)$$

При тази постановка на задачата (отговаряща на GLS и WLS) е възможно оценката $\hat{\theta}$, минимизираща (2.84), да се определи аналитично

(еднократно – без да е необходима итеративна процедура). Когато моделът е линеен, остатъкът е също зависи линейно от параметрите и следователно целевата функция (десният израз в (2.84)) е квадратична по отношение на θ . Това означава, че $\mathcal{F}(\theta)$ е унимодална и има минимум, когато $\nabla \mathcal{F}(\theta) = 0$. Нека $\mathcal{F}(\theta)$ от (2.84) се развие по следния начин:

$$\mathcal{F}(\theta) = (y - \Phi\theta)^T W (y - \Phi\theta) = y^T W y - 2\theta^T \Phi^T W y - \theta^T \Phi^T W \Phi \theta.$$

Тогава, след диференциране и приравняване на градиента $\nabla \mathcal{F}(\theta)$ на нула се получава:

$$\nabla \mathcal{F}(\theta)|_{\theta=\hat{\theta}} = -2\Phi^T W y - 2\Phi^T W \Phi \hat{\theta} = 0.$$

Оптималните оценки, удовлетворяващи това равенство, са:

$$\hat{\theta} = (\Phi^T W \Phi)^{-1} \Phi^T W y. \quad (2.85)$$

Както се вижда, откъсно на равенството фигурират само известни величини – наличните данни и тегловната матрица, която също се задава по съответни причини (хетероскедастичен остатък, нестационарност на системата, липсващи стойности, редуциране на данните и т.н.).

Ако остатъкът е автокорелиран, вместо W в (2.85) е подходящо да се използва Σ_e^{-1} . Тъй като $\hat{\theta}$ зависи от данните, в които се съдържа неопределеност, то и оценките може да се разглеждат като случайни величини. Когато $W = \Sigma_e^{-1}$, освен че в случая $\hat{\theta}$ е неизместена оценка, дисперсията ѝ $\sigma_{\hat{\theta}}^2 = (\Phi^T \Sigma_e^{-1} \Phi)^{-1}$ е минимална. Оценител, който осигурява такива оценки, се нарича „най-добър, линеен, неизместен оценител“ (BLUE – Best Linear Unbiased Estimator). Свойствата на оценките са дискутирани подробно в точка 3.2.2.4.

Идеята за най-добра оценка (неизместена и с минимална дисперсия) на параметър има смисъл, когато моделът е линеен. При нелинейните модели има случаи, когато може да се намерят изместени оценки с по-малка дисперсия от оценките, които са неизместени.

2.3.4.5 Метод на най-малките квадрати

При WLS и неговия нелинеен вариант се приема, че или остатъкът е хетероскедастичен, или по някаква причина е необходима тегловна матрица с различни елементи по диагонала. Когато няма предварителна информация за неопределеностите, влияещи на изхода, разумно е да се приеме, че e_k е ергодичен (ковариационната матрица не се изменя във времето, т.е. $\Sigma_{e,k} = \Sigma_e = \text{const}$). При изпълнението и на останалите приемания, валидни за WLS (и систематизирани на Фигура 2.34),

матрицата Σ_e в (2.83) става:

$$\Sigma_e = \text{diag}(\sigma_e^2, \sigma_e^2, \dots, \sigma_e^2) \in \mathcal{R}^{(N-n)\ell \times (N-n)\ell},$$

а $\sigma_e^2 = [\sigma_{e,1}^2 \ \sigma_{e,2}^2 \ \dots \ \sigma_{e,\ell}^2]$ е векторът, съдържащ дисперсията на компонентите на ергодичния остатък e_k . Така, ако няма информация за нестационарности в системата, не е нужно данните да се отчитат с различни тегла. В тези случаи тегловната матрица в (2.84) може да се представи като $W = \text{diag}(w, w, \dots, w) \in \mathcal{R}^{(N-n)\ell \times (N-n)\ell}$, за $w \in \mathcal{R}^\ell$. Така за (2.82) е в сила:

$$e^T W e = \sum_{i=1}^{\ell} w_i \sum_{k=n+1}^N e_{i,k}^2 \propto e^T e,$$

т.е. целевата функция се свежда до

$$\mathcal{F}(\theta) = e^T e. \quad (2.86)$$

С други думи, задачата не се променя, ако $W = I$. Така се получава методът на най-малките квадрати (LS) като частен случай на WLS. От (2.86) и аналогично на извеждането в предишната точка за оптималните оценки се получава:

$$\hat{\theta} = (\Phi^T \Phi)^{-1} \Phi^T y. \quad (2.87)$$

Когато се прилага LS, трябва да се имат предвид приеманията, извършени до достигане до този частен случай. Както се вижда от (2.87), за да се определят параметрите по LS, не е нужна нито априорна информация за параметрите, нито за разпределението на остатъка. Единственото, което е необходимо, са наличните данни за входно-изходните величини. Затова този (опростен) метод за оценяване на параметри е най-разпространен в практиката. Той е разгледан подробно в точки 3.2.2 и 3.3.2.

Интерпретацията на решението по LS, описана по-долу, е нужна, тъй като при нея се въвеждат две матрици със специални свойства, за които ще стане дума по-нататък в изложението.

Нека за улеснение моделът е SISO и изходът му се представи като

$$\hat{y} = \Phi \hat{\theta} = \Phi (\Phi^T \Phi)^{-1} \Phi^T y = M y, \quad (2.88)$$

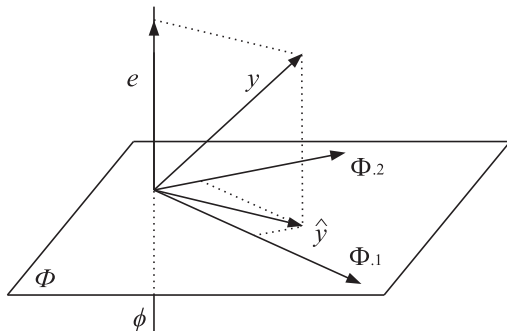
а остатъкът се развие по следния начин:

$$e = y - \hat{y} = (I - M)y = \Phi^\perp y. \quad (2.89)$$

В двете представяния са въведени матриците M и $\Phi^\perp$. Матрицата $M = \Phi (\Phi^T \Phi)^{-1} \Phi^T$ има свойството да проектира $N - n$ мерното прост-

ранство във факторното пространство, което е дефинирано от стълбовете на Φ . (Наистина лесно се проверява, че $M\Phi = \Phi$, т.е., както се очаква при тази проекция, стълбовете на Φ се проектират в себе си.) В (2.88) M проектира y в споменатото пространство, а получената проекция е именно $\hat{y}$. Това е интерпретацията на модела, записан във вида (2.88).

На Фигура 2.40 е представен случай, при който стълбовете на Φ са с размерност $N - n = 3$, а броят им (броят на факторите) е $z = 2$, т.е. $\Phi \in \mathcal{R}^{3 \times 2}$. Двата стълба $\Phi_{.1}$ и $\Phi_{.2}$ на матрицата на данните образуват база в $\mathcal{R}^2$ (това е равнината, в която те лежат – факторното пространство Φ). Тъй като $\hat{y}$ е проекция на y във Φ , то $\hat{y}$ лежи в тази равнина. Колкото по-информативна е комбинацията от факторите, толкова равнината ще е по-близка до y . В граничния случай, когато факторите съдържат цялата информация за описание на изхода (това означава, че няма остатъчна неопределеност), y ще лежи във факторната равнина. На фигурата $y \in \mathcal{R}^3$ е вектор, който не лежи в тази равнина, което показва, че той не може изцяло да бъде описан с двата фактора.



Фигура 2.40. Изходът на модела е проекция във факторното, а остатъкът – в ортогонално на факторното пространство. Пример за $\Phi \in \mathcal{R}^{3 \times 2}$.

Втората матрица $\Phi^\perp = I - \Phi(\Phi^T\Phi)^{-1}\Phi^T$ има свойството да проектира $N - n$ мерното пространство в ортогонално пространство на факторното. (Наистина лесно се проверява, че $\Phi^\perp\Phi = 0$, т.е. $\Phi^\perp$ проектира стълбовете на Φ в ортогонално за тях пространство и затова тяхната проекция е 0.) В (2.89) $\Phi^\perp$ проектира y в това ортогонално пространство, а получената проекция е остатъкът e . Матрицата $\Phi^\perp$ се използва в събспейс методите [148, 145, 125, 114] за оценяване на параметрите на динамични модели, описани в пространство на състоянието. На Фигура 2.40 ортогоналното допълнение на факторното (двумерно) пространство е правата ϕ (едномерно пространство). Векторът e лежи именно на

тази права. В такъв случай, колкото по-голямо е нивото на неопределеността в y , толкова посоката на този вектор е по-близка до правата ϕ . В граничния случай, когато y съдържа само неопределеност, а това означава, че факторите са напълно неинформативни, векторът y ще лежи на правата ϕ .

2.3.5 Еднократно и итеративно оценяване

В зависимост от вида на целевата функция може да се разграничат два вида решение на задачата за оценяване. В единия случай $\mathcal{F}(\theta)$ позволява аналитично оптималните параметри да се изразяват чрез данните. Известно е, че в оптимума на една функция (нека е скаларна с векторен аргумент) градиентът ѝ се нулира. Понякога от $\nabla \mathcal{F}(\theta) = 0$ е възможно оптималните оценки да се изразят чрез наличните наблюдения. В предните две точки са разгледани такива случаи. При тях целевата функция е квадратична, а моделът е линеен по отношение на параметрите и оценяването на θ се свежда до прилагане на GLS, WLS или LS.

В много случаи обаче не е възможно оптималните параметри да се определят аналитично. Ако ковариационната матрица на e_k или матрицата W не е известна или ако моделът е нелинейно параметризиран, се използват итеративни алгоритми, при които един вариант е многократно да се приложи LS или WLS. Това е т.нар. линеен подход за оценяване, на който е отделено внимание в Трета глава. Той се нарича линеен, защото се приема, че моделът е линеен по параметри. Дори зависимостта $\hat{y}_k(\hat{\theta})$ да не е линейна, има случаи, при които оценките сходят към търсените оптимални стойности.

Въпреки това невинаги е удачно или възможно да се приложи линейният подход за оценяване. Например нелинейността в модела по отношение на параметрите може да не позволява (дори фиктивно) отделянето и изразяването им като функция на данните. Тогава задачата за оценяване се решава с нелинейните варианти на споменатите методи или с по-общи методи като ML, като се прибегва до числената им реализация с методи за оптимизация. При тях отново итеративно параметрите се актуализират, докато не се изпълни съответно условие за край.

Понякога дали решението е аналитично или итеративно, зависи и от използвания метод за оптимизация. Например методът за оптимизация Нютон-Рафсън оценява параметрите с използване на градиента и Хесиана на целевата функция. Ако моделът е линеен, а целевата функция – квадратична, този метод достига до оптималното решение още в

първата итерация, т.е. реално оценяването не е итеративен процес. От тази гледна точка решението по LS в Трета глава може да се разглежда като реализация на метода Нютон-Рафсън, приложен към квадратичната целева функция. В случая целевата функция, нейният градиент и Хесиан (виж точка 3.3.2.2) за модел с вектор на параметрите са:

$$\mathcal{F}(\theta) = e^T e,$$

$$g = -2\Phi^T y + 2\Phi^T \Phi \theta,$$

$$H = 2\Phi^T \Phi.$$

Оценките по метода Нютон-Рафсън в i -тата итерация се актуализират като

$$\theta^{(i+1)} = \theta^{(i)} - H^{-1}g.$$

За първа итерация

$$\begin{aligned} \theta^{(1)} &= \theta^{(0)} - \frac{1}{2}(\Phi^T \Phi)^{-1}(-2\Phi^T y + 2\Phi^T \Phi \theta^{(0)}) \\ &= \theta^{(0)} + (\Phi^T \Phi)^{-1}\Phi^T y - \theta^{(0)} \\ &= (\Phi^T \Phi)^{-1}\Phi^T y. \end{aligned}$$

От друга страна, ако за оценяването на параметрите на същия модел се използва градиентен метод от първи ред [1, 34, 127] (оценките се обновяват с използване само на градиента), например:

$$\theta^{(i+1)} = \theta^{(i)} - s^{(i)}g,$$

като s е параметър за допълнителна промяна на стъпката, тогава в общия случай ще са необходими множество итерации за намиране на оценки, които да са задоволително близки до оптималните.

В настоящия труд не са представени методите за оценяване на параметрите на нелинейни (по параметри) модели.

2.3.6 Избор на структурата на модела

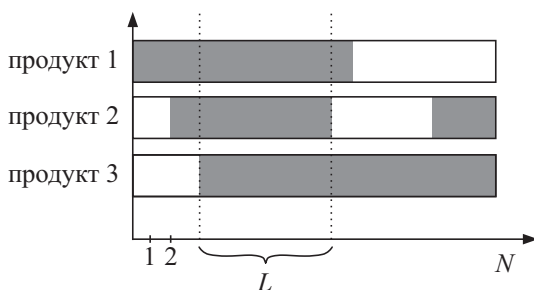
Тази дейност подробно е разгледана в Трета глава, а тук е маркирана само идеята и особеностите, свързани с уточняване на структурата на модела.

Изборът на подходяща структура на модела е важна стъпка в идентификацията. Тя се усложнява, когато системата е многомерна, особено при голям брой потенциални фактори, неинформативни, мултиколинearни сигнали и др. Основният стремеж при избора на структурата е

моделът да е достатъчно точен от гледна точка на целта, за която се разработва, и същевременно описанието да бъде лесно за използване (възможно по-просто). Понякога структурата на модела зависи и от други изисквания, наложени от приложната област. Например банките изискват модели на поведението на клиентите си, в които участват колкото е възможно повече характеристики (входове), които да не се купуват от кредитни бюра. Също така банките биха отхвърлили модел, ако някои фактори участват в него по нелогичен начин, от гледна точка на бизнеса. Това налага допълнителна промяна на структурата до получаването на модел, който логично може да бъде обяснен. Всъщност това е пример за използването на априорна информация, с която се налагат ограничения при избора на подходяща структура, които са наложени не от статистически, а от бизнес съображения.

В техниката също може да се наложат допълнителни изисквания, ако наблюдението на някои величини е свързано с повече средства, с нежелано влошаване на продукцията и т.н. Тогава също би следвало да се търси компромисен модел, който по възможност да не включва такива величини.

Освен изискванията моделът да е достоверен, опростен, икономически изгоден и невлизаш в противоречие с априорната информация, понякога върху структурата се налагат и ограничения, зависещи от наличните данни. Има случаи, когато стойностите на някои величини са достъпни в подинтервали на изборния



Фигура 2.41. Интервал на наблюдение и налични данни за три продукта. Интервалът, в който наборът е пълен, е означен с L .

интервал на наблюдение. Например входно-изходните сигнали, свързани с продуктите в пазарна система, са налични от момента на пускане на продуктите на пазара до тяхното снемане (Фигура 2.41). Така в набора от данни може да се окаже, че за отчитане на взаимовръзките между някои величини има малко данни и те може да отпаднат или максимално допустимият ред на полиноми да се ограничи. (Попълването на липсващи стойности е решение на проблема, представен на Фигура 2.41, но

то трябва да се прилага внимателно поради внасянето на допълнителна неопределеност с въведените ориентировъчни стойности.)

Структурата на модела се дефинира с помощта на т.нар. структурни параметри. За един многомерен модел структурни параметри са индексите на входно-изходните сигнали от набора, които са включени в модела, степените на полиномите в полиномните матрици (ако моделът е динамичен), както и чистите закъснения по различните канали. За статичен модел регресорите са и входове, но за динамичен модел даден вход или изход може да отговаря на няколко регресора в модела, ако участва в различни моменти от времето или (при общия вид с вектор на параметрите) ако влияе на повече от един изход.

За определяне на структурните параметри се използват различни техники, като обикновено се извършва оценяване на параметри на модели с различна структура. Такива са стъпковите алгоритми, разглеждани в следващата глава. При тях многократно се извършва оценяване на параметрите на модели с различни множества от регресори и с подходящи тестове се избира текущата най-удачна структура. Този процес продължава до изпълнението на определени условия за спиране.

На практика стъпковите алгоритми включват в себе си дейности, свързани със следващия (последен етап) от идентификацията. На всяка стъпка, след като се формира модел, на определено ниво се следи неговото качество, като тук акцентът е не толкова върху достоверността на модела като цяло, а по-скоро върху подобрението/влошаването на достоверността при промяна на структурата на модела. Въпреки че част от тестовите за значимост, провеждани при избора на структурата, се използват и на етапа на валидацията, тъй като стъпковите алгоритми служат за уточняване на структурните параметри, тези тестове се разглеждат като част от етапа на изграждането на модел, а не като елементи от същинската валидация.

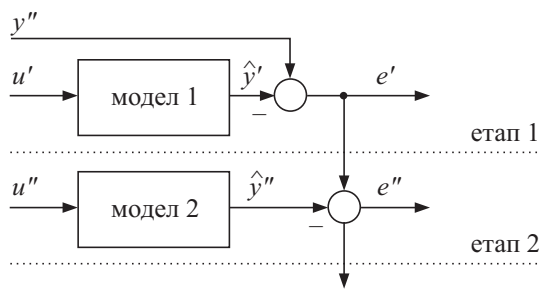
Съществуват много разновидности на стъпковите алгоритми. Ако алгоритъмът стартира с „празен“ модел (без фактори) и на всяка следваща стъпка се добавя фактор, стъпковата регресия е „права“. Ако алгоритъмът стартира с „пълния“ модел (отчитащ всички фактори) и на всяка стъпка се премахва фактор, стъпковата регресия е „обратна“. Обикновено двата варианта водят до различни крайни модели.

Правата и обратната стъпкови регресии са частни случаи на „комбинираната“ регресия. При нея последователно се извършват опити както

за добавяне, така и за премахване на фактори. Причината за проверка дали част от добавените (премахнати) в предишни стъпки фактори трябва да се премахнат (добавят отново), е евентуалната мултиколинеарност в данните – особеност, която е характерна за многомерните системи и е обсъдена подробно в точка 3.5.7.1.

2.3.7 Поетапно моделиране

Понякога изборът на структурата на модела се извършва на няколко етапа. Целта обикновено е да се получи модел със структура, удовлетворяваща изискванията, наложени от конкретната приложна област, които са в разрез със статистическата значимост на входните величини. При поетапното моделиране последователно се описват отделни аспекти от системата. В първия етап се намира максимално точен модел, зависещ от тези входове, които са най-предпочитани от гледна точка на използването на модела. На всеки следващ етап уточненият до момента модел се разширява с нови въздействия, които са все по-малко предпочитани. Така се дава възможност на част от входовете (дори и по-малко значими) да участват в модела, а други (дори и значими) по възможност да не участват.



Фигура 2.42. Поетапно моделиране

За да се реши тази задача, на първия етап изходите на системата се сравняват с тези на модела (Фигура 2.42 – модел 1), а за входове се избират най-предпочитаните въздействия (те са означени с u'). На всеки следващ етап обаче стремежът е изходите на модела да се доближат до остатъците от предишния етап. Причината за това е, че остатъците от даден етап съдържат тази част от поведението на системата, която още не е, но предстои тепърва да бъде описвана.

Ако моделът е нелинеен, поэтапното моделиране се извършва, като задачата за оптимизация се дефинира само по отношение на новите параметри от текущия етап, а останалите вече намерени параметри участват в алгоритъма като известни константи.

Пример. *Поетапно моделиране в банковия сектор*

В точка 2.3.6 е споменато, че банките предпочитат модели, които използват, доколкото е възможно, характеристиките на кредитоискателите, които не струват пари за финансовите институции. Те се получават, когато потенциален клиент попълва формуляр при кандидатстването за кредит. От друга страна, данните за тези величини често са недостатъчни за постигане на желана степен на достоверност на модела. Затова, в допълнение на безплатните характеристики на кандидатите, се използват и данни, които се предоставят от кредитни бюра.

Обикновено бюро характеристиките са значително по-информативни, тъй като включват в себе си макроикономически показатели, както и информация за кандидатите, която не е достъпна за банката. С помощта на данните от кредитни бюра банката може чувствително да подобри оценката на кредитния риск, който би поела, ако одобри съответен кредит, а това води до взимането на по-прецизни решения при отпускането на кредити.

Напълно възможно е част от входните данни за моделите, формиращи дадена бюро характеристика, да са силно корелирани с достъпните за банката данни от кандидатстването. Това означава, че при един равнопоставен избор на фактори бюро характеристиката (агрегирала в себе си информацията от различни фактори) със сигурност би участвала в модела, което е за сметка на по-малко значимите безплатни характеристики. Така банката би получила модел, използването на който е свързано с ненужни допълнителни плащания.

За да се осигури по-голям шанс на безплатните и предпочитани от банката характеристики да участват в модела, се използва споменатото поэтапно моделиране. На първия етап се използват изцяло данните, събрани от кандидатстването. На втория етап се въвеждат по-малко предпочитани характеристики и т.н. Така в крайния модел ще участват само тези платени характеристики, които значително подобряват модела и същевременно информацията в тях не се припокрива с тази в безплатните характеристики. ◇

2.4 Валидация на модела

След получаване на модел се пристъпва към последния етап на идентификацията, който е валидация на модела. При него се определя доколко е близко поведението на получения до момента модел до това на системата. Също така на този етап се взима решение дали моделирането да се прекрати. Ако текущият резултат не е удовлетворителен, се определя с коя дейност да се продължи. В чисто експерименталния подход за моделиране оценката на достоверността се постига с помощта на наличните данни. Съществуват много статистически показатели (зависещи от данните), с които се оценява качеството на различни аспекти на модела. Причината за това многообразие е стремежът да се осигури достатъчно информация за това, какво в модела (и дали) трябва да се подобри и съответно да се вземе решение откъде в цялата схема на идентификацията да се продължи.

Когато изследователят има очаквания, базирани на предварителна информация за обекта, тя също трябва да се вземе предвид. Тази информация е много полезна и дори е наложително да се използва особено когато данните не са подходящи за моделиране.

В темата по-долу са представени различни показатели на качеството, които се основават на експерименталните данни. Понеже показателите на качеството са статистически величини, често те се наричат накратко статистики. Тези величини може да се разделят на:

- обобщени показатели;
- частични показатели.

Обобщените показатели характеризират качеството на модела като цяло. Такива статистики са представени в следващата точка. Също така там са описани и тестове, базирани на корелационните свойства на остатъка. Показателите от втората група (разгледани в точка 2.4.4) характеризират точността на получените оценки поотделно, както и значимостта на всеки фактор, участващ в модела. Към тази група спадат и статистики, които оценяват степента на евентуалната връзка между факторите.

В следващите две подточки се описват някои важни моменти, свързани с валидацията на линейно параметричните модели.

2.4.1 Разпределение на остатъка

Отново, както при предишния етап, така и при изложението на етапа на валидация, ще се използва общото описание на модел с вектор на параметрите.

Методът LS се прилага, когато моделът е линеен по отношение на параметрите. Но допълнителното приемане, че остатъкът има нормално разпределение, не е нужно за валидността на метода. Въпреки това често се приема, че при изграждането на линеен модел (и съответно прилагането на LS) остатъкът е нормален, т.е.

$$e \sim \mathcal{N}(0, \Sigma_e). \quad (2.90)$$

Приемането за нормалност не е изискване, което се налага от теоретична гледна точка [101]. От друга страна, често предположението, че e_k е Гаусов, е напълно удовлетворително в практиката, тъй като обикновено обобщената неопределеност зависи от много и различни процеси. В тази насока централната гранична теорема [16] гласи, че разпределението на сумата от случайни величини с произволни разпределения кло-ни към нормално разпределение, когато броят на величините нараства. (Има някои специфични разпределения, които не изпълняват условието на Линдберг [7] и сумата не сходяща към Гаусово разпределение, но те нямат отношение към разглежданията в монографията.)

При някои разпределения на e_k , едно от които е Гаусовото, LS максимизира функцията на правдоподобие и оценките са BLUE, т.е., освен че са неизместени, имат минимална дисперсия (виж точка 3.2.2.4).

Друга особеност на случая с Гаусов остатък е, че показателите на качеството, проверката на хипотези, изследването на корелационните свойства на остатъка и др. дейности, свързани с анализа на модела, са ясно дефинирани и лесни за тълкуване. Затова, ако моделът е линеен, по-долу ще се приема, че e_k е с нормално разпределение.

2.4.2 Модел със свободен член

Тъй като разгледаните статистики и тестове се прилагат по един и същи начин за всеки изход на модела, за опростяване на изложението се приема, че моделите са MISO.

Ако данните са стандартизирани (виж предварителната обработка на данните, точка 2.2.2.4), в резултат на това се получава т.нар. стандартизиран модел. След оценяването на параметрите моделът се дестандартизира, като при това, освен че параметрите се изменят, в MISO

модела се добавя параметърът θ_0 на свободния член, с който се отчита първоначалната изместеност на сигналите. Това означава, че спрямо стандартизирания модел в дестандартизирания има една степен на свобода вповече, с която се отчита описанието на споменатата изместеност (показана на Фигура 2.18). Линейният MISO модел със свободен член е

$$y_k = \theta_0 + \theta_1 \varphi_{1,k} + \dots + \theta_{p-1} \varphi_{p-1,k} + e_k = \varphi_k^T \theta + e_k, \quad (2.91)$$

като векторите $\theta \in \mathcal{R}^p$ и $\varphi \in \mathcal{R}^p$ (за MISO модели $z = p$) са:

$$\theta = [\theta_0 \ \theta_1 \ \dots \ \theta_{p-1}]^T \text{ и } \varphi = [1 \ \varphi_{1,k} \ \dots \ \varphi_{p-1,k}]^T.$$

В монографията се използват и стандартизираното представяне (най-вече, когато става дума за оценяване на параметри), и дестандартизираното (обикновено в етапа на валидацията, тъй като има показатели, които приемат различни стойности в зависимост от това, дали моделът е стандартизиран или не). Понеже (2.91) описва взаимовръзките между първоначалните величини, то статистиките за оценка на достоверността на модела трябва да се изчисляват за това представяне. Изключение прави анализът на значимостта на факторите (точка 2.2.3.3), при който може да е по-подходящо да се използва стандартизираното описание на системата.

2.4.3 Обобщени показатели

В следващите две подточки са разгледани най-често срещаните обобщени показатели на качеството за линейни модели.

Общото между начините за валидация, описани в тази точка, е, че те оценяват като цяло близостта на модела до наличните данни за системата и са подходящи за вземане на решение дали моделът е достатъчно достоверен, или е необходимо допълнително да се подобри. Въпреки това те не са достатъчни за прекратяването на цикъла на идентификация. Възможно е например качеството на модела като цяло да е удовлетворително, тъй като в описанието участва подходяща група от фактори, но също така друга част от факторите може да не допринася значимо за описанието на изхода (изолирането на тези фактори става с помощта на статистики като тези, разгледани в точка 2.4.4). Тогава идентификацията е желателно да продължи с допълнително опростяване на структурата на модела.

2.4.3.1 Абсолютни показатели

Целева функция

Най-често абсолютните показатели са свързани с целевата функция, която в долните разглеждания е

$$\mathcal{F}(\theta) = e^T e$$

(както и претегленият ѝ вариант). Тя зависи от броя N на наблюденията и тъй като в случая е сума от неотрицателни стойности, то с нарастване на N и стойността ѝ неограничено нараства, т.е.

$$\lim_{N \rightarrow \infty} \mathcal{F}(\theta) = \infty.$$

Поради зависимостта на $\mathcal{F}(\theta)$ от N не е възможно да се сравнят модели, получени с набори от данни с различна дължина. Това нежелано свойство на $\mathcal{F}(\theta)$ може да се избегне, ако вместо сумата от квадрата на грешката се използва дисперсията на остатъка σ_e^2 .

Дисперсия на остатъка

Дисперсията на e_k характеризира определено статистическо свойство на остатъка и не зависи от броя на наблюденията в набора. Тъй като се работи с краен брой данни, вместо σ_e^2 се определя нейна оценка. При ергодичен остатък, колкото повече са наблюденията, толкова по-близка е тази оценка до действителната стойност на дисперсията, т.е.

$$\lim_{N \rightarrow \infty} \hat{\sigma}_e^2 = \sigma_e^2.$$

Тази статистика (както и останалите показатели) се изчислява с помощта на наличните данни. В записите по-долу изчислената оценка на дисперсията ще се бележи със σ_e^2 (вместо със $\hat{\sigma}_e^2$), освен ако не е нужно да се наблегне на факта, че се работи с краен брой наблюдения. И така дисперсията на остатъка (изчислена с наличните данни) е

$$\sigma_e^2 = \frac{1}{v} \mathcal{F}(\theta) = \frac{1}{N - n - p} \sum_{k=n+1}^N e_k^2, \quad (2.92)$$

където v са степените на свобода на остатъка за интервала на наблюдение, p е броят на параметрите, а n е максималното закъснение на сигнал (брой тактове) при формирането на изхода. Тъй като моделът е MISO, то p е броят на параметрите в описанието на изхода. Ако системата е ММО, то за качеството по i -тия изход $p \leftarrow p_i$.

Колкото по-малка е стойността на σ_e^2 , толкова по-голяма е близостта на модела до наличните данни. Въпреки че дисперсията на остатъка не нараства с увеличаване на броя на наблюденията, тя е абсолютен показател, чийто недостатък (описан в долния пример) е зависимостта му от дисперсията на изхода.

Пример. *Зависимост на σ_e^2 от дисперсията на изхода*

Тук е по-удобно системата да е с два изхода. Нека първият изход да се изменя в широки граници и дисперсията му е $\sigma_{y,1}^2 \approx 10^6$, а $y_{2,k}$ се изменя в тесен диапазон и е с дисперсия $\sigma_{y,2}^2 \approx 10^{-8}$. Това означава, че по-големи стойности на $\sigma_{e,1}^2$, спрямо $\sigma_{e,2}^2$, биха били приемливи (примерно $\sigma_{e,1}^2 < 25 \times 10^2$ отговаря на ниво на остатъчната неопределеност максимум 5% от това на изхода. От друга страна, приемливите стойности за $\sigma_{e,2}^2$ са значително по-малки – за да се гарантира ниво на остатъчната неопределеност максимум 5% от стойностите на изхода, е необходимо $\sigma_{e,2}^2 < 25 \times 10^{-12}$. $\diamond$

Следователно, особено когато системата е с повече изходи, е необходимо показателят на качеството да е унифициран така, че да не зависи от дисперсията на конкретния изход. Освен това в задачата на идентификацията често се налага да се сравняват модели, получени от различни експерименти и/или с различна структура. В тези случаи е подходящо да се използват относителни показатели като тези, разглеждани в следващата точка.

2.4.3.2 Относителни показатели

Показателите от тази група се формират така, че да не зависят нито от броя на наблюденията, нито от диапазоните на изменение на изходите. Също така те позволяват да се сравнява достоверността на модели с различна структура. Понеже стойностите им са относителни, често се формират в проценти. Може би най-разпространеният представител на относителните показатели е коефициентът на определеност (наричан още коефициент на детерминация), както и някои от многобройните му модификации. Затова акцентът в изложението е върху този показател и някои от най-често срещаните му варианти.

Коефициент на определеност

Коефициентът на определеност е величина, която показва каква част от вариацията на изхода се описва от регресионния модел. Често в литературата той се бележи с R^2 или с ρ^2 , тъй като е равен на квадрата

на корелационния коефициент на Пиърсън между изхода на системата и на модела. Въпреки че общоприетият символ на коефициента на Пиърсън е ρ , по-долу вместо ρ^2 за коефициента на определеност ще се използва най-разпространеният му запис, а именно R^2 .

В текста се използват следните означения:

- SST (Total Sum of Squares) е сумата от квадратите на центрирания изход на системата, т.е.

$$\text{SST} = \sum_{k=n+1}^N (y_k - \bar{y})^2$$

и отразява вариацията на изхода на системата;

- SSR (Regression Sum of Squares) е

$$\text{SSR} = \sum_{k=n+1}^N (\hat{y}_k - \bar{y})^2$$

и отразява вариацията (степената на изменение) на изхода на модела (понякога тази сума се бележи с SSE от Explained Sum of Squares);

- SSE (Error Sum of Squares) се формира като

$$\text{SSE} = \sum_{k=n+1}^N (y_k - \hat{y}_k)^2$$

и отразява вариация, която не е описана от модела, но се съдържа в изхода на системата (понякога тази сума се бележи с SSR от Residual Sum of Squares).

С използване на тези статистики коефициентът на определеност се дефинира като:

$$R^2 = 1 - \frac{\text{SSE}}{\text{SST}}. \quad (2.93)$$

В някои източници, вместо (2.93), R^2 се записва като:

$$R^2 = \frac{\text{SSR}}{\text{SST}}, \quad (2.94)$$

като при прехода от (2.93) към (2.94), се използва зависимостта:

$$\text{SSE} = \text{SST} - \text{SSR}. \quad (2.95)$$

В SAS и други софтуери за изчисляване на R^2 за линейни модели се използва (2.94). От друга страна, формула (2.93) се използва за оценка на достоверността както на линейни, така и на нелинейни модели.

Често (2.93) и (2.94) водят до различни (макар и близки) резултати. Причина за това е нарушаването на условията (две на брой), при които (2.95) е в сила. Едното условие е средната стойност на $\hat{y}_k$ да съвпада с $\bar{y}$, или с други думи в модела да има свободен член, който да гарантира това. За линейни модели такъв свободен член може без проблем да се въведе, но невинаги той може да се включи в структурата на нелинейните модели.

Другото условие може да се определи, ако SST се развие по следния начин:

$$\text{SST} = \sum_{k=n+1}^N (y_k - \bar{y})^2 = \sum_{k=n+1}^N (y_k - \hat{y}_k + \hat{y}_k - \bar{y})^2 = \sum_{k=n+1}^N (e_k + \hat{y}_k - \bar{y})^2,$$

или след разкриване на скобите:

$$\begin{aligned} \text{SST} &= \sum_{k=n+1}^N e_k^2 + \sum_{k=n+1}^N (\hat{y}_k - \bar{y})^2 + 2 \sum_{k=n+1}^N e_k(\hat{y}_k - \bar{y})^2 \\ &= \text{SSE} + \text{SSR} + 2 \sum_{k=n+1}^N e_k(\hat{y}_k - \bar{y})^2. \end{aligned}$$

Ако остатъкът e_k (неотчетеното от модела поведение на системата) и изходът на модела $\hat{y}_k$ (отчетените аспекти от нейното поведение) не са корелирани, то последната сума е пренебрежимо по-малка от първите две и с нарастване на N в граничния случай се достига до (2.95). Тази липсата на корелация между e_k и $\hat{y}_k$ е търсено свойство, което е в основата на корелационните тестове на остатъка, разгледани в 2.4.3.3.

Връзката между коефициента на определеност и корелационния коефициент на Пиърсън може да се определи по сходен начин като в разглежданията, свързани с (2.95). Нека числителят на

$$\rho_{\hat{y}y} = \frac{\sum_{k=n+1}^N (\hat{y}_k - \bar{y})(y_k - \bar{y})}{\sqrt{\sum_{k=n+1}^N (\hat{y}_k - \bar{y}) \sum_{k=n+1}^N (y_k - \bar{y})}}$$

се развие по следния начин:

$$\sum_{k=n+1}^N (\hat{y}_k - \bar{y})(y_k - \hat{y}_k + \hat{y}_k - \bar{y}) = \sum_{k=n+1}^N (\hat{y}_k - \bar{y})e_k + \sum_{k=n+1}^N (\hat{y}_k - \bar{y})^2.$$

Отново, ако e_k и $\hat{y}_k$ не са корелирани, то първата сума е пренебрежимо

по-малка от втората. Така, когато $N \rightarrow \infty$, за $\rho_{\hat{y}y}^2$ може да се запише:

$$\begin{aligned}\rho_{\hat{y}y}^2 &= \frac{(\sum_{k=n+1}^N (\hat{y}_k - \bar{y})^2)^2}{\sum_{k=n+1}^N (\hat{y}_k - \bar{y})^2 \sum_{k=n+1}^N (y_k - \bar{y})^2} \\ &= \frac{\sum_{k=n+1}^N (\hat{y}_k - \bar{y})^2}{\sum_{k=n+1}^N (y_k - \bar{y})^2} = \frac{\text{SSR}}{\text{SST}} = R^2.\end{aligned}$$

R^2 се изменя между 0 и 1, като, колкото е по-голям, толкова по-достоверен е моделът. Тъй като зависи единствено от данните, а не от модела, от (2.93) се вижда, че $R^2 = 1$ се получава, когато $\text{SSE} = 0$, т.е. остатъкът е нула за интервала на наблюдение, което е равносилно на пълно съвпадение на изхода на модела с този на системата ($\text{SSR} = \text{SST}$). Този случай не е реалистичен поради наличието на неопределености, които не се влияят от качеството на модела. Затова, ако все пак показателят има стойност 1, изследователят трябва да провери дали моделът не е преоразмерен (точка 2.4.5). В другия граничен случай, когато $R^2 = 0$, от (2.93) се вижда, че $\text{SSE} = \text{SST}$. Това означава, че всички фактори в модела са напълно незначими и съответните параметри на модела (всички с изключение на свободния член θ_0 , ако такъв е въведен) са равни на 0. В междинния случай, например ако $R^2 = 0.8$, то може да се каже, че 80% от измененията (вариацията) на изхода са отчетени от регресионния модел. Казано по друг начин, 80% от изменението на y_k се обяснява с помощта на факторите в модела.

Коригиран R^2 и показател на отчетената вариация

По-прецизна характеристика на достоверността на модела от R^2 е т.нар. коригиран R^2 (означен по-долу с R_a^2 от adjusted R^2). Тъй като за валидация се използва крайна извадка, стойността на R^2 съдържа неопределеност. В този смисъл разликата между R^2 и R_a^2 е, че последният се формира така, че да нараства само ако при добавянето на нови фактори подобрението в модела нараства повече, отколкото се очаква от влиянието на неопределеността. Поради по-консервативния характер на R_a^2 той може да приема отрицателни стойности, като винаги $R_a^2 \leq R^2$.

Показателят R_a^2 се изчислява по следния начин:

$$R_a^2 = 1 - \frac{\sigma_e^2}{\sigma_y^2} = 1 - \frac{\text{SSE}/v_e}{\text{SST}/v_y} = 1 - \frac{v_y}{v_e}(1 - R^2),$$

като степените на свобода са $v_e = N - n - p$ и $v_y = N - n - 1$.

С нарастване на броя на наблюденията N се очаква влиянието на неопределеността върху вариациите в стойността на R^2 да намалява. В граничния случай, ако e_k е ергодичен, R^2 не зависи от неопределеността. Освен това за отношението на степените на свобода в последния израз, е в сила:

$$\lim_{N \rightarrow \infty} \frac{v_y}{v_e} = 1.$$

Тогава за R_a^2 се получава (както се и очаква) $R_a^2 = R^2$. От друга страна, ако N намалява, то се очаква влиянието на неопределеността върху стойността на R_a^2 да расте. Отношението $\frac{v_y}{v_e}$ също нараства, което означава, че $R_a^2 < R^2$ и колкото по-чувствителна е стойността R^2 към случайните вариации в данните, толкова по-консервативен става R_a^2 .

Понякога вместо R_a^2 се използват процентните стойности на показателя, като в някои източници [146] тази статистика се нарича показател на отчетената вариация (VAF – Variance Accounted For), т.е.

$$\text{VAF} = \max \{0, R_a^2\} \times 100\%.$$

Ако $\sigma_e^2 \geq \sigma_y^2$, то $R_a^2 < 0$ и достоверността на модела е толкова малка, че се приема за 0%.

В [82] е предложен вариант на VAF, при който се отчита степента на неопределеност в текущия вектор или матрица на регресорите, когато предварително са попълвани липсващи стойности или са обработвани нехарактерни наблюдения.

2.4.3.3 Корелационни тестове на остатъка

Разгледаните до момента показатели на качеството са приложими както за динамични, така и за статични системи. За разлика от тях описаните по-долу тестове [11, 148, 71, 93, 129] изследват корелационната зависимост между стойностите на сигнали в различни наблюдения и имат смисъл, ако моделът е динамичен.

Автокорелационен тест на остатъка

При този тест първо се изчислява автокорелационната функция на остатъка

$$R_{e,i} = \text{cor}(e_k, e_{k-i})$$

за $i = \overline{0, s}$, като практическо правило [11] за максималното отместване е $s \in [25, N - p]$. Според централната гранична теорема разпределение-то на сумата от случайни величини с произволни разпределения клони

към нормално разпределение. Тогава, ако остатъкът е бял Гаусов шум с нулева средна стойност и независещ от входа, то нормираният вектор

$$r_e = \frac{\sqrt{N}}{R_{e,0}} [R_{e,1} \ R_{e,2} \ \dots \ R_{e,s}]^T$$

има асимптотично Гаусово разпределение с нулева средна стойност и ко-вариационна матрица $\Sigma_{r_e} = I_s$ [146]. Проверката дали r_e има нормално разпределение с вероятност 0.95, е еквивалентна на

$$|r_{e,i}| \leq 1.96 \text{ за } i = \overline{1, s}.$$

Липсата на значима корелация между стойностите на e_k в различни моменти от времето означава, че цялата зависимост на изхода от предисторията на входно-изходните сигнали е отчетена от модела. По този начин между стойностите на остатъка в различните моменти от времето не остават значими зависимости (те се съдържат в $\hat{y}_k$).

Взаимнокорелационен тест между остатъка и входа

Може да се извърши тест, подобен на предишния, но с използване на входните сигнали и остатъка. Тук се формират m взаимнокорелационни функции между e_k и $u_{j,k}$ (за ММО модели $e_k \in \mathcal{R}^\ell$ и ако всички m входове влияят на всички изходи, се формират $m\ell$ взаимнокорелационни функции). За опростяване индексът j на текущия вход се пропуска. Текущата взаимнокорелационна функция е $R_{eu,i} = \text{cor}(e_k, u_{k-i})$ за $i = \overline{1, s}$ (отново $s \in [25, N - p]$) и се формира векторът

$$r_{eu} = \frac{\sqrt{N}}{\sqrt{R_{e,0}R_{u,0}}} [R_{eu,1} \ R_{eu,2} \ \dots \ R_{eu,s}]^T.$$

Отново тестът се свежда до проверка дали елементите на вектора r_{eu} попадат в определени граници. Ако e_k е бял Гаусов шум с нулева средна стойност и независещ от входа, то нормираният вектор r_{eu} е с нулева средна стойност, и ако

$$|r_{eu,i}| \leq 1.96 \text{ за } i = \overline{1, s},$$

то с вероятност 0.95 r_{eu} има нормално разпределение.

Липсата на значима корелация между e_k и u_k означава, че влиянието на входовете на системата върху изхода (изходите) се описва коректно от модела, така че между стойностите на остатъка и на входния сигнал в различните моменти от времето не остават значими зависимости (те се съдържат в $\hat{y}_k$).

2.4.4 Частични показатели на качеството

2.4.4.1 Стандартна грешка

Оценка на параметър и доверителен интервал

Тъй като оценката $\hat{\theta}$ зависи от данните, а в тях има неопределеност, то елементите на вектора $\hat{\theta}$ са случайни величини. В такъв случай векторът θ^* на оптималните (неизвестни) параметри се намира в околност на изчислената оценка $\hat{\theta}$, като, колкото по-точна е тя, толкова околността е по-малка. Обикновено се приема, че елементите на $\hat{\theta}$ са нормално разпределени и ако не са изместени, то $\hat{\theta} \sim \mathcal{N}(\theta^*, \Sigma_{\hat{\theta}})$. В точка 3.2.2.4 е засегнат въпросът за изместеността и е изведена ковариационната матрица на оценките, която е

$$\Sigma_{\hat{\theta}} = \sigma_e^2 (\Phi^T \Phi)^{-1}.$$

При валидацията на модела от интерес е коренът от диагоналните елементи на $\Sigma_{\hat{\theta}}$

$$\sigma_{\hat{\theta}} = (\text{diag } \Sigma_{\hat{\theta}})^{\frac{1}{2}} = \sigma_e (\text{diag}(\Phi^T \Phi)^{-1})^{\frac{1}{2}}. \quad (2.96)$$

Тъй като $\hat{\theta}$ е вектор, то и $\sigma_{\hat{\theta}}$ е вектор с размерност p (в точка 3.2.2.4 на следващата глава е разгледан случаят, когато параметрите са подредени в матрицата Θ и тогава $\sigma_{\hat{\theta}}^2 \in \mathcal{R}^{z \times \ell}$). Елементите на $\sigma_{\hat{\theta}}$ са стандартните грешки на оценките на параметрите. Със $\sigma_{\hat{\theta}}$ може да се дефинират доверителните интервали, в които се намират оптималните търсени стойности на параметрите. Например при нормално разпределени оценки за оптималните параметри е в сила

$$\theta_i^* \in [\hat{\theta}_i - 1.96\sigma_{\hat{\theta},i}, \hat{\theta}_i + 1.96\sigma_{\hat{\theta},i}] \quad (2.97)$$

с вероятност 0.95. Ако елементите на $\sigma_{\hat{\theta}}$ са твърде големи спрямо съответните оценки, то тези оценки са неточни и идентификацията трябва да продължи. В икономическите системи например (където нивото на неопределеност е голямо) е прието да се проверява дали интервалът (2.97) включва нулата. Ако за даден параметър е изпълнено:

$$[\hat{\theta}_i - 1.96\sigma_{\hat{\theta},i} < 0 < \hat{\theta}_i + 1.96\sigma_{\hat{\theta},i}],$$

то съответният фактор се премахва от модела. Включването на нулата в доверителния интервал означава, че оценката може да е с обратен знак, т.е. толкова е неточна, че дори не е сигурно дали отразява коректно общата тенденция (посоката на тренда) между съответния фактор и изхода.

Едно предимство на стандартизирането на факторите по време на предварителната обработка на данните е, че анализът на модела, поотделно за всеки фактор, е значително по-лесен. Причината за това е следната. Тъй като стандартизираните фактори са центрирани и са с единични стандартни отклонения, елементите по диагонала на матрицата $(\Phi^T \Phi)^{-1}$, която участва в (2.96), са равни. i -тият диагонален елемент на матрицата на Фишер е $[\Phi^T \Phi]_{ii} = N - 1$, а при отсъствие на мултико-линеарност диагоналните елементи на обратната на $\Phi^T \Phi$ са:

$$[\text{diag}(\Phi^T \Phi)^{-1}]_i = \frac{1}{N - 1}.$$

На практика това е идеализиран случай – при крайна извадка между тези диагонални елементи винаги има разлика, която зависи от конкретните данни. От друга страна, дисперсията на остатъка σ_e по даден изход (в изложението моделите са MISO) е скалар. Тогава и елементите на $\sigma_{\hat{\theta}}$, определени от (2.96), също са сходни, т.е.

$$\sigma_{\hat{\theta}} = \sigma_e (\text{diag}(\Phi^T \Phi)^{-1})^{\frac{1}{2}} \approx \frac{\sigma_e}{\sqrt{N - 1}}.$$

Въпреки че в болшинството от случаите факторите не са напълно независими (особено при отчитането на голям брой фактори), анализът на модела е значително по-лесен, когато се основава на стойностите на стандартизираните оценки и на техните стандартни грешки. В този случай, колкото по-голяма е стойността на дадена оценка, толкова по-значим е съответният фактор. Също така вариацията между елементите на $\sigma_{\hat{\theta}}$ показва доколко факторите са независими. На практика колкото по-силна е връзката между факторите, толкова по-големи са стандартните грешки. Така изследователят може лесно да се ориентира за значимостта на факторите и дали има чувствителна взаимовръзка между тях, а това е полезно знание, когато се взима решение за следващите стъпки в моделирането. Освен това унифицирането на факторите улеснява и автоматизирането на тази част от цикъла на идентификация.

2.4.4.2 Степен на линейна зависимост между факторите

Когато в модела участват много фактори (и особено ако той е с много входове), е удачно да се следи доколко тези величини са зависими помежду си. Идеята е, че всеки фактор трябва да носи уникална информация за изхода. Ако информацията в даден фактор се припокрива с тази в останалите фактори, то той е ненужен и трябва да се премахне от модела. За тази цел се използва статистиката VIF (Variance Inflation

Factor), която се изчислява за всеки фактор поотделно. В началото се формират z модела (колкото са факторите и свободният член θ_0 , ако той фигурира в крайния модел). Изходът на i -тия модел е i -тият фактор, а регресори са всички останали фактори, т.е. текущият модел е

$$\varphi_{i,k} = \varphi_{-i,k}^T \theta_i + e'_{i,k}, \quad (2.98)$$

като $\varphi_{-i,k} \in \mathcal{R}^{z-1}$ е редуцираният вектор на регресорите, съдържащ всички фактори, без i -тия, $\theta_i \in \mathcal{R}^{z-1}$ е съответният вектор на параметрите, а $e'_{i,k}$ е остатъкът, който всъщност е тази част от поведението на i -тия фактор, която е уникална, тъй като не е описана от останалите фактори. Колкото по-качествен е моделът, толкова по-ненужен е i -тият фактор, и обратното – ако след изчисляване на $\hat{\theta}_i$ се получи недостоверен модел, това означава, че информацията, която носи i -тият фактор, не се съдържа в останалите фактори и $\varphi_{i,k}$ е удачно да участва в модела. За да се изчисли VIF, първо се определя коефициентът на определеност, който за i -тия модел е

$$R_i^2 = \frac{SSR_i}{SST_i},$$

където

$$SST_i = \sum_{k=n+1}^N (\varphi_{i,k} - \bar{\varphi}_i)^2,$$

$$SSR_i = \sum_{k=n+1}^N (\varphi_{-i,k}^T \hat{\theta}_i - \bar{\varphi}_i)^2,$$

след което се определя стойността на VIF:

$$VIF_i = \frac{1}{1 - R_i^2}.$$

Както се вижда, ако $R_i^2 = 0$, което е идеалният вариант ($\varphi_{i,k}$ е напълно независим от останалите фактори), то $VIF_i = 1$. Колкото по-зависим е факторът, толкова по-голяма е стойността на R_i^2 и съответно на VIF_i . В граничния случай, когато $R_i^2 = 1$, $VIF_i = \infty$.

Квадратният корен на VIF носи информация за това, колко пъти по-голяма е стойността на стандартната грешка от тази, която тя би имала, ако факторът не беше зависим от останалите фактори. Например, ако $VIF_i = 3.24$, то $\sqrt{3.24} = 1.8$ и съответно стандартната грешка на оценката е почти 2 пъти по-голяма, отколкото когато съответният фактор няма обща информация с останалите фактори.

Има различни прагови стойности за VIF, над които се приема, че даден фактор в модела е ненужен, които зависят най-вече от приложната област. Например $VIF = 3$ е често срещан праг във финансите. Също се среща $VIF = 5$, а в някои случаи е прието прагът да е $VIF = 10$.

Реципрочната стойност на VIF също се използва за оценка на степента на зависимост между факторите и се предлага от софтуери, като SAS и SPSS. Тази статистиката е наречена толеранс и е

$$\text{Tolerance}_i = \frac{1}{VIF_i} = 1 - R_i^2.$$

Изменя се между 0 и 1 и често е предпочитана вместо VIF поради по-директната връзка с R_i^2 .

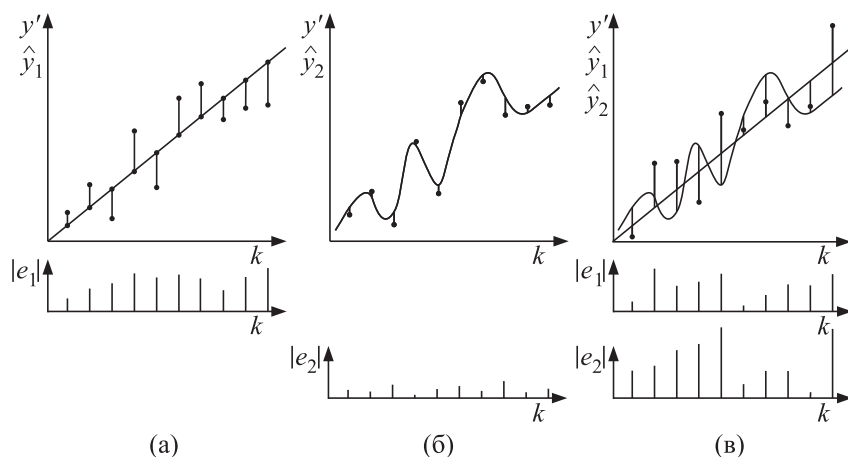
2.4.5 Преоразмеряване и универсалност на модела

При уточняването на структурата на модела съществува следната тенденция: колкото повече параметри се добавят, толкова изходът на модела се доближава до данните, които се използват за оценяване на параметрите. Причина за това е, че с всеки добавен параметър нараства възможността на описанието да отрази по-прецизно данните. В граничния случай, когато броят на параметрите е равен на броя на наблюденията ($\dim \theta = N$), ще се постигне пълно съвпадение на изхода на модела с изходните данни (тогава $\mathcal{F}(\theta) = 0$). Но на практика пълно съвпадение на поведението на модела с това на системата не е постигнато. Ако $\hat{y}_k = y_k$ за целия интервал на наблюдение, това означава, че моделът описва както полезната информация в данните, така и конкретния ефект на шума от измерването, влиянието на околната среда и т.н., асоциирани с използвания набор от данни.

На Фигура 2.43 са представени стойности на зашумен нарастващ линейно във времето сигнал. Удачен модел, който го описва, също би имал изход, изменящ се линейно (въпреки че в случая $\mathcal{F}(\theta) > 0$). Такъв модел би отразявал детерминираната съставка, без да се стреми да опише моментните стойности на смущението. От друга страна, ако моделът е преоразмерен, то изходът му ще зависи от конкретната реализация на обобщената неопределеност (случайна величина) в набора от данни.

На Фигура 2.43 е представен и такъв случай, като моделите са полиноми по степените на времето. Като цяло остатъците, свързани с линейния модел, са по-големи, отколкото тези на втория (преоразмерен) модел (на фигурата случаи (а) и (б)) и съответно $\mathcal{F}_1(\theta) > \mathcal{F}_2(\theta)$. Но ако двете описания се приложат към нови данни (случай (в)), а тогава реа-

лизацията на случайната величина ще е друга, за целевата функция на двата модела е в сила $\mathcal{F}'_1(\theta) \approx \mathcal{F}'_1(\theta)$ и $\mathcal{F}'_2(\theta) < \mathcal{F}'_2(\theta)$. Причината е, че прецизната оптимизация на втория модел по отношение на неопределеността в първия набор не допринася за качеството на модела, когато той е приложен към новите данни. Освен това поради липсата на връзка с новото проявление на неопределеността качеството му допълнително се влошава. В резултат на това общовалидността на преоразмерения модел е значително по-малка в сравнение с тази на първия модел.



Фигура 2.43. Преоразмеряване на модела. Остатъци, свързани с достоверния (а) и преоразмерения (б) модел. Сравнение с други данни (в).

Освен измамното подобрене на модела, дължащо се на преоразмеряването, друг ефект от изкуственото нарастване на броя на параметрите е ненужното усложняване на структурата на модела, което води до по-трудното му използване, а понякога и до повече разходи.

2.4.6 Модифицирани показатели на качеството

Един вариант за оценяване на това, дали моделът е преоразмерен, е използването на модифицирани показатели на качеството. Съществува голямо разнообразие от такива статистики. За тях е характерно, че зависят както от целевата функция, така и от допълнителен наказателен член, който е свързан с броя на параметрите. При малък брой параметри доминира стойността на $\mathcal{F}(\theta)$, а с увеличаване на броя им допълнителният член нараства, т.е. увеличава се наказанието върху модела. В

началото усложняването на модела води до значително нарастване на достоверността му и въпреки нарастването на наказателния член, модифицираният показател намалява. Но от определен брой параметри, с увеличаване на броя им, стойностите на този показател започват да нарастват, тъй като ефектът от подобрението в модела (намаляването на $\mathcal{F}(\theta)$) става по-малък от нарастването на допълнителния член, както е показано на Фигура 2.44. В този случай се приема, че описанието започва да се преоразмерява. Така моделът с оптимална структура е този, който отговаря на минимума на конкретния показател. По-долу са дадени някои често използвани в практиката модифицирани показатели. Всеки от тях води до оптимален модел с различна структура и е въпрос на избор кой вариант ще се използва. Тези статистики не се следят самостоятелно, а се комбинират с описаните по-горе начини за валидация на модела.

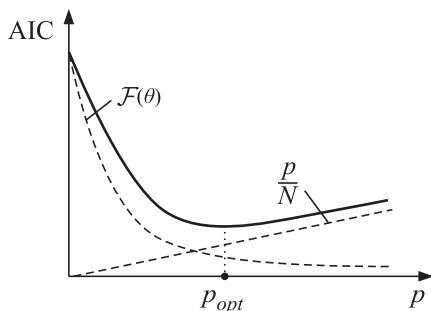
2.4.6.1 Информационен критерий на Акайке

Показателят AIC (Akaike Information Criterion) се дефинира по различен начин. Например в [11, 126, 148] той е представен съответно като:

$$\begin{aligned} \text{AIC} &= \ln \mathcal{F}(\theta) + \ln \left(1 + \frac{2p}{N} \right), \\ \text{AIC} &= \ln \mathcal{F}(\theta) + \frac{2p}{N}, \\ \text{AIC} &= \mathcal{F}(\theta) + \frac{p}{N}. \end{aligned} \tag{2.99}$$

В [57] подробно са изследвани особеностите на AIC и връзката му с други модифицирани показатели, а в [135, 141] са предложени варианти за малък брой наблюдения.

За да се използва правилно AIC и въобще показателите от тази група, удобно е предварително да се формира пълната матрица на данните Φ_{full} , съдържаща всички потенциални регресори. Ако моделът е статичен, редовете на Φ_{full} са N . При динамични модели, ако полиномите са от ред максимум n_{max} , то изходът на модела може да се формира в $N - n_{max}$ момента. След това, за да се определи $\mathcal{F}(\theta)$, се избира подмножество от стълбове на Φ_{full} , които отговарят на избраните до момента регресори. Така се гарантира, че независимо от това, колко е максималното закъснение n на сигнал в текущия модел, $\mathcal{F}(\theta)$ ще се формира на базата на $N - n_{max}$ стойности. В този случай сравняването на показате-



Фигура 2.44. Модифицираният показател на качеството AIC според формула (2.99)

нията на AIC за различни модели ще е коректно. Това важно най-вече когато броят на наблюденията е малък.

AIC често води до избор на по-сложни модели, тъй като наказателният му коефициент е по-слабо чувствителен към броя на параметрите.

2.4.6.2 Бейсов информационен критерий

Показателят BIC (Bayesian Information Criterion), наричан още Шварц критерий (SC – Schwarz Criterion), също има различни записи. Един от тях е

$$\text{BIC} = \ln \mathcal{F}(\theta) + \frac{p}{N} \ln N,$$

като отново $N \leftarrow N - n_{\max}$, ако моделът е динамичен. Особеност на този показател е, че може да приема отрицателни стойности. Въпреки това идеята му е същата като на AIC. BIC е по-рестриктивен по отношение на структурата на модела от AIC и минимумът му се достига при модел с по-малък брой параметри [106].

2.4.6.3 Статистика на Малоу

Този, както и останалите показатели от тази група, се използва при стъпковите методи за избор на структура на модела. Статистиката на Малоу (Malloy) Cp е

$$\text{Cp} = \frac{\mathcal{F}(\theta)}{\sigma_{e,full}^2} + 2p - N$$

($N \leftarrow N - n_{max}$ за динамични модели). $\sigma_{e,full}^2$ е средноквадратичната стойност на остатъка за модела $\hat{y}_{full,k} = \Phi_{full}\hat{\theta}$, съдържащ всички фактори, които потенциално може да участват за описание на изхода. В началото, с нарастване на p , първият член в Ср намалява по-бързо от втория (третото събираемо е константа). Правилото при работа с тази статистика е, че усложняването на модела трябва да се прекрати, когато Ср спадне под броя на факторите, добавени в модела, т.е., когато

$$Cp < p - 1.$$

Понякога Ср приема отрицателни стойности и в такъв случай става неизползваема. Причината е, че когато факторите не са много информативни, стойностите на $\mathcal{F}(\theta)$ и $\sigma_{e,full}^2$ са сходни и тогава тяхното отношение е по-малко от $|2p - N|$. По тази причина статистическите софтуери предоставят множество статистики, характеризиращи както модела като цяло, така и оценките, и факторите поотделно.

2.4.7 Кросвалидация

Друг вариант за избягване на преоразмеряването на модела е да се използват различни данни за оценяването и за проверката за неговата достоверност. Крайният модел обаче (особено при малък брой наблюдения) се препоръчва да се построи с използване на всички данни. Така за оценяване на параметрите му се използва цялата налична информация в данните.

2.4.7.1 Еднократна кросвалидация

За динамични системи се е установило правилото $\frac{2}{3}$ от данните да се използват за намиране на подходящи параметри на модела, а останалата $\frac{1}{3}$ да се използва за проверка на достоверността на модела. По този начин се гарантира, че валидацията на модела не зависи от данните, участващи във формирането на оценките. Това разделно изграждане и валидиране на модела се нарича кросвалидация.

Ако системата е статична, то подходящо е двата набора от данни да се формират на случаен принцип, като съотношението на броя на наблюденията за оценяване и за валидация може да е различно. Например за $N \gg 1$ може да се приеме $\frac{4}{5}N$ от наблюденията да се използват за оценяване и $\frac{1}{5}N$ от наблюдения – за валидация.

Понякога в процеса на моделиране се използват три набора от данни. Единият е за оценяване на параметрите (за обучение). Вторият набор е

тестов – с него, ако оценителят е реализиран с итеративна процедура, се изчисляват статистиките, които участват в критериите за спиране на оптимизацията, а ако се търсят структурните параметри със стъпков метод, този набор е нужен за вземане на решение за прекратяване на подобряването на структурата. Третият набор се използва за валидация на крайния модел.

Обикновено стойностите на показателите на качеството, определени от набора за оценяване, са по-оптимистични в сравнение със стойностите, получени от данните за валидация. Причината е, че с първия набор моделът се оптимизира така, че изходът му да се доближава максимално до тези данни. Ако разликата между стойностите на даден показател, получени с двата набора, е чувствителна, това означава, че моделът не е формиран коректно. Други причини за това (освен преоразмеряването) може да са: наличие на поведение в данните за валидация, което не се наблюдава в данните за формиране на модела; недостатъчно данни за разделно изграждане и валидиране на модела. Тогава може да се използват модифицираните показатели или да се приложи една от техниките, изложени по-долу.

2.4.7.2 Многократна кросвалидация

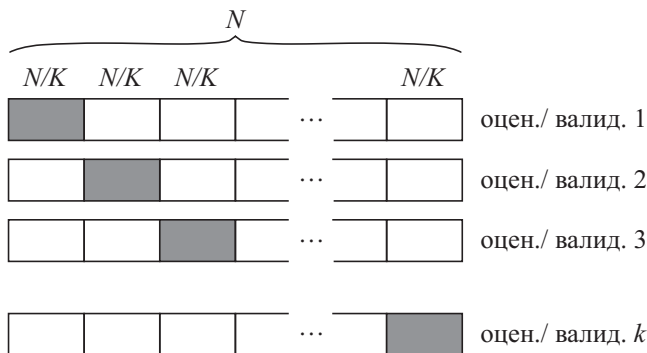
Разделянето на интервала на наблюдение на две или три части не е подходящо при набори от данни с малко наблюдения. За тези случаи е удачно да се използва многократна кросвалидация, при която данните се използват и за изграждане на модел, и за валидация.

К-кратна кросвалидация

Един вариант за проверка на универсалността на модела е целият набор от входно-изходни данни, означен с $\{u, y\}$ да се разделя на K на брой поднабора, които по-долу се записват като $\{u, y\}_i$ за $i = \overline{1, K}$ [140]. При този метод, наречен K -кратна кросвалидация, се изграждат K модела, като i -тият модел се построява с наблюденията от $K - 1$ поднабора от данни, в които не участва $\{u, y\}_i$, а валидацията му се извършва с наблюденията от $\{u, y\}_i$ (виж Фигура 2.45).

Общата целева функция $\mathcal{F}(\theta)$ се формира като средна стойност на показателите $\mathcal{F}(\hat{\theta}_{-i})$, получени за всеки един от моделите, т.е.

$$\mathcal{F}(\theta) = \frac{1}{K} \sum_{i=1}^K \mathcal{F}(\hat{\theta}_{-i}). \quad (2.100)$$



Фигура 2.45. Поднабори за оценяване на параметри (светли) и за валидация (тъмни) при K -кратната кросвалидация

$\hat{\theta}_{-i}$ са оценките на параметрите, получени с данните, от които е премахнат i -тият поднабор.

Този подход има предимството, че всички наблюдения се използват по равен брой пъти за оценяване (K на брой) и за валидация (по веднъж). Това се постига за сметка на многократното прилагане на тези две дейности.

Както беше споменато в началото на точка 2.4.7, с кросвалидацията се проверява дали модел със зададена структура би бил универсален, ако се построи с използване на наличните данни. В този смисъл (2.100) отразява именно това. Например с нарастване на броя на параметрите от един момент настъпва преоразмеряване. Тъй като i -тият модел не е оптимизиран за проявленията на неопределеността в набора $\{u, y\}_i$ (тези данни не се използват за формирането му), то членовете в сумата на (2.100) започват да нарастват, което е индикация за започналото преоразмеряване.

LOO кросвалидация

В граничния случай, когато $K = N$ (или $K = N - n$ за динамични системи), се получава т. нар. LOO (Leave One Out) кросвалидация [44]. При този вариант, когато моделът е линеен, е възможно показателят да се изведе аналитично, с което се избягва многократното оценяване и валидация.

За извеждане на показателя с избягване на явното оценяване на параметри нека $\mathcal{F}(\theta)$ за LOO кросвалидацията се представи по следния начин:

$$\begin{aligned}\mathcal{F}(\theta) &= \frac{1}{N-n} \sum_{k=n+1}^N \mathcal{F}(\hat{\theta}_{-k}) \\ &= \frac{1}{N-n} \sum_{k=n+1}^N (y_k - \varphi_k^T \hat{\theta}_{-k})^2 \\ &= \frac{1}{N-n} \sum_{k=n+1}^N (y_k - \varphi_k^T (\Phi_{-k}^T \Phi_{-k})^{-1} \Phi_{-k}^T y_{-k})^2, \quad (2.101)\end{aligned}$$

където Φ_{-k} и y_{-k} са матрицата на данните и векторът от стойностите на изхода за интервала на наблюдение, от които са премахнати съответно φ_k и y_k , а $\hat{\theta}_{-k}$ са оценките на параметрите, получени на базата на Φ_{-k} и y_{-k} . За вектора $\Phi_{-k}^T y_{-k}$ е в сила:

$$\Phi_{-k}^T y_{-k} = \Phi^T y - \varphi_k y_k. \quad (2.102)$$

Нека се положи $P = (\Phi^T \Phi)^{-1}$ и $P_{-k} = (\Phi_{-k}^T \Phi_{-k})^{-1}$. Лесно може да се провери (виж [78]), че връзката между P и P_{-k} е

$$P_{-k} = P + \frac{P \varphi_k \varphi_k^T P}{1 - \varphi_k^T P \varphi_k}.$$

Тази зависимост се използва за преобразуване на вектора $\varphi_k^T (\Phi_{-k}^T \Phi_{-k})^{-1}$, участващ в (2.101). За транспонирания му вариант се получава:

$$(\Phi_{-k}^T \Phi_{-k})^{-1} \varphi_k = P_{-k} \varphi_k = \frac{(\Phi^T \Phi)^{-1} \varphi_k}{1 - \varphi_k^T (\Phi^T \Phi)^{-1} \varphi_k}. \quad (2.103)$$

Замествайки (2.102) и (2.103) в (2.101), за текущия член на сумата се получава:

$$\begin{aligned}y_k - \varphi_k^T \hat{\theta}_{-k} &= y_k - \frac{(\Phi^T \Phi)^{-1} \varphi_k}{1 - \varphi_k^T (\Phi^T \Phi)^{-1} \varphi_k} (\Phi^T y - \varphi_k y_k) \\ &= \frac{y_k - \varphi_k^T (\Phi^T \Phi)^{-1} \Phi^T y}{1 - \varphi_k^T (\Phi^T \Phi)^{-1} \varphi_k}.\end{aligned}$$

Числителят на последния израз е i -тият елемент на вектора

$$\Phi^\perp y = (I - \Phi (\Phi^T \Phi)^{-1} \Phi^T) y.$$

Матрицата $\Phi^\perp$, както и ефектът ѝ върху y са дефинирани в точка 2.3.4.5

при интерпретацията на LS. Знаменателят е i -тият елемент на диагонала на $\Phi^\perp$. Следователно векторът на остатъците за целия интервал на наблюдение може да се запише така:

$$\begin{bmatrix} y_{n+1} - \varphi_{n+1}^T \hat{\theta}_{-(n+1)} \\ y_{n+2} - \varphi_{n+2}^T \hat{\theta}_{-(n+2)} \\ \vdots \\ y_N - \varphi_N^T \hat{\theta}_{-N} \end{bmatrix} = \Phi_d^{\perp-1} \Phi^\perp y,$$

където $\Phi_d^\perp$ е диагонална матрица, за която $\text{diag } \Phi_d^\perp = \text{diag } \Phi^\perp$. Отчитайки последния израз, показателят $\mathcal{F}(\theta)$ добива вида:

$$\mathcal{F}(\theta) = \frac{1}{N} y^T \Phi^\perp \Phi_d^{\perp-2} \Phi^\perp y.$$

Както се вижда, изразът за $\mathcal{F}(\theta)$ не съдържа оценки на параметрите, а само налични входно-изходни данни, с което се избягват N -те на брой оценявания на параметрите.

Има и други подходи за валидация, като Jackknife и Bootstrapping. Те са много ефективни при недостатъчен брой данни, но с тях трябва да се внимава. Името на подхода Bootstrapping идва от произведението „Приключенията на барон Мюнхаузен“ на Рудолф Ерих Распе и е свързано с момента, когато баронът се издърпва от блатото, като се хваща за косата си. По подобен начин, когато данните са недостатъчни, с Bootstrapping изкуствено се генерират голям брой извадки с желан брой наблюдения. Но тъй като първоначалните данни за поведението на системата са малко, крайните резултати трябва да се тълкуват внимателно. За повече информация може да се ползва [112].

Глава 3

Оценяване на параметри с линеен подход. Структура на модела

3.1 Параметри и структура на модела

В тази част на монографията се разглежда етапът на изграждане на модели, който включва дейностите по оценяването на параметрите и избора на структурата на модела. Трета глава засяга методи за изграждане на модел, при които на определено ниво се приема, че описанието е линейно параметризирано. Има група нелинейни по параметри модели, които при определени допускания може да се разглеждат като линейно параметризирани. Пример за такъв модел е ARMAX

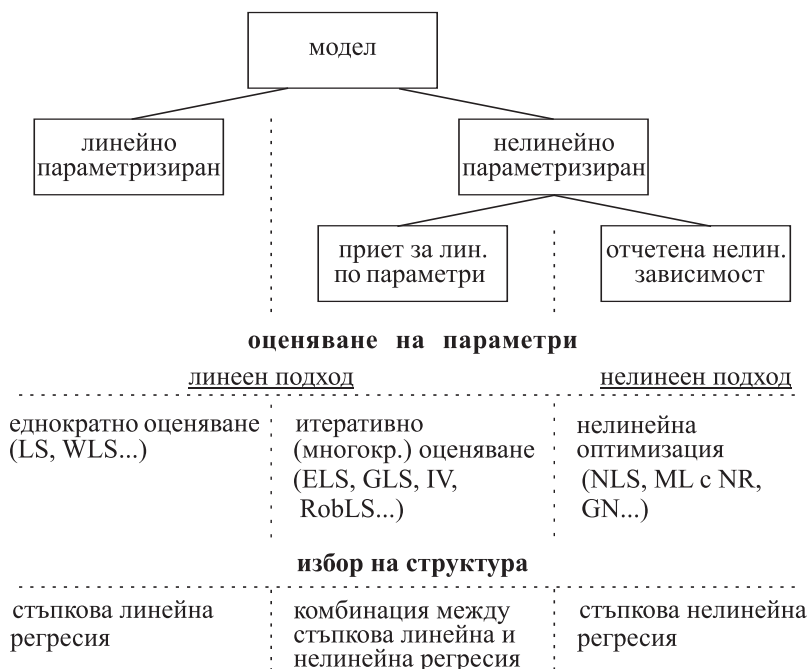
$$A(q^{-1})y_k = B(q^{-1})u_k + C(q^{-1})e_k. \quad (3.1)$$

Тъй като остатъкът е неотчетеното от модела поведение на системата, то e_k зависи от конкретните параметри на модела. С други думи, както се и очаква, при промяна на параметрите се променя и неотчетеното поведение на обекта, т.е. $e_k = e_k(\theta)$. За (3.1) се получава

$$A(q^{-1})y_k = B(q^{-1})u_k + C(q^{-1})e_k(\theta).$$

Последното събираемо е нелинейна функция на параметрите. Въпреки това има методи за оценяване на параметри, при които се приема, че e_k не зависи от θ , и при това положение ARMAX изкуствено се разглежда като линейно параметризиран модел.

Линейните по параметри модели са разгледани подробно в точка 2.1.3.2 от предишната глава. Причината за разграничаването на два вида модели, а именно линейни и нелинейни по параметри (Фигура 3.1), е принципната разлика между методите както за оценяване на параметри, така и за уточняване на структурата им.



Фигура 3.1. Модели и методи за оценяване на параметри

Ако моделът е или линеен, или нелинеен, но може да се запише в линейно параметризиран вид, оценяването се свежда до еднократно прилагане на LS или WLS.

Когато моделът е нелинеен по параметри, но условно може да се приеме, че връзката между y_k и θ е линейна, се разграничават два подхода за оценяване. Единият се свежда до многократно (итеративно) прилагане на LS, като тук моделът изкуствено се приема за линейно параметризиран (това е линейният подход за този клас модели). Такива са например методът на инструменталните променливи (IV – Instrumental Variable), методът на разширените най-малки квадрати (ELS – Extended LS)), методът на обобщените най-малки квадрати и др.

Третата възможност е моделът да е нелинеен и да е невъзможно или неудачно да се представи като линейна функция на параметрите. В този случай се прилага т.нар. нелинеен подход за оценяване [5, 92]. Той ще бъде разгледан в допълнителен труд. При него моделът се интерпретира като нелинеен по параметри, какъвто е в действителност,

и тогава оценяването е свързано с прилагането на методи за нелинейна оптимизация. Пример за този вариант е методът на максималното правдоподобие (ML – Maximum Likelihood), реализиран с метод за оптимизация като Нютон-Рафсън (NR – Newton-Raphson) [151, 42, 56] или реализацията на нелинейните най-малки квадрати (NLS – Nonlinear LS) с метода Гаус-Нютон (GN – Gauss-Newton) [11, 146, 143].

По отношение на методите за избор на структурните параметри също има принципна разлика, зависеща от това, дали параметрите се оценяват еднократно или итеративно. Във втория случай методите за уточняване на структурните параметри са значително по-сложни. В монографията са засегнати методи за избор на структура, приложими за линейно параметризирани модели.

В следващите точки са изложени методите за оценяване както от първата група (LS и WLS), така и от втората (GLS, ELS, IV и робастният метод на най-малките квадрати (RobLS – Robust LS)). След представянето на методите са разгледани и числените проблеми, които може да възникнат при изпълнение на съответните алгоритми. Обсъдени са и източниците на тези проблеми, както и ефективни числени реализации на методите за оценяване на параметри. Главата завършва с методи за уточняване на структурата на линейно параметризирани модели. Разгледани са структурните параметри, характерни за многомерните системи. Също така е изложена идеята на стъпковите методи [64, 109, 133] за избор на структура. Освен това са представени и високоефективни алгоритми за права, обратна и комбинирана стъпкови регресии.

3.1.1 Оценяване на линейно параметризирани модели

В точка 2.4.1 е споменато, че ако функцията на правдоподобие L_θ не е предварително известна, често се приема, че тя отговаря на нормално разпределение, т.е.

$$L_\theta = \frac{1}{(2\pi)^{\frac{\ell N}{2}} \det \Sigma_e} e^{-\frac{1}{2} e^T \Sigma_e^{-1} e}. \quad (3.2)$$

В много приложения това приемане е удовлетворително. Когато остатъкът e е с Гаусово разпределение, оценките, получени по LS, са ефективни, а оценителят е BLUE (виж още точка 3.2.2.4). В противен случай дори остатъкът да е БШ, ако не е Гаусов, няма гаранция, че дисперсията на грешката на оценките е минимална. Поради тези и други причини

(споменати в темата за валидация) в Трета глава се приема, че ϵ е с нормално разпределение. Така методът ML, минимизиращ в общия случай L_θ , се свежда до LS, WLS или до GLS, ELS, IV и др., когато се наблюдава корелация между стойностите на остатъка. Дори когато в данните се съдържат нехарактерни стойности, водещи до нарушаване на приемането за нормалност, с помощта на робастни оценители е възможно, макар и итеративно, да се избегне изместването на оценките, дължащо се на тези нехарактерни отчети.

3.1.2 Оценяване на един MIMO или множество MISO модели

Методите за оценяване на параметри са представени за MIMO модели. За разлика от валидацията, където към всеки изход може да се подходи поотделно, на етапа на същинското моделиране е удачно да се работи с многоизходови модели. Така, вместо многократно да се извършва оценяване на MISO подмодели, директно се формира многоизходовото описание. В [120] са обсъдени някои особености на подхода с декомпозицията на MISO подмодели, както и са споменати част от предимствата от едновременното описание на всички изходи. Когато моделът се използва за предсказване, по-добри резултати се получават, ако моделът е MIMO. Причината е, че предисторията на всички изходи осигурява повече информация за поведението на системата в сравнение със случая, когато се използва предисторията само на един от изходите.

Освен споменатото в [120] предимство друга причина за разглеждането на системата като многоизходова е възможността структура на модела коректно да се определи при отчитане на ограничения като минимален брой входове, максимална размерност на полиноми в описанието, а също и когато има различни изисквания за достоверност по всеки изход. При работа с MIMO модели по естествен начин се следи изменението на поведението в модела по всички изходи, докато при оценяването на MISO модели се загубва цялата картина.

Когато добавянето на входове е свързано със средства, за да не се оскъпи излишно употребата на модела, се търси описание с възможно по-малко входове. В такъв случай, ако подобрението на модела по първия изход се постига с включване на определен вход, но подобрение по другите изходи не се наблюдава, то е удачно да се потърси друг вход, който, макар и не най-подходящ за описание на първия изход, като цяло

да максимизира близостта между изходите на модела и тези на системата.

В банковия сектор понякога се разработват модели на поведението на клиент според различни стратегии на банката, като всеки изход отразява поведението му при съответна стратегия. Тъй като данните за част от входовете се купуват, банката би предпочела описание, което използва минимален брой закупувани входни величини. За построяването на такива модели е необходимо системата да се разглежда като ММО, а не да се декомпозира на MISO подсистеми. Причината е, че достоверността на модела трябва да се максимизира като цяло – по всички изходи. Но ако при изграждането му не се отчете това изискване, т.е. оптимизацията на структурата се извършва независимо по всеки изход, тогава броят на входовете може ненужно да нарасне, като сходни величини се използват за описание на отделните изходи.

Когато изследваната система е с голям брой изходи, наборът от данни се състои от много наблюдения, и особено когато идентификацията се автоматизира, използването на ММО модел може значително да ускори цялостния процес на моделиране.

Ефектът от оценяването на ММО описание, а не на множество MISO модели, се подсилва при реализацията на оценителите с технологии за паралелни изчисления като Open CL (Open Computing Language) и CUDA (Computing Unified Device Architecture) [91], с които е възможно да се използват графичните процесори (GPU) за целите на идентификацията. CUDA е архитектура, която позволява на различни приложения да използват процесори като CPU и GPU за изчислителните операции, а Open CL е свободният (нелицензиран) вариант на CUDA. GPU са оптимизирани за работа с масиви. Голяма част от обработката на данните при изпълнение на алгоритмите за ММО модели може да се извършва паралелно и поради наличието на стотици ядра в тези процесори, времето за оценяване на параметрите би могло драстично да намалее.

Много важно условие за ефективното използване на GPU е обработката да се декомпозира на голям брой независими, но задължително еднотипни дейности. LS и неговите обобщения са подходящи за това, тъй като оценките на параметрите са резултат от умножението на големи матрици, съдържащи наличните данни. Всеки елемент от резултантните матрици се формира независимо от останалите елементи като сума от голям брой произведения на съответни елементи на първона-

чалните матрици. В допълнение паметта на съвременните GPU е от порядъка на гигабайти, което позволява зареждането на големи масиви от данни в паметта на GPU. С ефективно използване на глобалната и бързата локална памет на процесора е възможно изчислителният процес допълнително да се ускори. Възможността процесът на оценяване да се разпаралели, е още една причина да се предпочита изграждането на MIMO, а не на множество MISO модели.

Недостатък на описаната технология е зависимостта от конструктивните особености на използваното GPU. До тези ограничения лесно може да се стигне, когато броят на данните е значителен. В този смисъл при проектирането на даден алгоритъм, за максимизиране на бързодействието, трябва да се отчитат специфичните за конкретното GPU ограничения. Освен това, въпреки възможността от разпаралеляване на изчислителните операции в оценителите, поради несложните действия, които се извършват паралелно в LS, WLS и др., времето за формиране на оценките не намалява в степента, в която би намаляло при изграждането на някои нелинейни модели, където паралелно се изпълняват голям брой сложни операции. Причината е, че подготовката на GPU за паралелната обработка отнема време и в някои случаи цялостното време за формиране на модел може да е дори съизмеримо с това, когато се използва CPU.

3.2 Методи за модел с матрица на параметрите

Характерно за представянето на моделите в линейно параметризиран вид е, че параметрите може да се отделят от регресорите, и при квадратична целева функция за изчисляването на оптималните оценки да не се използват итеративни реализации на методите за оценяване. Например беше показано, че методът NR, приложен за описаната постановка на задачата, намира оптималните оценки още в първата итерация, при това независимо от началните условия. Това значително опростява процеса на изграждане на модел както по време на оценяването на параметри, така и при избора на структурата му. В тази и в следващата точка се разглеждат методи, изведени за общото представяне на модела както с матрица, така и с вектор на параметрите. Също се обръща внимание на особеностите на задачата за оценяване при всяко от двете представяния.

3.2.1 Модел в общ вид с матрица на параметрите

По-долу е прието, че изходът в даден момент е вектор с ℓ компоненти. Тогава линейният регресионен модел може да се представи в общия вид:

$$y_k = \Theta^T \varphi_k + e_k. \quad (3.3)$$

При този запис параметрите на модела са подредени в матрицата $\Theta \in \mathcal{R}^{z \times \ell}$, а векторът $\varphi_k \in \mathcal{R}^z$ съдържа регресорите, описващи изхода в k -тия такт. В [67, 5, 92, 152, 154, 66] се използва такова представяне на регресионния модел.

Нека за представяне на MIMO системата се използва ARX модел

$$A(q^{-1})y_k = B(q^{-1})u_k + e_k. \quad (3.4)$$

Полиномните матрици са $A(q^{-1}) \in \mathcal{R}_{na}^{\ell \times \ell}$ и $B(q^{-1}) \in \mathcal{R}_{nb}^{\ell \times m}$. Във Втора глава са дадени две преобразувания на (3.4) от полиномен в общ вид, като се стартира от текущото описание на системата (в k -тия такт). След достигане до общия вид моделът се обобщава за интервала на наблюдение с цел получаване на матрично уравнение, в което в явен вид се обвързват данните, търсените параметри и остатъкът, който е необходим за формулиране на целта, заложена в оценителя.

Преобразуванията в общ вид във Втора глава са интуитивни от гледна точка на протичането на процесите във времето. По-долу, отново от (3.4), се достига до общото представяне, но начинът, по който се формират матриците на данните, е удобен за реализацията на оценителя. В този случай се работи с индекси и се избягва използването на цикли. Това е удобно при работа с GPU.

Нека двете страни на (3.4) се транспонират. В резултат на това се получава:

$$y_k^T A^T(q^{-1}) = u_k^T B^T(q^{-1}) + e_k^T. \quad (3.5)$$

Нека също се въведат матриците на изходните данни $Y \in \mathcal{R}^{N-n \times \ell}$, на входните данни $U \in \mathcal{R}^{N-n \times m}$ и на остатъка $E \in \mathcal{R}^{N-n \times \ell}$, които са:

$$Y = [y_{n+1} \ y_{n+2} \ \dots \ y_N]^T, \quad (3.6)$$

$$U = [u_{n+1} \ u_{n+2} \ \dots \ u_N]^T, \quad (3.7)$$

$$E = [e_{n+1} \ e_{n+2} \ \dots \ e_N]^T, \quad (3.8)$$

където $n = \max(na, nb)$ е броят предишни тактове, в които са необходими данни за формирането на изхода в текущия момент. В такъв случай

ARX моделът може да се запише за интервала на наблюдение като:

$$YA^T(q^{-1}) = UB^T(q^{-1}) + E. \quad (3.9)$$

Последният израз представлява система от $N - n$ уравнения от вида (3.4), за $k = \overline{n+1, N}$. Нека също се дефинират матриците:

$$Y_{-i} = q^{-i}Y = [y_{n+1-i} \ y_{n+2-i} \ \dots \ y_{N-i}]^T, \quad (3.10)$$

$$U_{-i} = q^{-i}U = [u_{n+1-i} \ u_{n+2-i} \ \dots \ u_{N-i}]^T, \quad (3.11)$$

$$E_{-i} = q^{-i}E = [e_{n+1-i} \ e_{n+2-i} \ \dots \ e_{N-i}]^T, \quad (3.12)$$

за $i = \overline{1, n}$. Имайки предвид представянето на полиномните матрици като полиноми с матрични коефициенти по степените на q^{-1} , зависимостта (3.9) може да се запише като:

$$Y + Y_{-1}A_1^T + \dots + Y_{-na}A_{na}^T = U_{-1}B_1^T + \dots + U_{-nb}B_{nb}^T + E.$$

След изразяване на Y за модела, записан за интервала на наблюдение, се получава:

$$Y = \Phi\Theta + E. \quad (3.13)$$

В това общо представяне матриците на параметрите $\Theta \in \mathcal{R}^{z \times \ell}$ и на данните $\Phi \in \mathcal{R}^{N-n \times z}$ са:

$$\Theta = [A_1 \ A_2 \ \dots \ A_{na} \ B_1 \ B_2 \ \dots \ B_{nb}]^T,$$

$$\Phi = [-Y_{-1} \ -Y_{-2} \ \dots \ -Y_{-na} \ U_{-1} \ U_{-2} \ \dots \ U_{-nb}]. \quad (3.14)$$

Размерността $z = \ell na + m \ nb$ отговаря на броя на регресорите, необходими за формирането на всеки от изходите на модела в даден момент от времето, а броят на параметрите е $p = z\ell$.

Когато ефективната реализация на метода изисква работа с масиви, е удачно, работейки с индекси, да се извлекат матриците U_{-i} и Y_{-i} , а формирането на матрицата Φ да се извърши съгласно (3.14).

Представеното описание на модела с матрица на параметрите е поинтуитивно, в смисъл че преобразуването на модела в общ вид следва същите стъпки като за SISO модели. Освен това формирането на матрицата Φ за този случай се извършва значително по-лесно в сравнение с представянето с вектор на параметрите (точка 3.3), тъй като тя зависи от матриците U_{-i} и Y_{-i} , които се формират с просто изместване на данните за входа и изхода със съответен брой тактове.

Обикновено оценителите, изведени за това представяне, се използват за бързо получаване на модел, с цел изследователят да се ориентира в това, какво може да се извлече от данните. Също модел с такава структура е удачен, когато се търси описание с колкото се може по-малко входове. Така, ако се добави нов вход, то е логично той да се използва за описанието на всички изходи.

Основният недостатък на представянето на модела с матрица на параметрите е по-малката свобода при избора на структурните параметри (редът на полиномите в описанието). За конкретната структура на матрицата на параметрите се вижда, че редът на полиномите по отношение на всеки вход е nb , а полиномите, отчитащи авторегресията по отношение на изходите, са от ред na . С други думи при така конструираната матрица не е възможно да се формира модел с полиноми, в рамките на дадена полиномна матрица, от различен ред.

В точка 2.1.3.3 е даден вариант с матрица на параметрите, при който има по-голяма свобода, като ограничението върху структурата е, че влиянието на предисторията на даден вход или изход върху стойностите на изходите се описва с полиноми от еднакъв ред. Например влиянието на j -тия вход върху всички изходи задължително се описва с полиноми от ред nb_j .

Независимо от конкретния начин на представяне винаги съществува известно структурно ограничение, което води до недостатъчно параметри по някои канали и/или до наличието на излишни параметри по други канали. Следователно някои взаимовръзки в системата няма да са до-описани или структурата на модела ненужно ще е усложнена, като в някои случаи това може да доведе до преоразмеряване на модела, а понякога и до лоша обусловеност на задачата за оценяване. Този недостатък не възниква при SISO модели, но когато описанието е с много входове и изходи, проблемите може да доведат (ако не са взети предпазни мерки) дори до невъзможност да бъде намерен адекватен модел. В такъв случай едно решение е да се използва оценител, изграден на базата на модел с вектор на параметрите, който е разгледан в точка 3.3.

3.2.2 Метод на най-малките квадрати

По-долу е представена постановката на задачата на LS за оценяване параметрите на модел, представен с матрица на параметрите. Методът е изведен и са анализирани свойствата на оценките.

3.2.2.1 Показател на качеството

При най-малките квадрати се минимизира сумата от квадратите на остатъците. Когато системата е с много изходи, целевата функция (показател на качеството) е сума от остатъците по отношение на всеки един от изходите, т.е.

$$\mathcal{F}(\Theta) = \sum_{k=n+1}^N \sum_{i=1}^{\ell} e_{i,k}^2.$$

Ако се използва матрицата на остатъците, дефинирана в (3.12), то i -тият диагонален елемент на $E^T E$ е сумата от квадратите на i -тия остатък, или $[E^T E]_{ii} = \sum_{k=n+1}^N e_{i,k}^2$. Така за $\mathcal{F}(\Theta)$ може да се използва по-удобният за извежданията запис:

$$\mathcal{F}(\Theta) = \text{tr}(E^T E),$$

където $\text{tr}(\cdot)$ е следа на матрица. От (3.13) матрицата на остатъка може да се изрази чрез известните матрици на данните Y и Φ и търсената матрица на параметрите Θ . Получава се:

$$E = Y - \Phi\Theta.$$

В литературата, свързана с идентификацията на системи, често се използва следният запис на целевата функция:

$$\mathcal{F}(\Theta) = \|E\|_F^2 = \|Y - \Phi\Theta\|_F^2.$$

В него отпада остатъкът, който е неизвестен преди оценяването на параметрите, и така $\mathcal{F}(\Theta)$ явно зависи от търсените параметри и наличните данни, което е стъпка към определянето на оптималните оценки. В горния запис $\|\cdot\|_F$ е Фробениус (Frobenius) норма на матрица. Нека A е произволна матрица с реални елементи. В такъв случай $\|A\|_F$ е квадратен корен от сумата от квадратите на елементите на A , т.е. $\|A\|_F = \sqrt{\text{tr}(A^T A)}$. Записът с тази норма се използва при представянето на някои числени реализации на методите за оценяване на параметри.

В долните извеждания целевата функция се представя с по-краткия запис $\mathcal{F}(\Theta) = \text{tr}(E^T E)$. Тогава критерият, заложен в LS за модел с матрица на параметрите, ще се записва по следния начин:

$$\min_{\Theta} \mathcal{F}(\Theta) = \min_{\Theta} \text{tr}(E^T E). \quad (3.15)$$

3.2.2.2 Оценки по LS

Определянето на матрицата на параметрите Θ , при която показателят на качеството има минимум, се извършва, като първо $\mathcal{F}(\Theta)$ се представи във вид, удобен за диференциране, после се определи градиентната матрица $\nabla\mathcal{F}(\Theta)$ и накрая се изрази $\hat{\Theta}$, при която $\nabla\mathcal{F}(\Theta)$ се нулира.

Целевата функция зависи от матричен аргумент и е следа на матрица. Затова в извеждането по-долу се въвежда понятието градиентна матрица и се дават някои зависимости, свързани с диференцирането на следа на матрица.

Градиентна матрица

Аналогът на първата производна на скаларната функция $\mathcal{F}(\Theta)$ по отношение на матричния аргумент Θ се нарича градиентна матрица (в някои източници се нарича градиент). Елементите на тази матрица са първите частни производни на целевата функция по отношение на параметрите, т.е. ij -тият елемент е $\frac{\partial\mathcal{F}}{\partial\theta_{ij}}$.

Градиентната матрица има вида:

$$\nabla\mathcal{F}(\Theta) = G = \begin{bmatrix} \frac{\partial\mathcal{F}}{\partial\theta_{11}} & \frac{\partial\mathcal{F}}{\partial\theta_{12}} & \cdots & \frac{\partial\mathcal{F}}{\partial\theta_{1\ell}} \\ \frac{\partial\mathcal{F}}{\partial\theta_{21}} & \frac{\partial\mathcal{F}}{\partial\theta_{22}} & \cdots & \frac{\partial\mathcal{F}}{\partial\theta_{2\ell}} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial\mathcal{F}}{\partial\theta_{z1}} & \frac{\partial\mathcal{F}}{\partial\theta_{z2}} & \cdots & \frac{\partial\mathcal{F}}{\partial\theta_{z\ell}} \end{bmatrix}.$$

Диференциране на следа на матрица

При определянето на елементите на G се използват някои зависимости, валидни за диференцирането на следа на матрица, зависеща от матричен аргумент. Нека са дадени матриците $A \in \mathcal{R}^{m \times m}$, $B \in \mathcal{R}^{m \times n}$ и $X \in \mathcal{R}^{m \times n}$. Тогава са в сила съотношенията:

$$\nabla_X \operatorname{tr}(B^T X) = B, \quad (3.16)$$

$$\nabla_X \operatorname{tr}(X^T B) = B, \quad (3.17)$$

$$\nabla_X \operatorname{tr}(X^T A X) = (A + A^T)X. \quad (3.18)$$

Свойствата (3.16) и (3.17) лесно се проверяват, тъй като транспонирането на матрицата не променя главния ѝ диагонал (за произволна матрица C , $\text{tr} C = \text{tr} C^T$). Ако A е симетрична, какъвто е случаят при диференцирането на $\mathcal{F}(\Theta)$, (3.18) се свежда до

$$\nabla_X \text{tr}(X^T A X) = 2AX. \quad (3.19)$$

По-долу е описан начинът, по който се определя оптималната матрица $\hat{\Theta}$ като функция на наличните входно-изходни данни.

Извеждане на LS

Представянето на $\mathcal{F}(\Theta)$ във вид, удобен за диференциране по отношение на Θ , се получава по следния начин:

$$\begin{aligned} \mathcal{F}(\Theta) &= \text{tr}(E^T E) \\ &= \text{tr}((Y - \Phi\Theta)^T (Y - \Phi\Theta)) \\ &= \text{tr}(Y^T Y) - \text{tr}(\Theta^T \Phi^T Y) - \text{tr}(Y^T \Phi\Theta) + \text{tr}(\Theta^T \Phi^T \Phi\Theta). \end{aligned}$$

След диференциране на $\mathcal{F}(\Theta)$, понеже първият член на горната сума не зависи от параметрите на модела, то

$$\nabla_{\Theta} \text{tr}(Y^T Y) = 0.$$

От свойствата (3.16)–(3.17) за втория и третия член следва, че

$$\nabla_{\Theta} \text{tr}(\Theta^T \Phi^T Y) = \nabla_{\Theta} \text{tr}(Y^T \Phi\Theta) = \Phi^T Y.$$

Тъй като последният член е квадратичен по отношение на параметрите, от (3.19) се получава:

$$\nabla_{\Theta} \text{tr}(\Theta^T \Phi^T \Phi\Theta) = 2\Phi^T \Phi\Theta.$$

Тогава градиентната матрица G добива вида:

$$G = -2\Phi^T Y + 2\Phi^T \Phi\Theta.$$

За определяне на оптималните параметри изразът за G се приравнява на нула, т.е.

$$G|_{\Theta=\hat{\Theta}} = 0_{z \times \ell} \Leftrightarrow -\Phi^T Y + \Phi^T \Phi\hat{\Theta} = 0_{z \times \ell}.$$

Матрицата $\hat{\Theta}$, за която G се нулира, е

$$\hat{\Theta} = (\Phi^T \Phi)^{-1} \Phi^T Y. \quad (3.20)$$

Тя съдържа оценките на параметрите, определени по метода LS.

3.2.2.3 Свойства на оценките

В общия случай важни свойства на оценките са: състоятелност (консистентност), неизместеност, ефективност, устойчивост, сходимост, следене, изглаждане и др. Последните две свойства се изследват, когато оценките се получават от итеративни, а последните четири – от рекурсивни алгоритми. В случая поведението на оценките се разглежда с нарастване на броя на итерациите или на дискретните моменти от времето. Тъй като в настоящата тема се разглежда нерекурсивно оценяване на параметри, ще бъдат изследвани свойствата, които имат отношение към блочния (нерекурсивен) подход, а именно състоятелност, неизместеност и ефективност.

Състоятелност

Оценката $\hat{\theta}_{ij}$ е състоятелна, ако с нарастване на броя на наблюденията, тя сходяща към оптималната стойност θ_{ij}^* на параметъра, т.е.

$$\lim_{N \rightarrow \infty} \hat{\theta}_{ij} = \theta_{ij}^*.$$

Неизместеност

Неизместена оценка е тази, при която независимо от броя на наблюденията, използвани за формирането ѝ, е в сила:

$$M\{\hat{\theta}_{ij}\} = \theta_{ij}^*.$$

Възможно е една изместена оценка да е състоятелна, както и обратното – несъстоятелна оценка да е неизместена. С долните два примера се показват такива случаи, както и допълнително се изясняват понятията състоятелност и изместеност.

Пример. Състоятелна оценка, която е изместена

Нека x_k е ергодичен сигнал с математическо очакване m_x , а векторът $x_c = [x_{c,1} \ x_{c,2} \ \dots \ x_{c,N}]^T$ съдържа центрираните стойности на x_k , които са $x_{c,k} = x_k - m_x$. Тогава оценката на дисперсията

$$s_x^2 = \frac{1}{N} x_c^T x_c$$

е изместена (тя би била неизместена, ако знаменателят е $N - 1$, понеже е загубена една степен на свобода при центрирането на сигнала), но е състоятелна, защото с нарастване на броя на наблюденията s_x^2 сходяща към истинската дисперсия, или

$$\lim_{N \rightarrow \infty} s_x^2 = \sigma_x^2.$$

Обяснението е, че при голямо N намаляването на знаменателя с една степен на свобода е пренебрежимо спрямо броя на наблюденията. $\diamond$

Пример. *Несъстоятелна оценка, която е неизместена*

Нека ε е случаен сигнал с математическо очакване $m_\varepsilon = 0$. Тогава оценката

$$s_x^2 = \frac{1}{N-1} x_c^T x_c + \varepsilon$$

на дисперсията на x_k от предишния пример е неизместена, понеже, независимо от броя на наблюденията, използвани за формирането ѝ,

$$M\{s_x^2\} = \sigma_x^2 + m_\varepsilon = \sigma_x^2.$$

Но тази оценка не е състоятелна, тъй като не схожда към истинската дисперсия. По-точно е изпълнено:

$$\lim_{N \rightarrow \infty} s_x^2 = \sigma_x^2 + \varepsilon. \quad \diamond$$

Ефективност

Ако оценките имат минимална дисперсия, то те са ефективни. Тъй като се работи с краен брой наблюдения, интерес в идентификацията представляват най-вече свойствата неизместеност и ефективност, понеже, ако входно-изходните величини са ергодични, неизместените оценки по LS са и състоятелни. В темата за оценяването на параметри във Втора глава беше въведена идеята за най-добър, линеен, неизместен оценител (съкратено BLUE от Best Linear Unbiased Estimator). Този оценител осигурява неизместени оценки, които имат минимална дисперсия. По-долу са разгледани по-подробно свойствата неизместеност и ефективност.

3.2.2.4 Анализ на оценките по LS

Грешка в оценяването на параметрите

При анализа на свойствата на оценките, извършен по-долу, първо се дефинира грешката $\tilde{\Theta}$ в оценяването на параметрите. След това се изследва влиянието ѝ върху свойствата на оценките. Анализът завършва с уточняване на условията, при които оценките са неизместени и ефективни. За тази цел е необходимо да се определят математическото очакване $m_{\tilde{\Theta}}$ и дисперсията $\sigma_{\tilde{\Theta}}^2$ на грешката в оценяването.

Нека решението (3.20) по LS се развие по следния начин:

$$\begin{aligned}\hat{\Theta} &= (\Phi^T \Phi)^{-1} \Phi^T Y \\ &= (\Phi^T \Phi)^{-1} \Phi^T (\Phi \Theta^* + E) \\ &= \Theta^* + (\Phi^T \Phi)^{-1} \Phi^T E.\end{aligned}\tag{3.21}$$

С Θ^* е означена неизвестната матрица на оптималните неизместени параметри. Второто събираемо в последното равенство е матрицата:

$$\tilde{\Theta} = (\Phi^T \Phi)^{-1} \Phi^T E,\tag{3.22}$$

която е всъщност допуснатата грешка в оценяването на параметрите.

Неизместеност на $\hat{\Theta}$

За да се определи връзката между математическото очакване на оценките и неизвестните оптимални стойности, трябва да се уточни влиянието на грешката $\tilde{\Theta}$ върху $m_{\hat{\Theta}}$. В резултат на това ще се определят условията, при които $\hat{\Theta}$ е неизместена. Това се постига, като двете страни на (3.21) се усреднят. Тъй като Θ^* са оптималните параметри, те не зависят от конкретните данни и тогава $m_{\Theta^*} = \Theta^*$. Така, с използване на (3.22), се стига до равенството:

$$m_{\hat{\Theta}} = \Theta^* + m_{\tilde{\Theta}}.\tag{3.23}$$

Условието оценките да бъдат неизместени, е математическите им очаквания да съвпадат с оптималните параметри, т.е. е необходимо $m_{\hat{\Theta}} = \Theta^*$, което е равносилно на $m_{\tilde{\Theta}} = 0$. Ако сигналът e_k не е корелиран с регресорите, т.е. е бял шум, то за математическото очакване на грешката в оценяването на параметрите се получава:

$$m_{\tilde{\Theta}} = M\{(\Phi^T \Phi)^{-1} \Phi^T E\} = M\{(\Phi^T \Phi)^{-1} \Phi^T\} m_E.$$

Тъй като елементите на Φ са стойностите на входа и изхода за интервала на наблюдение, между тях съществува корелация. Това означава, че $M\{(\Phi^T \Phi)^{-1} \Phi^T\} \neq 0$. Тогава, за да бъде $m_{\tilde{\Theta}} = 0$, е необходимо $m_E = 0$. От горните разсъждения следва, че LS осигурява неизместено оценяване, когато e_k е центриран бял шум.

Ефективност на $\hat{\Theta}$

Нека ковариационната матрица на параметрите от i -тия стълб на $\hat{\Theta}$ е $\Sigma_{\Theta, i}$, а ковариационната матрица на грешката в оценяването на тези параметри е $\Sigma_{\tilde{\Theta}, i}$. Връзката между двете матрици, при условие че

оценките са неизместени, може да се изрази по следния начин:

$$\Sigma_{\hat{\Theta}.i} = M\{(\hat{\Theta}.i - m_{\hat{\Theta}.i})(\hat{\Theta}.i - m_{\hat{\Theta}.i})^T\} \quad (3.24)$$

$$= M\{(\hat{\Theta}.i - \Theta_{.i}^*)(\hat{\Theta}.i - \Theta_{.i}^*)^T\} \quad (3.25)$$

$$= M\{\tilde{\Theta}.i \tilde{\Theta}.i^T\} = \Sigma_{\tilde{\Theta}.i}. \quad (3.26)$$

От (3.22) за ковариационната матрица на грешката се получава:

$$\Sigma_{\hat{\Theta}.i} = M\{(\Phi^T \Phi)^{-1} \Phi^T E_i E_i^T \Phi (\Phi^T \Phi)^{-1}\}.$$

Ако e_k е бял шум, регресорите (елементи на Φ), които формират i -тия изход на модела, са независими от остатъка на i -тия изход (участващ при формирането на вектора E_i). Тогава ковариационната матрица на грешката $\tilde{\Theta}$ може да се запише като:

$$\Sigma_{\tilde{\Theta}.i} = M\{(\Phi^T \Phi)^{-1} \Phi^T \Sigma_{E_i} \Phi (\Phi^T \Phi)^{-1}\}.$$

Матрицата $\Sigma_{E_i} = M\{E_i E_i^T\} \in \mathcal{R}^{N-n \times N-n}$ е диагонална, когато e_k е центриран бял шум. Ако освен това e_k е ергодичен, то диагоналните елементи на Σ_{E_i} са равни (нека стойността им е $\sigma_{e,i}^2$). Тогава

$$\Sigma_{E_i} = \sigma_{e,i}^2 I_{N-n} \quad (3.27)$$

и за $\Sigma_{\tilde{\Theta}.i}$ се получава:

$$\Sigma_{\tilde{\Theta}.i} = M\{(\Phi^T \Phi)^{-1}\} \sigma_{e,i}^2.$$

При достатъчно голям брой наблюдения операцията „математическо очакване“ може да се пропусне, защото елементите на $\Phi^T \Phi$ са резултат от сумирането на произведенията между факторите за интервала на наблюдение, при което влиянието на случайната компонента, намираща се в данните, намалява. Матрицата $F = \Phi^T \Phi$ се нарича информационна матрица или матрица на Фишер. Нека нейната обратна е $P = (\Phi^T \Phi)^{-1}$. В LS не се разчита на априорна информация за дисперсията на неопределеността, затова подходяща оценка на $\sigma_{e,i}^2$ е средноквадратичната стойност на остатъка за интервала на наблюдение. Тъй като често се приема, че σ_e^2 не се променя в рамките на експеримента поради връзката $\Sigma_{\hat{\Theta}.i} = \Sigma_{\tilde{\Theta}.i} \propto P$, за $i = \overline{1, \ell}$, в литературата е прието P да се нарича ковариационна матрица на параметрите.

Интерес за анализа на оценките представляват диагоналните елементи на $\Sigma_{\tilde{\Theta}.i}$, които са всъщност дисперсиите на елементите на $\tilde{\Theta}.i$. Диагоналът на $\Sigma_{\tilde{\Theta}.i}$ е векторът

$$\sigma_{\tilde{\Theta}.i}^2 = \sigma_{\Theta.i}^2 = \text{diag } P \sigma_{e,i}^2.$$

Ако горният израз се обобщи за всички стълбове на матрицата на оценките, за дисперсията в оценяването на параметрите се получава:

$$\sigma_{\hat{\Theta}}^2 = \text{diag } P \sigma_e^{2T}. \quad (3.28)$$

В последния запис $\text{diag } P$ е z -мерен, σ_e^2 е ℓ -мерен вектор и съответно дисперсията е $z \times \ell$ мерна матрица, като $[\sigma_{\hat{\Theta}}^2]_{ij}$ е дисперсията на ij -тия елемент на $\hat{\Theta}$. В някои софтуери вместо дисперсията се представя стандартното отклонение $\sigma_{\hat{\Theta}}$, наричано още стандартна грешка (разгледано в точка 2.4.4.1). В LS се минимизира целева функция, която е сумата от квадратите на остатъците по всеки изход, т.е. критерият е

$$\min_{\Theta} \text{tr}(E^T E).$$

Нека e_k е ергодичен центриран бял шум, а също нека е и Гаусов. Тогава е в сила $\text{tr}(E^T E) \propto \sigma_e^2$ и тогава LS води до

$$\min_{\Theta} \sigma_e^2. \quad (3.29)$$

При фиксиран набор от данни и структура на модела $\text{diag } P$ в (3.28) е константен вектор. От друга страна, дисперсията на остатъчната грешка се влияе от стойностите на оценките на параметрите. Тъй като според (3.29) оценките по LS минимизират елементите на σ_e^2 , от (3.28) следва, че LS води и до

$$\min_{\Theta} \sigma_{\hat{\Theta}}^2, \quad (3.30)$$

или с други думи оценките са ефективни.

От анализа на оценките до момента следва, че ако e_k е ергодичен, центриран бял Гаусов шум, оценките са неизместени и ефективни, а LS оценителят е BLUE. Ако остатъкът не е ергодичен, $\Sigma_{E,i}$ остава диагонална и WLS, разгледан в следващата точка, е BLUE оценител. Това означава, че при подходящ избор на тегловната матрица оценките по WLS са ефективни. Ако $\text{cov}(e_{k_1}, e_{k_2}) \neq 0$ за $k_1 \neq k_2$, $\Sigma_{E,i}$ престава да е диагонална и BLUE оценител е GLS.

Както беше уточнено във Втора глава, приемането за нормално разпределение на e_k се извършва по-скоро по практически съображения. Често остатъкът е сума от различни случайни величини и тогава от централната гранична теорема следва, че разпределението му се доближава до Гаусовото. На практика съществуват и други разпределения, освен нормалното, при които $\hat{\Theta}$ са BLUE.

В заключение, има случаи, при които оценките по LS може да са неизместени, но да не са с минимална дисперсия или да са с минимална

дисперсия, но да са изместени, а също и да са неизместени и ефективни, при остатък, който не е Гаусов.

3.2.2.5 Примери

С долните примери се представят свойствата на оценките по LS, когато се нарушават някои от условията, при които те са BLUE. Също така се представя ефектът от стандартизацията на данните върху оценките на параметрите, както и улесняването в този случай на анализа на модела. В последния пример е засегнат въпросът за влиянието на мултиколинеарността в данните върху оценките.

Пример. *Изместеност на оценките при цветен остатък*

Изследването на ефекта от нарушаването на условието остатъкът да е бял шум, се извършва със симулация, а не с реални данни. Така оптималните стойности на параметрите, които се търсят с LS, са известни (зададат се в началото на симулацията). По този начин може за конкретни генерирани данни директно да се определи изместеността в оценяването на параметрите.

Сценарият на симулацията е следният. С помощта на ARX и ARMAX модели се симулира поведението на система, като се генерират входно-изходни данни. Нека всички полиноми на тези модели са от първи ред, а полиномните матрици са:

$$A(q^{-1}) = I + \begin{bmatrix} 0.3 & 0.2 \\ -0.5 & -0.8 \end{bmatrix} q^{-1}, \quad B(q^{-1}) = 0 + \begin{bmatrix} 0.8 & -0.9 \\ 0.9 & 0.6 \end{bmatrix} q^{-1},$$

$$C(q^{-1}) = I + \begin{bmatrix} 0.4 & -0.9 \\ 0.5 & -0.8 \end{bmatrix} q^{-1}.$$

Когато се изследва случаят с бял остатък, се приема, че $C(q^{-1}) = I$, т.е. моделът за генериране на данни е ARX и между стойностите на остатъчната грешка не се въвежда корелация. От друга страна, когато полиномите на $C(q^{-1})$ са от първи ред, остатъкът е цветен. Формират се $N = 300$ наблюдения, като първите 200 се използват за оценяване на параметри, а останалите 100 – за валидация на модела. Входният сигнал е бял шум с равномерно разпределение и стойности между 0 и

1 или накратко $u_{i,k} \sim \mathcal{U}(0,1)$ за $i = \overline{1,2}$. Изчисленото детерминирано поведение на системата y_d се зашумява със сигнал, който е центриран бял Гаусов шум, чието стандартно отклонение е 25% от това на y_d , или

$$e_k \sim \mathcal{N}(0, \frac{1}{16} \Sigma_{y_d}).$$

Моделът, чиито параметри се оценяват в примера, е ARX и отново полиномите му са от първи ред. Резултатът от прилагането на LS към двата генерирани набора от данни е представен в Таблицы 3.1 и 3.2. Векторът на параметрите θ , чиито стойности са в първата колона на таблиците, е със следната структура:

$$\theta = [(\text{vec } A_1^T)^T \quad (\text{vec } B_1^T)^T \quad (\text{vec } C_1^T)^T]^T.$$

С $\text{vec}(\cdot)$ е означено векторизирането на матрица. При това преобразуване стълбовете на матрицата се подреждат във вектор, т.е., ако $M \in \mathcal{R}^{m \times n}$, то

$$\text{vec } M = [M_{\cdot 1}^T \quad M_{\cdot 2}^T \quad \dots \quad M_{\cdot n}^T]^T \in \mathcal{R}^{mn}.$$

Таблица 3.1. Параметри, оценки, абсолютна и относителна грешка при остатък – бял шум

θ_i	$\hat{\theta}_i$	$\theta_i - \hat{\theta}_i$	$\frac{\theta_i - \hat{\theta}_i}{\theta_i} \times 100\%$
0.3	0.3093	-0.0093	-3.1013
0.2	0.1911	0.0089	4.4530
-0.5	-0.5376	0.0376	-7.5229
-0.8	-0.8196	0.0196	-2.4473
0.8	0.8236	-0.0236	-2.9472
-0.9	-0.9142	0.0142	-1.5747
0.9	0.8501	0.0499	5.5421
0.6	0.6463	-0.0463	-7.7107

Показателят VAF за двата изхода при бял и цветен остатък е съответно:

$$\text{VAF}_w = \begin{bmatrix} 94.06 \\ 92.40 \end{bmatrix} \quad \text{и} \quad \text{VAF}_c = \begin{bmatrix} 84.10 \\ 89.40 \end{bmatrix},$$

като стойностите са в проценти.

Както от таблиците, така и от стойностите на VAF се вижда, че когато остатъчната грешка е цветен шум, поради неподходящата структура

Таблица 3.2. Параметри, оценки, абсолютна и относителна грешка при цветен остатък

θ_i	$\hat{\theta}_i$	$\theta_i - \hat{\theta}_i$	$\frac{\theta_i - \hat{\theta}_i}{\hat{\theta}_i} \times 100\%$
0.3	0.2504	0.0496	16.5209
0.2	0.2439	-0.0439	-21.9457
-0.5	-0.5890	0.0890	-17.7946
-0.8	-0.7805	-0.0195	2.4391
0.8	0.8261	-0.0261	-3.2590
-0.9	-0.8625	-0.0375	4.1666
0.9	0.8488	0.0512	5.6873
0.6	0.6881	-0.0881	-14.6827
0.4	—	—	—
-0.9	—	—	—
0.5	—	—	—
-0.8	—	—	—

на ARX модела оценките са изместени и достоверността на модела се влошава. Изместеността е най-очевидна, когато е изразена с относителната грешка в проценти (последната колона в таблиците).

Един вариант за подобряване на модела е в него да се въведе формиращ филтър. Пример за това е ARMAX моделът, като тук филтърът е от тип МА. Параметрите на този разширен модел може да се оценят по метода на разширените най-малки квадрати (виж точка 3.2.5). Другият вариант е да се потърси ARX модел с по-подходяща структура. Този случай, когато e_k е цветен шум, е разгледан по-долу.

Ако се увеличат степените на полиномите на ARX модела, поради добавените степени на свобода може да се очаква, че качеството на модела ще се повиши. Наистина при нарастване на степените na , nb и nc от 1 до 10, както се вижда в Таблица 3.3, VAF първоначално нараства. Това показва, че когато за даден ARX модел остатъкът няма свойствата на бял шум, преди да се потърси модел от по-сложен тип, може да се потърси по-подходяща структура на ARX модела. Основната причина за този избор е, че параметрите на ARX модела се оценяват без използването на итеративни процедури, понеже този модел е линейно параметризиран. От друга страна, ако е въведен формиращ филтър, моделът вече не е линейно параметризиран (виж точка 3.2.5) и тогава параметрите се оценяват с итеративни методи.

Отново от Таблица 3.3 се вижда, че от $na = nb = nc = 3$, с нарастване на степените на полиномите, стойностите на VAF престават да нарастват. Причина за това е преоразмеряването на модела. Например

Таблица 3.3. VAF за двата изхода на ARX модели с различни степени на полиномите при цветен остатък

na, nb, nc	VAF _{c,1} %	VAF _{c,2} %
1	84.1023	89.4019
2	93.9599	93.4162
3	94.3147	93.5294
4	94.2673	93.4458
5	94.1914	93.4587
6	94.6706	93.1904
7	94.6170	93.1022
8	94.6812	92.6668
9	94.4583	92.0572
10	94.4554	92.4021

за последния модел, при който всички полиноми са от 10-и ред, броят на параметрите на всеки MISO ARX модел е $p_i = 40$, а наблюденията за определяне на $\hat{\Theta}$ са 200 и за изчисляване на VAF са 100. За сравнение, когато полиномите са от 3-и ред, $p_i = 12$. Ето защо при този вариант за подобряване на качеството на модела трябва да се внимава с добавянето на нови параметри. $\diamond$

Пример. *Кредитен риск: стандартизация и анализ на значимостта на факторите*

Целта на този пример е да се покаже ефектът от стандартизацията на данните върху оценките на LS, както и удобството, когато анализът се извършва със стандартизирания модел. За целта се използват данни, предоставени от компанията „Експириън“ [87]. В примера се използват 9 независими характеристики (фактори), описващи $N = 50283$ кандидати за кредит. Тези данни са отчетени в момента на кандидатстването и тъй като отразяват статиката в поведението на кандидатите, то и моделът ще е статичен. Част от факторите са предоставени от кандидатите, а друга част е закупена от кредитно бюро. За всеки кандидат е формирана зависима променлива y_k , която е вероятността кредитоискателят да е подходящ за кредит. Тук k е индексът на поредния клиент. Ако клиентът е одобрен от банката, то на базата на поведението му и на набор от предварително определени правила се изчислява y_k . Например, ако кредитополучателят е пропуснал три вноски в рамките на година, може да се приеме за лош платец и $y_k = 0$. В противен случай $y_k = 1$ (добър платец). От друга страна, ако клиентът е отхвърлен, т.е. не му е предоставен кредит, то за него са налични само независимите харак-

теристики. Тъй като поведението му не е наблюдавано, y_k не може да се определи директно. Затова се използва техника, с която се оценява каква би била вероятността кандидатът да е добър, ако банката му беше предоставила кредит. Тази дейност е извършена предварително, т.е. в използвания набор няма липсващи данни за y_k .

Тъй като системата е с един изход, то моделът е с вектор на параметрите. В Таблица 3.4 са представени резултатите от прилагането на LS. В първия случай са използвани данните без предварителна обработка (оценките са $\hat{\theta}$), а във втория случай факторите са стандартизирани (получените оценки са означени с $\hat{\theta}_s$). Моделът е със свободен член, т.е. първите елементи на $\hat{\theta}$ и $\hat{\theta}_s$ са съответно $\hat{\theta}_0$ и $\hat{\theta}_{s,0}$ (в таблицата тези оценки са означени като Intercept). Също така са добавени и колони, съдържащи стандартните грешки $\sigma_{\hat{\theta}}$ и $\sigma_{\hat{\theta}_s}$.

Таблица 3.4. Оценки на параметрите по LS при нестандартизирани и стандартизирани данни

фактор	$\hat{\theta}$	$\sigma_{\hat{\theta}}$	$\hat{\theta}_s$	$\sigma_{\hat{\theta}_s}$
Intercept	-0.0023	0.0001	0.6705	0.0027
AppScore	0.0014	0.0001	0.1525	0.0052
BurScore	-0.0006	0.0001	0.0319	0.0055
Age	1×10^{-7}	2×10^{-8}	0.0056	0.0029
GrAnInc	4×10^{-6}	3×10^{-7}	0.0116	0.0030
Loan	4×10^{-7}	2×10^{-7}	0.0115	0.0028
NPayms	-1×10^{-5}	6×10^{-6}	-0.0144	0.0030
SearchL6M	0.0005	0.0010	0.0046	0.0028
ERRef	-0.0713	0.0027	-0.0293	0.0029
AppDat	-1×10^{-5}	1×10^{-6}	-0.0357	0.0027

Стойностите на VAF са съответно:

$$\text{VAF} = 13.49\% \text{ и } \text{VAF}_s = 17.75\%.$$

Подобряването на достоверността на стандартизирания модел се дължи както на центрирането на факторите, така и на мащабирането им. От таблицата се вижда, че оценката $\hat{\theta}_4$, отговаряща на значимия фактор ‘годишен доход’, означен в таблицата с ‘GrAnInc’, е от порядъка на 10^{-6} . Причината за това не е слабо влияние на дохода на кандидатите върху y_k , а по-скоро големите стойности на ‘GrAnInc’. Средната стойност на този фактор е 14198.5, а стандартното отклонение е 9294.2. Понеже y_k е вероятност, то е очевидно, че параметърът $\hat{\theta}_4$ трябва да приема малки стойности, за да компенсира големите стойности на фактора. От стандартизираните оценки се вижда, че този фактор е практически

по-значим от възрастта на клиентите и броя на кандидатстванията им за кредит през последните шест месеца (съответно факторите ‘Age’ и ‘SearchL6M’). Така от стойностите на стандартизираните оценки лесно може да се определи значимостта на факторите.

От стандартните грешки се вижда, че когато факторите са стандартизирани, точността в оценяването на параметрите е сходна, за разлика от първия случай, при който $\hat{\theta}_i$, свързани с факторите с по-слаба вариация (на практика изменящи се в по-малък диапазон), са по-неточни. Основната причина е мащабирането на факторите, което е част от стандартизацията. $\diamond$

Пример. *Кредитен риск: мултиколинеарност и логични трендове*

Нека към данните, използвани в предишния пример се добави като фактор и доходът през последните 6 месеца, означен с ‘Gr6MInc’. Този фактор съдържа почти идентична информация с тази в ‘GrAnInc’ отразяващ дохода за една година. Това означава, че с използване на новия набор в модела ще участват два фактора, единият от които е излишен, тъй като не допринася значимо за описание на вариацията на изхода. Такива случаи невинаги са очевидни особено когато се работи с много фактори. Например възможно е близка зависимост до линейната да се наблюдава между слабо корелирани фактори (такава възможност се дискутира в последната тема на точка 2.2.3.3).

Таблица 3.5. Оценки на параметрите LS при нестандартизирани и стандартизирани данни

фактор	$\hat{\theta}_s$	$\sigma_{\hat{\theta}_s}$
Intercept	0.6705	0.0027
AppScore	0.1525	0.0052
BurScore	0.0319	0.0055
Age	0.0056	0.0029
GrAnInc	0.2659	0.2287
Gr6MInc	-0.2544	0.2287
Loan	0.0115	0.0028
NPayms	-0.0144	0.0030
SearchL6M	0.0046	0.0028
ERRef	-0.0293	0.0029
AppDat	-0.0357	0.0027

Оценките на параметрите и стандартните им грешки са дадени в Таблица 3.5. Вижда се, че ‘Gr6MInc’ влияе обратно пропорционално на изхода, откъдето може да се направи грешният извод, че колкото

по-голям е шестмесечният доход на кандидат, толкова по-малка е вероятността той да е добър платец. Освен това стандартната грешка, отразяваща точността в определянето на оценките, е с два порядъка по-голяма за $\hat{\theta}_{s,i}$, отговарящи на мултиколинеарните фактори, от $\sigma_{\hat{\theta}_s}$ на останалите оценки.

Мултиколинеарността от гледна точка на числените проблеми е разгледана в точка 3.4, а от гледна точка на значимостта на факторите – в точка 3.5.7. $\diamond$

3.2.3 Метод на претеглените най-малки квадрати

Едно обобщение на LS може да се получи чрез отчитане на стойностите на остатъка с различно тегло в оптимизационната процедура. Това е идеята, заложена в метода на претеглените най-малки квадрати (WLS – Weighted Least Squares) [146]. Свойствата на оценките по WLS зависят от избора на теглата, с които се отчитат квадратите на остатъците в целевата функция.

При изложението на метода е удобно да се използват тензори и поради това те са разгледани по-долу. Така решението на WLS за многомерния случай остава компактно като запис [77, 76]. Друга важна предпоставка за употребата на тензори е, че процесорите на видеокартите, чието приложение беше споменато в началото на главата, са хардуерно оптимизирани за работа с тензори до трети ред (за целта се използват масиви с един, два или три индекса). Причината за това е, че тези процесори са разработени за визуализация на движение от тримерното пространство. Когато се реализират алгоритми за оценяване на параметри на нелинейни модели, понякога по естествен път се достига до тензори от по-висок ред. Например един от най-разпространените методи за обучение на невронни мрежи с обратно разпространение на грешката, реализиран с метода за оптимизация Нютон-Рафсън, се описва с тензори от пети ред, като изкуствено редът може да се сведе до трети или втори ред. Причината е, че при диференциране два пъти на векторна функция по матрица, съдържаща параметрите от даден слой на мрежата, се получава тензор от пети ред. Оценителят значително се опростява, ако се използва методът GN, при който се извършва еднократно диференциране, и резултатът е тензор от трети ред.

При реализацията на WLS се достига именно до тензори от трети ред, което е предпоставка за получаването на високоефективни оцени-

тели, базирани на този метод. По-долу е представен WLS, като първоначалното преобразуване на модела в общ вид с матрица на параметрите е извършено в точка 3.2.1.

3.2.3.1 Показател на качеството

В приложението към монографията са дефинирани действията с тензори, които по-долу се извършват. Когато моделът е с повече от един изход, то и остатъците в даден момент са повече от един. Затова е необходимо да се въведе отделна тегловна матрица W_i за всеки остатък, свързан със съответен изход. В постановката на WLS целевата функция не включва взаимните произведения $e_{i,k_1}e_{i,k_2}$ за $k_1 \neq k_2$, а също и $e_{i,k}e_{j,k}$ за $i \neq j$ и задачата за оценяване се опростява, тъй като матриците W_i за $i = \overline{1, \ell}$ са диагонални. За коректната работа на оценителя диагоналните елементи на W_i трябва да са неотрицателни, тъй като имат смисъл на тегла, отразяващи степента на влияние на $e_{i,k}^2$ върху целевата функция.

Нека тензорът от трети ред (порядък) $W \in \mathcal{R}^{N-n \times N-n \times \ell}$, даден на Фигура 3.2, съдържа ℓ -те на брой тегловни матрици, като $W_{..i} = W_i$. В такъв случай, по аналогия със стандартния LS, но с въведени тегла на квадратичните стойности на остатъците, целевата функция може да се представи компактно като:

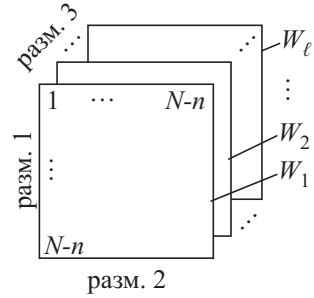
$$\mathcal{F}(\Theta) = \text{sum}(W \bar{\circ}_1^3 E \bar{\circ}_2^3 E). \quad (3.31)$$

За онагледяване на действията, които се извършват в (3.31), те са по-

$$E \quad W \quad E \quad = \quad T_1 = W \bar{\circ}_1^3 E \quad E \quad = \quad T_2 = W \bar{\circ}_1^3 E \bar{\circ}_2^3 E$$

Фигура 3.3. Формиране на стойността на целевата функция

казани графично на Фигура 3.3. За опростяване на изразите е въведено т.нар. свиващо повекторно умножение, означено с ' $\bar{\circ}$ '. То е въведе-



Фигура 3.2. Тегловният тензор W в WLS

но в монографията с цел опростяване на записите и е дефинирано в Приложението. Първото умножение е между тензора W и матрицата $E \in \mathcal{R}^{N-n \times \ell}$. Резултатът е тензорът $T_1 = W \circ_1^3 E \in \mathcal{R}^{1 \times N-n \times \ell}$. Както се вижда от размерността на T_1 , свиването е извършено по първата размерност. При следващото, отново свиващо, повекторно умножение, но този път по втората размерност, резултатът е $T_2 = T_1 \circ_2^3 E \in \mathcal{R}^{1 \times 1 \times \ell}$.

Този тензор може да се разглежда като ℓ мерен вектор, елементите на който са сумите от претеглените квадрати на остатъците по съответните изходи. С други думи $\mathcal{F}(\Theta) = \text{sum } T_2$ е претеглената сума от квадратите на остатъците по всеки изход на MIMO модела.

3.2.3.2 Оценки по WLS

Отново, както при LS, за да се намерят оптималните параметри, е необходимо да се определят производните на $\mathcal{F}(\Theta)$. При диференциране на $\mathcal{F}(\Theta)$ по матрицата $\Theta \in \mathcal{R}^{z \times \ell}$ се получава градиентната матрица $G \in \mathcal{R}^{z \times \ell}$. С използване на матрицата на данните Φ , матрицата на остатъка E и тегловния тензор W , както и като се отчете, че остатъкът е линейна функция на параметрите, G може да се запише съкратено по следния начин:

$$G = 2W \circ_1 \Phi \circ_2^3 E. \quad (3.32)$$

Действието ‘ $\circ$ ’ е стандартното умножение на тензор с матрица (виж Приложението). За да се определят оптималните оценки на параметрите, нека изразът за G да се развие по следния начин:

$$\begin{aligned} G &= 2W \circ_1 \Phi \circ_2^3 (Y - \Phi \hat{\Theta}) \\ &= 2W \circ_1 \Phi \circ_2^3 Y - 2W \circ_1 \Phi \circ_2 \Phi \circ_2^3 \hat{\Theta} \\ &= 2W \circ_1 \Phi \circ_2^3 Y - 2F \circ_2^3 \hat{\Theta}. \end{aligned}$$

В последния израз е въведен тензорът $F = W \circ_1 \Phi \circ_2 \Phi$, който съдържа матриците на Фишер $F_{..i} = \Phi^T W_i \Phi$ за $i = 1, \ell$. Нека обратните на тези матрици формират тензора P , т.е. $P_{..i} = F_{..i}^{-1}$. Това са ковариационните матрици на грешките в оценяването на параметрите, съдържащи се в стълбовете на Θ .

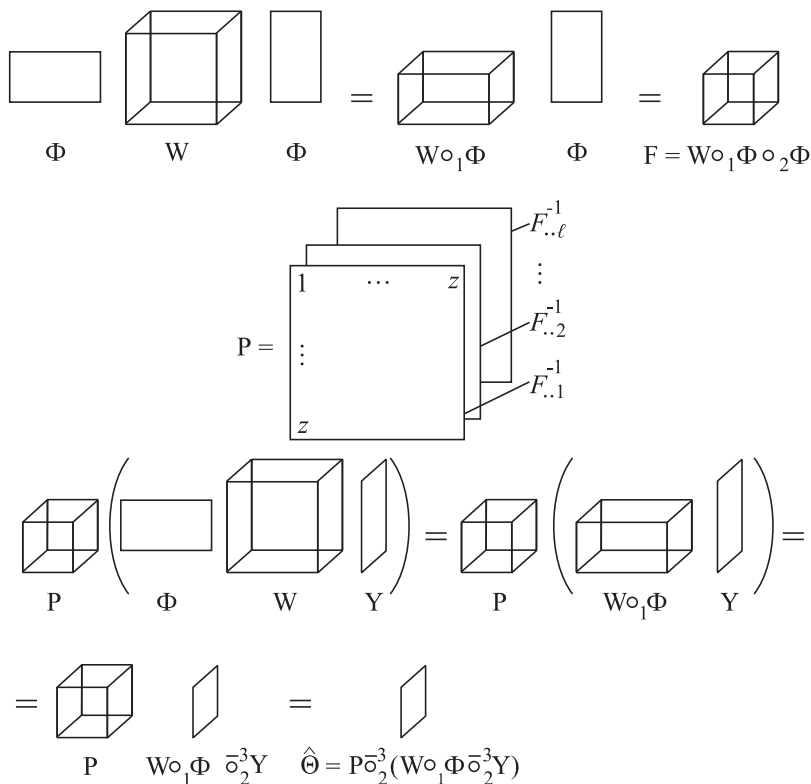
Търсените оптимални оценки на параметрите са тези, при които градиентната матрица се нулира, т.е.

$$G|_{\Theta=\hat{\Theta}} = 2W \circ_1 \Phi \circ_2^3 Y - 2F \circ_2^3 \hat{\Theta} = 0.$$

Така за оценките, минимизиращи $\mathcal{F}(\Theta)$, се получава:

$$\hat{\Theta} = P \bar{\sigma}_2^3 (W \circ_1 \Phi \bar{\sigma}_2^3 Y). \quad (3.33)$$

Действията, които се извършват в (3.33), са представени визуално на Фигура 3.4.



Фигура 3.4. Формиране на оценките по WLS

3.2.3.3 Приложения на метода

Ефективни оценки при неергодичен остатък

Когато остатъкът е ергодичен бял шум и нека е с Гаусово разпределение, LS е BLUE оценител. При нарушаване на изискването за ергодичност е възможно WLS да осигури ефективни оценки. За целта нека тегловният тензор да съдържа обратните на ковариационните матри-

ци $\Sigma_{E..i} = M\{E_i E_i^T\}$, за $i = \overline{1, \ell}$. Те не се свеждат до произведение на скалар и единична матрица, тъй като поради неергодичността на e_k $\sigma_{e,i,k}^2$ се изменя в рамките на интервала на наблюдение. Такъв случай е разгледан в точка 2.3.4.4. Въпреки това корелационната матрица на остатъка е диагонална, при положение че e_k е бял шум. За този случай нека $\sigma_{e,i}^2 = \text{diag } \Sigma_{E..i}$ е векторът, съдържащ стойностите на дисперсията на i -тия остатък за интервала на наблюдение.

Често $\sigma_{e,i}^2$ не са предварително известни и тогава информацията за тях се извлича от наличните данни. Съществуват итеративни алгоритми, при които параметрите на модела многократно се оценяват с WLS. При този вариант, на всяка итерация, за формиране на W се използва оценка на ковариационната матрица на остатъка, определен от предишната итерация.

Нестационарни системи

Друго важно приложение на WLS е, когато системата е нестационарна. При този вариант старите данни съдържат неактуална информация за моментното поведение на изследваната система. В този случай всяка от тегловните матрици се избира така, че елементите ѝ по диагонала да намаляват с отдалечаване на наблюденията от последния момент, за който има данни. Ако теглата намаляват експоненциално, се получава WLS с т.нар. експоненциален фактор на забравяне. При така формиран тегловен тензор оценките на параметрите са най-чувствителни към най-новите данни.

Нека вместо два индекса на диагоналните елементи на всяка от матриците W_i се използва един индекс, отговарящ на момента от време, в който е снето наблюдението, т.е. j -тият диагонален елемент $w_{i,jj}$ за $j = \overline{1, N-n}$ се замести с $w_{i,k}$, за $k = \overline{n+1, N}$. С други думи i -тата матрица в тегловния тензор е

$$W_{..i} = W_i = \text{diag}[w_{i,n+1} \ w_{i,n+2} \ \dots \ w_{i,N}].$$

Теглата при експоненциалния фактор на забравяне се формират по следния начин:

$$w_{i,k} = \lambda_i^{N-k} \text{ за } k = \overline{n+1, N} \text{ и } \lambda \in [0, 1].$$

В последния дискретен момент теглото, свързано с i -тия изход, е $w_{i,N} = 1$, а с намаляване на k $w_{i,k}$ също намалява. От друга страна, ако $\lambda = 1$, всички тегла са равни на единица и WLS се свежда до LS.

Редуциран набор с претеглени наблюдения

Това приложение на WLS обикновено се използва, когато първоначалните данни са твърде много и е необходимо те да се редуцират (виж точка 2.2.1.4). Такива случаи са типични за финансите, социологията, медицината, а понякога за пазарните системи. С нарастване на броя на наблюденията от един момент подобрението на достоверността на модела става нечувствително към нововъведените наблюдения.

Причината за редуцирането – дейност от предварителната обработка на данните, е да се намали времето за идентификация, като не се влоши чувствително качеството на крайния модел. Ако данните са от порядъка на много мега- или гигабайти, процесът на оценяване може да отнеме твърде много време, което би могло да се използва например за анализ на данните. В същата точка от предишната глава е обяснена причината за въвеждането на тегла, асоциирани с наблюденията в намаления набор. В случая диагоналните тегловни матрици съдържат споменатите тегла.

Пример. Социология – редуциране на данни и въвеждане на тегла

Нека даден изход на социална система да отразява удовлетвореността от живота на индивидите от конкретен географски район и нека той да приема стойностите 0 (неудовлетворен) и 1 (удовлетворен). Нека също 80% от изследваните индивиди са удовлетворени, а целите на изследването са насочени към изолиране на факторите, оказващи негативно влияние върху удовлетвореността. Тогава е желателно данните за всички неудовлетворени хора да участват в набора, а редуцирането да се извърши по отношение на данните за удовлетворените. Нека 50% от удовлетворените се премахнат. Тогава теглата на неудовлетворените индивиди са равни на 1, а на удовлетворените – на 2 (понеже на всеки удовлетворен индивид отговарят двама от първоначалния набор). ◇

Липсващи и/или нехарактерни стойности

Ако първоначалният набор от данни съдържа нехарактерни или липсващи стойности, е необходимо данните да се „почистят“ (виж точки 2.2.2.2 и 2.2.2.3). Обикновено при малък брой наблюдения на мястото на липсващите и на премахнатите нехарактерни стойности се въвеждат подходящи стойности. Това е необходимо за осъществяване на процеса на оценяване на модела и неговата валидация. Тази дейност е характерна за многомерните системи, където не е целесъобразно да се пренебрегват данните за всички величини в даден момент, ако някой сигнал

е с нехарактерна/липсваща стойност. Независимо как се формират въведените в набора стойности, те са изкуствено създадени на базата на останалите данни (не са получени директно от експеримент). Затова е коректно наблюдения, съдържащи такива стойности, да участват с по-малко тегло, когато моделът се изгражда/валидира.

В [17] е предложен начин за формиране на теглата в зависимост от процентното съотношение на изкуствено въведените стойности в наблюденията. С нарастване на броя на тези стойности в даден ред от матрицата на регресорите теглото, с което този ред влияе на оценките, намалява. Най-общо теглата $w_{i,k}$ се формират по следния начин:

$$w_{i,k} \begin{cases} = 1, & \text{когато } y_{i,k} \text{ и } \varphi_{i,k} \text{ са налични,} \\ \in (0, 1), & \text{когато } y_{i,k} \text{ и част от } \varphi_{i,k} \text{ са налични,} \\ 0, & \text{когато } y_{i,k} \text{ и/или } \varphi_{i,k} \text{ не са налични,} \end{cases}$$

като $\varphi_{i,k}$ е векторът на регресорите, описващи i -тия изход.

Пример. *Пазарна система: липсващи стойности в набора от данни*

За да се изследва влиянието на непълни входно-изходни наблюдения върху точността на оценките, се използват данни от търговската верига „Доминик“ (Dominick's Finer Foods). Тази верига от супермаркети е предоставила значителна база данни на Чикагския университет и може да бъде достъпена от [70].

Пълният набор съдържа данни за 29 категории от продукти в магазините на веригата. Данните отразяват седмичните наблюдения на продажбите, цените на дребно, намаленията, както и промоционалните дейности (вид на промоцията, вид на рекламата и начин на излагане на продукта в магазина). Отделните продукти се разпознават с уникален код на продукта (UPC – Unit Product Code). В примера е използван поднабор от данни от конкретен магазин за категорията "Зърнени храни". Извадката обхваща почти четиригодишен период ($N = 199$ седмици).

При подготвянето на набора, подходящ за идентификация, за всяка от седмиците се формира вектор на наблюдение, съдържащ всички входно-изходни данни за всеки продукт. Някои продукти не са присъствали на пазара в началото на изследвания период, други са премахнати преди края, а също така има случаи, когато някои продукти са били изчерпвани. Така в набора се появяват значимо количество липсващи стойности (за конкретните данни над 16%). За да може данните да се използват, е необходимо липсващите стойности да се попълнят по подходящ начин. В примера са определени средните стойности на входовете и изходите и те се въвеждат на мястото на липсващите.

В примера са построени два модела – единият с използване на LS, а другият с WLS. $\frac{2}{3}N$ от данните се използват за определяне на параметрите и структурата на модела, а с останалата $\frac{1}{3}N$ са определени стойностите на VAF.

Системата е декомпозирана на MISO подсистеми, за които се определят модели. Причината да не се използва MIMO модел, е, че наборът не е почистен от продукти, за които данните са недостатъчни. Структурата на описанията е избрана с използване на стъпков метод като този, описан в 3.5.7. Максималните степени на полиномите са $na = 1$ и $nb = 2$.

Таблица 3.6. Стойности на VAF, получени по LS и по WLS, с отчитане на липсващите стойности

UPC	VAF _{LS} %	VAF _{WLS} %
1600062680	22.2508	23.8177
1600062750	70.5849	71.6071
1600065700	68.8959	69.2772
1600065720	30.1016	30.4272
1600065850	36.4797	38.4380
1600065890	57.9704	59.8511
1600065960	47.3226	47.0468
1600066310	45.6047	47.7463
1600067440	58.7668	61.3108
3000002580	24.6647	25.7683
3000006430	63.6439	64.7744
3000006500	7.5587	7.6201
3000006560	17.3699	19.6529
3800000535	44.7967	45.2638
3800001011	34.4605	35.0779
3800001210	12.5017	13.1122
3800001912	37.8300	40.6329
3800003700	14.5637	14.6655
3800004900	62.6096	63.1103
3800013900	48.6490	49.4104
4300010521	38.1991	38.1616
4300011171	15.6990	23.7164
4300011650	23.9698	24.3698
4300011760	15.1110	17.4288

В WLS теглата са формирани така, че да намаляват с нарастване на липсващите стойности в дадено наблюдение, и са нула, когато липсва стойност за съответния изход. Освен попълването на липсващите стойности данните са стандартизирани. Също така от набора са премахнати мултиколинеарните фактори. Формирането на допълнителни промен-

ливи, както и намирането на по-подходящи стойности за na и nb , не е извършвано. Въпреки това от резултатите, представени в Таблица 3.6, се вижда тенденцията на увеличаване на VAF, когато се използват тегла, зависещи от информацията за броя на липсващите стойности във всяко наблюдение, както за това, кой сигнал не е наличен. $\diamond$

Комбинирано приложение на WLS

Възможно е горните приложения да се комбинират, например, ако W'_1 е тензор, формиран с цел отчитане на нестационарност, а с W''_2 се отразява наличието на липсващи и нехарактерни стойности в данните, то резултантният тегловен тензор, който участва в WLS, е

$$W = W'_1 * W''_2,$$

като символът “*” е поелементно умножение, в случая на тензори.

Тегловният тензор (или тензори) в WLS е разреден, т.е. съдържа голям брой нулеви стойности, и затова е подходящо той да се представя в паметта именно като разреден (sparse). Така не се заделя памет за нулевите елементи и съответно не се извършват действия с тях.

3.2.4 Метод на обобщените най-малки квадрати

В WLS се приема, че тегловните матрици W_i за $i = \overline{1, \ell}$ са диагонални. Но това приемане се нарушава, когато остатъкът е цветен шум. Оцветеността в e_k може да се дължи на нестационарно поведение на системата. Друг вариант е недостатъчен брой фактори за описание на изхода, като липсващите значими величини се причисляват към смущенията от околната среда. Поради това, че влиянието им не може да бъде описано от модела, то се добавя в остатъка. В резултат на това стойностите на e_k в различни моменти от време са корелирани, което води до недиагонални ковариационни матрици $\Sigma_{E,i}$. Оценителят, който осигурява неизместени и ефективни оценки, в този случай е GLS.

3.2.4.1 Показател на качеството

Когато тегловните матрици се зададат като $W_i = \Sigma_{E,i}^{-1}$, те не са диагонални, а това води в целевата функция до появата на (претеглени) взаимни произведения от вида $e_{i,k_1} e_{i,k_2}$ за $k_1 \neq k_2$. Въпреки това не се променя видът на целевата функция, нито настъпва изменение в израза за оптималното решение. Целевата функция е

$$\mathcal{F}(\Theta) = \text{sum}(W \circ_1^3 E \circ_2^3 E).$$

При първото повекторно умножение, което е между тензора W и матрицата $E \in \mathcal{R}^{N-n \times \ell}$, резултатът е тензорът $T_1 = W \circ_1^3 E \in \mathcal{R}^{1 \times N-n \times \ell}$. При следващото, повекторно умножение резултатът е $T_2 = T_1 \circ_2^3 E \in \mathcal{R}^{1 \times 1 \times \ell}$. Този тензор може да се разглежда като ℓ мерен вектор, елементите на който са сумите от претеглените произведения на остатъците по съответните изходи.

3.2.4.2 Оценки по GLS

Извеждането на оптималните оценки се извършва по същите стъпки, както при WLS, а оценките, минимизиращи $\mathcal{F}(\Theta)$, са:

$$\hat{\Theta} = P \circ_2^3 (W \circ_1 \Phi \circ_2^3 Y).$$

Тензорът P се състои от ковариационните матрици $P_{..i} = F_{..i}^{-1}$, а $F = W \circ_1 \Phi \circ_2 \Phi$ съдържа матриците на Фишер $F_{..i} = \Phi^T W_i \Phi$ за $i = \overline{1, \ell}$. Единствената разлика с WLS е структурата на тензора W .

Обикновено матриците $\Sigma_{E,i}$ не са известни. Затова GLS се реализира итеративно, като с всяка итерация $\text{cov}(e_{i,k_1}, e_{i,k_2})$ се описва все по-точно. Тъй като стойностите на e_k в различни моменти от време са корелирани, е подходящо да се използва AR модел по отношение на остатъка и с него да се опише тази зависимост [65, 56, 150]. Друг вариант е да се използва MA филтър, на входа на който постъпва фиктивен бял шум, а изходът е оцветеният остатък. Регресионният модел престава да е линейно параметризиран, понеже стойностите на остатъка, които зависят от модела, също участват в описанието като фактори. Въпреки това е възможно многократно да се приложи LS, като се приеме, че моделът е линеен по параметри. С нарастване на итерациите оценките стават неизместени. Методът, реализиращ тази идея, когато филтърът е MA, е ELS. Той е разгледан по-долу.

3.2.5 Метод на разширените най-малки квадрати

Когато се търси линеен модел, тъй като наличните данни са входно-изходните наблюдения, естествено е първо да се търси директната връзка между тях. Ако системата е динамична, то в модела се включва и предисторията им. С други думи удачно е първоначално LS да се приложи към ARX модел. Ако не се намери подходяща структура на модела, част от детерминираното поведение на системата остава неотразено от модела. Такъв е случаят, когато не всички значими величини за описанието на изхода са налични (те са причислени към влиянието

на околната среда). Друг вариант е, тъй като моделът е апроксимация на системата, някои аспекти от нейното поведение, като значими производни, нелинейности и др., да са пренебрегнати (това е несъответствие между структурата на модела и системата).

Неотчетеният от модела детерминиран аспект от поведението на системата се причислява към остатъка. По тази причина e_k няма свойствата на бял шум, което води до изместеност на оценките. В този случай цветният остатък ще се означава с $\tilde{e}_k$, а ARX моделът ще се записва като:

$$A(q^{-1})y_k = B(q^{-1})u_k + \tilde{e}_k. \quad (3.34)$$

Един от начините за осигуряване на неизместени оценки е в модела да се въведе формиращ филтър, с който да се опише оцветеността в остатъчната грешка. Този филтър се определя така, че на входа му да постъпва (фиктивен) случаен сигнал e_k , който е бял шум, а изходът на филтъра да е цветният остатък $\tilde{e}_k$. Така задачата за оценяване се свежда до търсене на параметрите на разширен (с формиращия филтър) модел, чиято структура позволява остатъчната грешка да се избели.

Независимо от вида на филтъра параметрите му също се оценяват. В тази подточка е разгледан методът на разширените най-малки квадрати (ELS), който се прилага за модели от тип ARMAX.

3.2.5.1 Описание на системата

Един от най-разпространените в практиката варианти за осигуряване на неизместени оценки е филтърът да бъде от тип пълзящо средно (MA), т.е. да е от вида:

$$\tilde{e}_k = C(q^{-1})e_k. \quad (3.35)$$

Разширеният ARX модел с MA филтър е всъщност ARMAX моделът:

$$A(q^{-1})y_k = B(q^{-1})u_k + C(q^{-1})e_k. \quad (3.36)$$

С представянето на полиномните матрици като полиноми по степените на q^{-1} , (3.36), може да се запише като:

$$\begin{aligned} y_k + A_1 y_{k-1} + \dots + A_{na} y_{k-na} &= B_1 u_{k-1} + \dots + B_{nb} u_{k-nb} \\ &+ C_1 e_{k-1} + \dots + C_{nc} e_{k-nc} + e_k. \end{aligned}$$

Отново нека параметрите се групират в матрица, а регресорите във вектор, както това беше направено в точка 3.2.1 за ARX модел. След изразяване на y_k от последния израз и отделяне на параметрите от рег-

ресорите за общата форма на модела в k -тия момент се получава:

$$y_k = \Theta^T \varphi_k + e_k, \quad (3.37)$$

където $\varphi_k \in \mathcal{R}^z$ и $\Theta \in \mathcal{R}^{z \times \ell}$ са:

$$\begin{aligned} \varphi_k &= [-y_{k-1}^T \ \cdots \ -y_{k-na}^T \ u_{k-1}^T \ \cdots \ u_{k-nb}^T \mid e_{k-1}^T \ \cdots \ e_{k-nc}^T]^T \\ &= [\varphi_{\text{ARX},k}^T \ \varphi_{\text{f},k}^T]^T, \end{aligned} \quad (3.38)$$

$$\Theta = [A_1 \ \cdots \ A_{na} \ B_1 \ \cdots \ B_{nb} \mid C_1 \ \cdots \ C_{nc}]^T = [\Theta_{\text{ARX}}^T \ \Theta_{\text{f}}^T]^T, \quad (3.39)$$

където броят на регресорите е $z = \ell na + m nb + \ell nc$. С разделянето на φ_k и Θ в (3.38) и (3.39) на две части се цели отделянето на ARX частта на разширения модел от тази на МА филтъра. Общото представяне на модела за интервала на наблюдение е

$$Y = \Phi \Theta + E. \quad (3.40)$$

Матриците Y , E и Φ се дефинират по същия начин, както в точка 3.2.1 (виж (3.6), (3.8) и (3.14)). Според разделянето на ARMAX модела на две части матрицата на данните Φ може да се представи като:

$$\Phi = [\Phi_{\text{ARX}} \ \Phi_{\text{f}}]. \quad (3.41)$$

Φ_{ARX} съдържа входно-изходните регресори, а Φ_{f} зависи от стойностите на неизвестния остатък на ARMAX модела, който трябва бъде бял шум.

Принципната разлика между ARX и ARMAX моделите е, че при последните φ_k зависи освен от наличните входно-изходни величини и от неизвестната величина e_k . Това означава, че Φ_{f} в (3.41) не може да се формира преди получаването на модел, и по тази причина решението за оптималната матрица на параметрите $\hat{\Theta}$ по LS не е приложимо за ARMAX модели. Наличието на e_k в матрицата на данните означава, че този сигнал също трябва да се оценява. В такъв случай, оценката на остатъка зависи както от източника на информация – входно-изходните данни, така и от конкретните оценки на параметрите и структурата на модела. Поради наличието на e_k във вектора на регресорите, изходът на ARMAX моделът не е линеен по отношение на параметрите.

За да се представи тази нелинейна зависимост, нека изходът на ARMAX модела $\hat{y}_k$ се запише като:

$$\hat{y}_k = \Theta^T \varphi_k = \Theta^T [\varphi_{\text{ARX},k}^T \ \varphi_{\text{f},k}^T]^T.$$

Част от елементите на φ_k са регресорите на филтъра (елементите на вектора $\varphi_{\text{f},k}$), които са предишни стойности на e_k . Но остатъчната грешка

зависи от параметрите, по-точно

$$e_k = y_k - \Theta^T \varphi_k.$$

Това означава, че когато моделът е ARMAX, $\varphi_k = \varphi_k(\Theta)$, откъдето следва, че изходът

$$\hat{y}_k = \Theta^T \varphi_k(\Theta)$$

е нелинейна функция на параметрите.

3.2.5.2 Показател на качеството

Видът на показателя на качеството и на критерия в ELS са като при LS, т.е. $\mathcal{F}(\Theta) = \text{tr}(E^T E)$, а критерият е

$$\min_{\Theta} \mathcal{F}(\Theta) = \min_{\Theta} \text{tr}(E^T E). \quad (3.42)$$

В началото на Трета глава беше споменато, че в практиката са се установили два подхода за оценяване на параметрите на модели като ARMAX. Тъй като при това описание $\hat{y}_k$ зависи нелинейно от параметрите, единият подход за минимизация на $\mathcal{F}(\Theta)$ е да се приложи метод на нелинейните най-малките квадрати (NLS – Non-linear Least Squares). Това води да необходимостта от използване на метод за оптимизация, който решава задачата (3.42).

Другият вариант е да се приложи описаният по-долу метод на разширените най-малки квадрати ELS. При него многократно се прилага стандартният LS и се изчислява остатъкът на ARMAX модела. Използването на LS предполага приемането, че изходът на модела е линеен по отношение на параметрите. Получените оценки, които в началото са изместени, се използват за формиране на сигнал, който е приближение на e_k . Това приближение се използва, за да се определи нова матрица на данните, към която отново се прилага LS. Така итеративно се получават параметри на модела, които водят до все по-точна оценка на e_k , а оттам и до по-малка изместеност на оценките на параметрите. С нарастване на итерациите, колкото по-удачно е избрана структурата на модела, толкова свойствата на остатъка стават по-близки до тези на белия шум.

Разликата между NLS (реализиран с метод за оптимизация) и итеративната процедура на ELS е следната. В първия случай на базата на информацията за целевата функция (стойността ѝ и/или наклона и/или кривината ѝ) в текущата точка от пространството на параметрите се избира посока на придвижване към оптимума. От друга страна, в

описания по-долу алгоритъм придвижването на оценките на параметрите към оптималните им стойности става без експлицитното използване на информацията за повърхнината на $\mathcal{F}(\Theta)$.

Въпреки че в ELS се приема, че нелинейният по параметри модел е линейно параметризиран, често ELS води до по-бързо определяне на оценките на параметрите от NLS.

3.2.5.3 Алгоритъм на ELS

Алгоритъмът на ELS за оценка на ARMAX модели се състои от описаните по-долу стъпки.

1. Инициализация

Задават се параметрите, които определят условията за спиране на итеративната процедура, например толеранси като τ_Θ (праг на промяна на параметрите), $\tau_{\mathcal{F}}$ (праг на промяна на целевата функция) и др., а също и максималният брой итерации i_{max} . Формира се матрицата Φ_{ARX} , отговаряща на структурата на ARX модела (съдържа само входно-изходните наблюдения). Параметрите на ARX модела се оценяват като:

$$\Theta_{ARX}^{(0)} = (\Phi_{ARX}^T \Phi_{ARX})^{-1} \Phi_{ARX}^T Y.$$

$\hat{\Theta}_{ARX}^{(0)}$ е изместена оценка поради цветния характер на остатъка $\tilde{e}_k$. Формира се матрицата E , съдържаща стойностите на оценения (неизвестен) остатък e_k , който трябва да е БШ. В инициализиращата стъпка E съвпада с матрицата $\tilde{E}$, зависеща от цветните остатъци, т.е.

$$E^{(0)} \equiv \tilde{E}^{(0)} = Y - \Phi_{ARX} \Theta_{ARX}^{(0)}. \quad (3.43)$$

2. LS за ARMAX модела

В i -тата итерация се формира вторият блок $\Phi_f^{(i)}$ в (3.41) и матрицата $\Phi^{(i)}$. Матрицата $\Phi_f^{(i)}$ съответства на МА модела, а пълната матрица на данните (отговаряща на ARMAX модела) в i -тата итерация е

$$\Phi^{(i)} = [\Phi_{ARX} \quad \Phi_f^{(i)}].$$

Φ_{ARX} съдържа наличните входно-изходни данни и не се актуализира. Затова в последното равенство означението Φ_{ARX} няма индекс на текущата итерация. От друга страна, $\Phi_f^{(i)}$ зависи от елементите на $E^{(i)}$, които се произчисляват на всяка итерация. Параметрите

на ARMAX модела в текущата итерация се формират като:

$$\Theta^{(i)} = (\Phi^{(i)T} \Phi^{(i)})^{-1} \Phi^{(i)T} Y.$$

3. *Оценяване на остатъка e_k*

Матрицата E в текущата итерация се изчислява като:

$$E^{(i)} = Y - \Phi^{(i)} \Theta^{(i)}. \quad (3.44)$$

4. *Проверка за спиране*

Избраното условие за спиране на итеративната процедура може да включва абсолютния критерий

$$\max(|\Theta^{(i)} - \Theta^{(i-1)}|) < \tau_{\Theta},$$

както и допълнителното условие $i > i_{max}$, в случай че в разумен брой итерации не се намери подходящо решение. Друг подходящ критерий, свързан с изменението на параметрите, е $\|\Theta^{(i)} - \Theta^{(i-1)}\|_F < \tau_{\Theta}$. Въпреки че изчисляването на целевата функция не е необходимо за изпълнението на итеративния алгоритъм, подходящ критерий е

$$|\mathcal{F}_{\Theta}^{(i)} - \mathcal{F}_{\Theta}^{(i-1)}| < \tau_{\mathcal{F}}$$

(за опростяване на записа, когато се описват итеративните алгоритми, се използва означението $\mathcal{F}^{(i)} = \mathcal{F}(\Theta^{(i)})$).

Ако някое от избраните условия е изпълнено, процедурата за оценяване на параметрите на ARMAX модела се преустановява. В противен случай се продължава от стъпка 2.

3.2.5.4 Особености на метода

От (3.43) се вижда, че при стартиране на процедурата $\Theta_f^{(0)} = 0$ и тогава $C(q^{-1}) = I$. Поради неоптималните стойности на параметрите на филтъра $C(q^{-1})$ в началото на процедурата сигналът $e_k^{(0)} = \tilde{e}_k$ е цветен шум.

Очаква се с нарастване на итерациите, елементите на матрицата $\Theta_f^{(i)}$ да сходят към оптималните параметри. Това води до доближаване свойствата на e_k до тези на белия шум и съответно до намаляване на изместеността на оценките $\Theta_{ARX}^{(i)}$.

Матрицата $\Theta_f^{(i)}$ зависи най-вече от e_k , който се оценява, а $\Theta_{ARX}^{(i)}$ от наличните входно-изходни данни. По тази причина оценката $\Theta_f^{(i)}$ сходя чувствително по-бавно от $\Theta_{ARX}^{(i)}$.

Една особеност при реализацията на метода е следната. Нека $n = \max(na, nb, nc)$ е броят предишни тактове, необходими за формиране на изхода в даден момент. В първата итерация ($i = 1$) първоначално матрицата $E^{(0)} \equiv \tilde{E}^{(0)}$ има $N - n$ реда, като първият ред е $(e_{n+1}^{(0)})^T$. Тъй като първите nc реда в $E^{(0)}$ са нужни за формиране на първия ред в матрицата $\Phi_f^{(1)}$, който е векторът на неизмеримите, но оценени фактори в модела $\varphi_{f,n+nc+1}^{(1)}$, то $\Phi_f^{(1)} \in \mathcal{R}^{N-n-nc \times \ell nc}$. Това значи, че броят на редовете на $\Phi^{(1)}$ в общото представяне на ARMAX модела за интервала на наблюдение са намалели с nc . Следователно в следващата итерация $E^{(1)} \in \mathcal{R}^{N-n-nc \times \ell}$. Отново първите nc реда на $E^{(1)}$ се използват за формиране на първия ред на $\Phi_f^{(2)}$, откъдето общото представяне на модела за интервала на наблюдение във втората итерация се състои от $N - n - 2nc$ реда.

В i -тата итерация броят на уравненията (редовете на $\Phi^{(i)}$) става $N - n - i \cdot nc$. Този ефект на намаляване размерността на $\Phi^{(i)}$ е нежелателен, защото при голям брой итерации може да възникнат проблеми, свързани с преоразмеряването на модела. Затова едно решение на този проблем, както при ELS, така и при останалите методи от този вид, е винаги $\Phi_f^{(i)}$, а оттам и $\Phi^{(i)}$ да се състоят от $N - n$ реда, а липсващите начални стойности на оценената остатъчна грешка на ARMAX модела да се приемат за нули.

В точка 3.2.2.5 е даден пример за получаване на изместени оценки по LS, когато остатъкът е цветен шум. В долния пример е представен ефектът от разширяването на ARX модела с формиращ филтър от тип MA, с който се отчита оцветеността в $\tilde{e}_k$.

Пример. Цветен остатък, ARX и ARMAX модел

За да може директно да се наблюдава изместеността на оценките, долните изследвания са извършени с помощта на симулация, в която оптималните стойности на параметрите са известни. Сценарият на тази симулация е подобен на този в примера от точка 3.2.2.5.

Данните се генерират с ARMAX модел със следните полиномни матрици:

$$A(q^{-1}) = I + \begin{bmatrix} 0.3 & 0.2 \\ 0.8 & -0.3 \end{bmatrix} q^{-1} + \begin{bmatrix} 0.1 & 0.4 \\ -0.5 & 0.7 \end{bmatrix} q^{-2},$$

$$B(q^{-1}) = 0 + \begin{bmatrix} 0.4 & -0.9 \\ -0.5 & -0.8 \end{bmatrix} q^{-1},$$

$$C(q^{-1}) = I + \begin{bmatrix} 0.9 & 0.1 \\ 0.5 & -0.8 \end{bmatrix} q^{-1} + \begin{bmatrix} 0.6 & -0.1 \\ 0.2 & 0.7 \end{bmatrix} q^{-2}.$$

Таблица 3.7. Параметри, оценки по LS, абсолютна и относителна грешка при цветен остатък

θ_i	$\hat{\theta}_i$	$\theta_i - \hat{\theta}_i$	$\frac{\theta_i - \hat{\theta}_i}{\theta_i} \times 100\%$
0.3	0.2026	0.0974	32.4788
0.2	0.2756	-0.0756	-37.7959
0.8	0.6629	0.1371	17.1364
-0.3	-0.1444	-0.1556	51.8625
0.1	0.0865	0.0135	13.4730
0.4	0.4465	-0.0465	-11.6162
-0.5	-0.4219	-0.0781	15.6243
0.7	0.6071	0.0929	13.2736
0.4	0.3902	0.0098	2.4522
-0.9	-0.9380	0.0380	-4.2199
-0.5	-0.5670	0.0670	-13.4084
-0.8	-0.7900	-0.0100	1.2442
0.9	—	—	—
0.1	—	—	—
0.5	—	—	—
-0.8	—	—	—
0.6	—	—	—
-0.1	—	—	—
0.2	—	—	—
0.7	—	—	—

Първите 200 наблюдения се използват за оценяване на параметрите, а останалите 100 – за валидация. Входните сигнали са $u_{i,k} \sim \mathcal{U}(0, 1)$ за $i = \overline{1, 2}$. Формира се остатъкът $e_k \sim \mathcal{N}(0, \frac{1}{16} \Sigma_{y_d})$, където y_d е детерминираният изход на ARX частта, т.е. зависещ от $A(q^{-1})$ и $B(q^{-1})$. Цветният остатък $\tilde{e}_k$ се получава, като стойностите на e_k допълнително се корелират с МА филтъра с матрицата $C(q^{-1})$.

С използване на LS се оценяват параметрите на ARX модел. Полиномите му са от същите редове като на полиномите в модела, използван за генериране на данните (акцентът в примера е върху оценяването на параметрите, а не върху избора на структурните параметри и затова е прието, че структурата е известна). Оптималните параметри, техните оценки, абсолютната и относителната грешка са дадени в Таблица 3.7. Също така се оценяват и параметрите на ARMAX модел, като за целта е използван ELS). Резултатът е даден в Таблица 3.8. Векторът на параметрите θ , чиито стойности са в първата колона на двете таблици, е със структура:

$$\theta = [(\text{vec } A_1^T)^T \quad (\text{vec } A_2^T)^T \quad (\text{vec } B_1^T)^T \quad (\text{vec } C_1^T)^T \quad (\text{vec } C_2^T)^T]^T.$$

Таблица 3.8. Параметри, оценки по ELS, абсолютна и относителна грешка при цветен остатък

θ_i	$\hat{\theta}_i$	$\theta_i - \hat{\theta}_i$	$\frac{\theta_i - \hat{\theta}_i}{\theta_i} \times 100\%$
0.3	0.2674	0.0326	10.8575
0.2	0.2560	-0.0560	-27.9991
0.8	0.7602	0.0398	4.9809
-0.3	-0.2701	-0.0299	9.9772
0.1	0.1286	-0.0286	-28.5785
0.4	0.4234	-0.0234	-5.8429
-0.5	-0.4836	-0.0164	3.2700
0.7	0.6583	0.0417	5.9579
0.4	0.4005	-0.0005	-0.1341
-0.9	-0.9042	0.0042	-0.4660
-0.5	-0.5510	0.0510	-10.2025
-0.8	-0.7458	-0.0542	6.7775
0.9	0.6874	0.2126	23.6240
0.1	0.0279	0.0721	72.1329
0.5	0.4912	0.0088	1.7687
-0.8	-0.6541	-0.1459	18.2420
0.6	0.3645	0.2355	39.2433
-0.1	-0.1514	0.0514	-51.4170
0.2	0.1063	0.0937	46.8341
0.7	0.3344	0.3656	52.2279

Показателят VAF за двата изхода на моделите, получени с LS и с ELS, е съответно:

$$\text{VAF}_{\text{LS}} = \begin{bmatrix} 87.91 \\ 80.83 \end{bmatrix} \quad \text{и} \quad \text{VAF}_{\text{ELS}} = \begin{bmatrix} 93.40 \\ 85.39 \end{bmatrix}.$$

В първата таблица се наблюдава изместеността на $\hat{\theta}$ на ARX модела, като обяснението за това (виж точка 3.2.2.4) е неподходящата структура на описанието. Един вариант за подобряване на модела е типът ARX да се запази, а степените на полиномите да нараснат, но обикновено това води до полиноми от висок ред. От друга страна, ако се потърси ARMAX модел, често броят на параметрите в крайното описание е по-малък. Цената за това е използването на итеративни методи за оценка, тъй като моделът става нелинейно параметризиран.

От втората таблица се вижда, че с добавяне на МА филтъра разликата между действителните параметри и техните оценки чувствително намалява.

Също така е приложен и автокорелационният тест на остатъка (точка 2.4.3.3). За ARX модела процентните съотношения на случаите, в които

$$|r_{e,i}| \geq 1.96 \text{ за } i = \overline{1, 49},$$

е съответно 12% и 24%. Тъй като и за двата изхода тези съотношения са повече от 5%, не може да се приеме, че остатъкът е бял шум. От друга страна, за ARMAX модела горното неравенство се изпълнява в 2% от случаите. Това означава, че ARMAX моделът отразява закономерните зависимости между входовете и изходите, което води до остатък, между стойностите на който не се наблюдават значими корелации. $\diamond$

3.2.6 Метод на инструменталните променливи

До момента е представен LS и неговите обобщения WLS и GLS, които при нарушаване на приемането за ергодичност и за липса на отговорност на остатъка, осигуряват ефективни оценки. Характерно за тези методи е, че се изменя целевата функция, като се добавят тегла и нови членове. Когато остатъкът е цветен шум, често се прибегва до описания по-горе метод ELS.

Различен подход за получаване на неизместени оценки се реализира с метода на инструменталните променливи (IV). За едномерни системи той е разгледан в [11, 12, 123]. При него $\mathcal{F}(\Theta)$ остава същата като при LS, а се изменя уравнението на модела, когато той е описан в общ вид.

3.2.6.1 Уравнение на модела

Нека отново с $\tilde{e}_k$ се означаи оцветеният остатък, който се получава, ако структурата на модела не е подходяща, ако системата е нестацио-

нарна или когато има фактори, причислени към влиянието от околната среда, но оказващи известно влияние върху поведението на обекта. Нека също матрицата $\tilde{E}$ да съдържа стойностите на тези остатъци по отделните изходи за интервала на наблюдение. Тогава уравнението на модела с използване на наличните данни е

$$Y = \Phi\Theta + \tilde{E}. \quad (3.45)$$

При IV (3.45) се умножава с матрица, която има същата размерност като матрицата на данните Φ . Нека това е $Z \in \mathcal{R}^{N-n \times z}$ и тогава уравнението на модела става:

$$Z^T Y = Z^T \Phi \Theta + Z^T \tilde{E}.$$

Стълбовете на Z са стойности на т.нар. инструментални променливи. Те се определят по такъв начин, че да са корелирани с данните, т.е. $Z^T Y \neq 0$ и $Z^T \Phi \neq 0$, но да не са корелирани с остатъка. От независимостта между стълбовете на Z и $\tilde{E}$ (виж графичната интерпретация на LS в точка 2.3.4.5) следва, че Z_i лежат в пространството, което е ортогонално на $\text{range } \tilde{E}$ (това е пространството, в което стълбовете на $\tilde{E}$ формират базис). Когато, в допълнение на това, остатъкът е центриран, $Z^T \tilde{E} = 0$ и тогава (3.45) се свежда до

$$Z^T Y = Z^T \Phi \hat{\Theta}.$$

Когато Z и матрицата на данните са от пълен ранг, а това означава, че обектът е възбуден по всички канали, $Z^T \Phi$ е добре обусловена и за оптималните оценки се получава:

$$\hat{\Theta} = (Z^T \Phi)^{-1} Z^T Y.$$

Независимостта на Z от $\tilde{E}$ осигурява неизместени оценки. Когато $Z = \Phi$, този метод съвпада с LS.

Нека моделът е ARX, а това означава, че Φ съдържа предисторията на изхода. Когато e_k е БШ, тогава Φ и E са независими, понеже регресорите до $k - 1$ -ия такт не зависят от e_k . С други думи даден ред на Φ не е корелиран със съответния ред на E и в резултат на това $\Phi^T E = 0$. Но когато остатък $\tilde{e}_k$ е ЦШ (т.е. $\exists i > 0$, за което $\text{cov}(\tilde{e}_k, \tilde{e}_{k-i}) \neq 0$), регресорите до $k - 1$ -я такт са корелирани с $\tilde{e}_k$, понеже както те, така и $\tilde{e}_k$ зависят от стойности на остатъка до $k - 1$ -я такт. Това означава, че $\Phi^T \tilde{E} \neq 0$, или с други думи $\text{range } \Phi$ и $\text{range } \tilde{E}$ не са ортогонални. Следователно матрицата на данните не е подходяща за инструментална матрица.

От друга страна, дисперсията на $\hat{\Theta}$ е пропорционална на $(Z^T \Phi)^{-1}$ (виж (3.28)) и затова, колкото по-корелирани са Z и Φ , толкова по-малка е дисперсията на оценките $\sigma_{\hat{\Theta}}^2$.

3.2.6.2 Алгоритъм на IV

В началото на темата за оценяването на модел с матрица на параметрите е показана структурата на Φ за ARX модел. Тя е

$$\Phi = [-Y_{-1} \quad -Y_{-2} \quad \dots \quad -Y_{-na} \quad U_{-1} \quad U_{-2} \quad \dots \quad U_{-nb}].$$

Матриците U_{-i} не зависят от $\tilde{E}$. Затова те са идеални инструментални променливи. Но от горните разглеждания стана ясно, че Y_{-i} зависят от $\tilde{E}$, и поради това те не са подходящи за стълбове на Z . Един вариант за попълване на лявата част на инструменталната матрица, като се гарантира, че стълбовете ѝ лежат в пространство, ортогонално на $\text{range } \tilde{E}$, е да се използва модел, с който Y_{-i} да се оценят само с помощта на матриците U_{-i} . От некорелираността между U_{-i} и $\tilde{E}$ следва, че и оценките $\hat{Y}_{-i}$ няма да зависят от $\tilde{E}$. Тази идея е залегнала в следния алгоритъм.

Алгоритъм на IV

1. Инициализация

Задават се условията за спиране, например толеранси на изменение на $\mathcal{F}(\Theta)$ или на оценките в съседни итерации. Също така се формира Φ и се прилага LS:

$$\Theta^{(0)} = (\Phi^T \Phi)^{-1} \Phi^T Y.$$

2. Оценяване на Y

В i -тата итерация се формира матрицата:

$$Y^{(i)} = [y_{n+1}^{(i)} \quad y_{n+2}^{(i)} \quad \dots \quad y_N^{(i)}]^T,$$

като текущият ред $Y^{(i)}$ е векторът:

$$y_k^{(i)} = \Theta^{(i-1)} \varphi'_k,$$

а векторът на регресорите, независещ от изхода на системата, е

$$\varphi'_k = [-y_{k-1}^{(i)T} \quad \dots \quad -y_{k-na}^{(i)T} \quad u_{k-1}^T \quad \dots \quad u_{k-nb}^T]^T.$$

3. Формиране на Z

$$Z^{(i)} = [-Y_{-1}^{(i)} \quad -Y_{-2}^{(i)} \quad \dots \quad -Y_{-na}^{(i)} \quad U_{-1} \quad U_{-2} \quad \dots \quad U_{-nb}].$$

4. *Оценяване на Θ по IV*

$$\Theta^{(i)} = (Z^{(i)T} \Phi)^{-1} Z^{(i)T} Y.$$

5. *Проверка за спиране*

Ако избраното условие за спиране е изпълнено, итеративното оценяване спира и се приема, че търсените параметри са $\hat{\Theta} = \Theta^{(i)}$. В противен случай $i \leftarrow i+1$ и се продължава от стъпка 2. Подходящо условие, свързано със стойността на целевата функция, е

$$|\mathcal{F}(\Theta^{(i)}) - \mathcal{F}(\Theta^{(i-1)})| < \tau_F.$$

Друго условие за спиране по отношение на изменението на оценките е

$$\max |\Theta^{(i)} - \Theta^{(i-1)}| < \tau_{\Theta}.$$

Тъй като първоначалните оценки по LS са изместени, то и $\hat{Y}^{(1)}$ е неточна оценка на изхода на системата. Въпреки това $Z^{(1)}$ и Φ са корелирани, а $Z^{(1)}$ и $\tilde{E}$ – не. Затова на всяка итерация оценката $\hat{Y}^{(1)}$ и съответно $\Theta^{(i)}$ е все по-малко изместена.

Съществуват много варианти за формиране на Z . Например тя може да се състави без използването на модел, с което се избягва допълнителното оценяване. Един вариант е

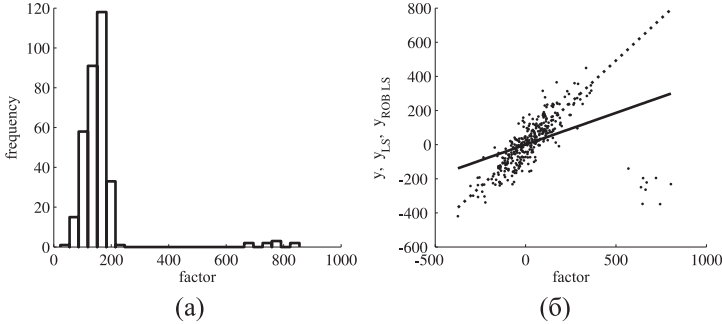
$$Z = [U_{-1} \ U_{-2} \ \dots \ U_{-na-nb}]^T.$$

Също така подобрение на описания алгоритъм може да се постигне, ако оцветеността на остатъка на модела от стъпка 4 допълнително се оцени [119]. Тъй като това е зависимост на стойностите на сигнала от предисторията им, за целта се използва AR модел.

3.2.7 Робастен метод на най-малките квадрати

LS е BLUE при някои разпределения на остатъка, едно от които е нормалното. Въпреки това този метод осигурява задоволителни оценки и в много практически задачи, при които оценките не са BLUE. От друга страна, има разпределения, при които моделите, получени по LS, не са приемливи – такива случаи са разгледани във Втора глава, в темата за предварителната обработка на данните, когато се наблюдават нехарактерни стойности. Често при тези случаи разпределенията имат т.нар. „тежки опашки“ (Фигура 3.5 (а)). Нехарактерните стойности в опашките може значително да влошат оценяването на параметрите. Едно решение е те да се премахнат на етапа на предварителната обработка

на данните. Невинаги това е лесно и тогава е подходящо да се използва целева функция, която не е чувствителна към решението на описания проблем. Това е довело до разработването на т.нар. робастни оценители на параметрите [38, 111] (Фигура 3.5 (б)). Те може да се разглеждат като обобщение [94] на метода на максималното правдоподобие.



Фигура 3.5. Хистограма на сигнал с т.нар. „тежка опашка“ (а); за същите данни: изход на модел, получен по LS (плътна линия) и по робастен LS (линия от точки) (б)

Нека се въведе функцията $\zeta_{i,k} = \zeta(e_{i,k})$, която има смисъл на принос на i -тия остатък в k -тия такт към стойността на целевата функция (в неявен вид това е приносът на данните: i -тият изход и факторите, асоциирани с k -тия такт). В общия случай тя трябва да притежава следните свойства:

$$\begin{aligned} \zeta(e_{i,k}) &\geq 0, \\ \zeta(0) &= 0, \\ \zeta(e_{i,k}) &= \zeta(-e_{i,k}), \\ \zeta(e_{i,k_1}) &\geq \zeta(e_{i,k_2}), \text{ когато } |e_{i,k_1}| \geq |e_{i,k_2}|. \end{aligned}$$

LS и WLS са частни случаи на робастния LS (RobLS). Когато функцията на приноса е $\zeta_{i,k} = e_{i,k}^2$, се достига до LS, а когато $\zeta_{i,k} = w_{i,k} e_{i,k}^2$ – до WLS.

Нека приносите в k -тия такт да са елементи на вектора $\zeta_k \in \mathcal{R}^\ell$, т.е.

$$\zeta_k = [\zeta_{1,k} \ \zeta_{2,k} \ \dots \ \zeta_{\ell,k}]^T.$$

Нека също първата производна на $\zeta_{i,k}$ е $\gamma_{i,k}$ и се въведе векторът:

$$\gamma_k = [\gamma_{1,k} \ \gamma_{2,k} \ \dots \ \gamma_{\ell,k}]^T,$$

съдържащ тези производни. Тези вектори се използват при извеждането на RobLS.

3.2.7.1 Оценки по RobLS

За да се намерят оптималните параметри, трябва да се определи $\nabla \mathcal{F}(\Theta)$. От общия вид на модела с матрица на параметрите:

$$y_k = \Theta^T \varphi_k + e_k,$$

лесно се вижда, че i -тата компонента на e_k зависи само от елементите на i -тия стълб Θ на матрицата на параметрите. Това е така, понеже $e_{i,k} = y_{i,k} - \Theta_{.i}^T \varphi_k$. Следователно производните на $\zeta(e_{i,k})$ по параметър, който не е от i -тия стълб на $\Theta_{.i}$, са равни на нула, а ненулеви са производните на $\zeta(e_{i,k})$ единствено по параметрите от i -тия стълб, т.е.

$$\frac{\partial \zeta_{i,k}}{\partial \theta_{jh}} \begin{cases} = 0, & \text{за } h \neq i, \\ \neq 0, & \text{за } h = i. \end{cases} \quad (3.46)$$

До момента индексът i отговаря на текущия изход, на даден стълб на Θ , а съответно и на текущия MISO подмодел. За да се запази това означение, с θ_{ji} ще се обозначи j -тият елемент от текущия i -ти стълб на матрицата Θ .

Ненулевите производни от (3.46) са:

$$\frac{\partial \zeta_{i,k}}{\partial \theta_{ji}} = \frac{\partial \zeta_{i,k}}{\partial e_{i,k}} \frac{\partial e_{i,k}}{\partial \theta_{ji}} = -\gamma_{i,k} \varphi_{j,k},$$

за $i = \overline{1, \ell}$. В последното равенство е отчетено, че остатъкът $e_{i,k} = y_{i,k} - \Theta_{.i}^T \varphi_k$ е линейната функция на параметрите и производната ѝ по θ_{ji} е $\frac{\partial e_{i,k}}{\partial \theta_{ji}} = -\varphi_{j,k}$.

Както при предишните методи, тъй като $\mathcal{F}(\Theta)$ е скалярна функция, то диференцирайки по матрицата $\Theta \in \mathcal{R}^{z \times \ell}$, се получава градиентната матрица $G \in \mathcal{R}^{z \times \ell}$. Нейният ji -ти елемент е

$$g_{ji} = \frac{\partial \mathcal{F}(\Theta)}{\partial \theta_{ji}} = \sum_{k=n+1}^N \frac{\partial \zeta_{i,k}}{\partial \theta_{ji}} = - \sum_{k=n+1}^N \gamma_{i,k} \varphi_{j,k},$$

а за i -тия i стълб се получава:

$$G_{.i} = \nabla_{\Theta_{.i}} \mathcal{F}(\Theta) = - \sum_{k=n+1}^N \gamma_{i,k} \varphi_k.$$

Нека се въведат отношенията:

$$w_{i,k} = \frac{\gamma_{i,k}}{e_{i,k}}, \quad (3.47)$$

откъдето $\gamma_{i,k}$ се изразява като $\gamma_{i,k} = w_{i,k}e_{i,k}$. Тогава i -тият стълб на G , изразен с $w_{i,k}$, е

$$G_{..i} = - \sum_{k=n+1}^N w_{i,k}e_{i,k}\varphi_k.$$

Обобщавайки израза за всички стълбове, градиентната матрица може да се запише като:

$$G = - \sum_{k=n+1}^N \varphi_k(w_k * e_k)^T.$$

Нека по аналогия с WLS се въведе тегловният тензор W , който се състои от матриците:

$$W_{..i} = \text{diag}(w_{i,n+1}, w_{i,n+2}, \dots, w_{i,N}). \quad (3.48)$$

С използване на тензора W и на матриците Φ и E , градиентната матрица може да се запише съкратено по следния начин:

$$G = W \circ_1 \Phi \bar{\circ}_2^3 E.$$

Този израз съвпада с (3.32), като тук, ако $\mathcal{F}(\Theta) = \text{sum}(W \bar{\circ}_1^3 E \bar{\circ}_2^3 E)$, то ненулевите елементи на W са $2w_{i,k}$ поради квадратичния характер на $\mathcal{F}(\Theta)$. Тогава аналогично G може да се преработи във вида:

$$G = W \circ_1 \Phi \bar{\circ}_2^3 Y - F \bar{\circ}_2^3 \hat{\Theta}.$$

Тензорът $F = W \circ_1 \Phi \circ_2 \Phi$ съдържа матриците на Фишер $F_{..i} = \Phi^T W_i \Phi$. Нека тензорът P да се състои от техните обратни матрици, като $P_{..i} = F_{..i}^{-1}$. Тогава от условието:

$$G|_{\Theta=\hat{\Theta}} = W \circ_1 \Phi \bar{\circ}_2^3 Y - F \bar{\circ}_2^3 \hat{\Theta} = 0$$

за оптималните оценки на параметрите се получава:

$$\hat{\Theta} = P \bar{\circ}_2^3 (W \circ_1 \Phi \bar{\circ}_2^3 Y).$$

Оценките съвпадат с тези по WLS, а това означава, че с въвеждане на теглата (3.47) задачата за оценяване на параметрите по RobLS се свежда до WLS. При тази постановка от (3.47) се вижда, че теглата зависят от остатъците, които пък зависят от оценките, а те от своя страна зависят от теглата. Затова робастният оценител се реализира като

итеративна процедура. На всяка итерация се формират теглата, зависещи от остатъците, получени от предишната итерация. След това се прилага стандартният WLS и получените оценки се използват за изчисляване на новите остатъци. В следващата итерация те се използват при формирането на нови тегла и т.н.

Има и други варианти на RobLS, с които се реализира идеята на робастното оценяване. Методът намира приложение в области като социология, пазарни системи и др., при които е характерно наличието на тежки опашки в разпределението на остатъка. Такъв алгоритъм е описан по-долу. При него теглата се формират, като

$$w_{j,k}^{(i)} = \min(1, \max(\delta, |e_{j,k}^{(i)}|^{-1})), \quad (3.49)$$

а δ е малко положително число, например $\delta = 10^{-8}$. Целта е теглата да намаляват, ако грешката расте, което е характерно за отдалечените наблюдения от общата тенденция. По този начин се потиска нежеланият ефект от нехарактерните наблюдения.

3.2.7.2 Алгоритъм на RobLS

1. Инициализация

Задават се условията за спиране, формира се Φ и се прилага LS:

$$\Theta^{(0)} = (\Phi^T \Phi)^{-1} \Phi^T Y.$$

Това е равносилно на WLS с тегловен тензор, за който $W_{..j}^{(0)} = I$.

2. Определяне на остатъка

В i -тата итерация се формира матрицата $E^{(i)}$:

$$E^{(i-1)} = Y - \Phi \Theta^{(i-1)}.$$

3. Формиране на тегловния тензор

Изчисляват се новите тегла $w_{j,k}^{(i)}$ според (3.49). Те са необходими за формиране на новия тегловен тензор $W^{(i)}$ от (3.48).

4. Оценяване на параметрите по WLS

Параметрите на модела се определят от

$$\Theta^{(i)} = P^{(i)} \bar{\circ}_2^3 (W^{(i)} \circ_1 \Phi \bar{\circ}_2^3 Y),$$

където тензорът $P^{(i)}$ се формира от матриците $P_{..j} = (\Phi^T W_j^{(i)} \Phi)^{-1}$.

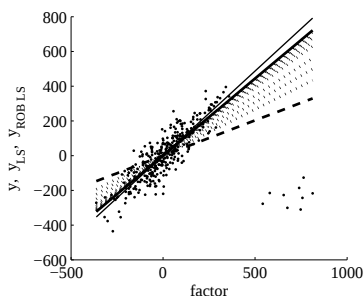
5. Проверка за спиране

Ако условието за спиране е изпълнено, итеративното оценяване се преустановява и се приема, че търсените параметри са $\hat{\Theta} = \Theta^{(i)}$. В противен случай $i \leftarrow i + 1$ и се продължава от стъпка 2.

В следващите два примера се представят възможностите на RobLS. В първия пример са използвани генерирани данни, а във втория се показва подобрението на модела спрямо този, получен по LS, за случая, когато в реални данни се наблюдават нехарактерни наблюдения.

Пример. LS и RobLS за данни с тежка опашка

За представяне на подобрението на модела с използване на RobLS са генерирани данни с разпределение, в което се наблюдава тежка опашка (т.е. се съдържат нехарактерни стойности). Изследваната система е едномерна, а моделът отчита връзката между изхода и един фактор. При така зададения сценарий резултатите се представят удобно в графичен вид.



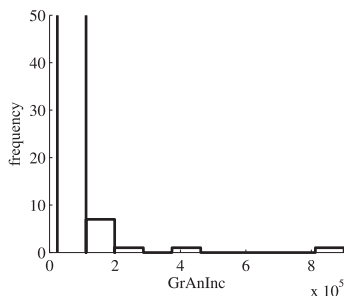
Фигура 3.6. Наблюдения (точки); модел, получен с LS и набора без нехарактерни стойности (плътна тънка линия); модел, получен с LS и набора с нехарактерните стойности (прекъсната линия); модели от междинните итерации на RobLS (линии от точки); модел по RobLS (плътна линия) след схождение на оценките

Формирани са 321 стойности на фактора, които имат нормално разпределение. Използван е LS, за да се оцени параметърът на модела. Също така към набора са добавени още 9 нехарактерни стойности и са приложени LS и RobLS.

Данните, както и моделите, формирани по LS и RobLS, са показани на Фигура 3.6, като за робастния вариант са дадени и моделите от междинните итерации. Тъй като всички модели са линейни, то и зависимостите между изхода и фактора са прави линии. Първоначалният модел от инициализиращата стъпка на RobLS съвпада с LS, приложен към същите данни. Вижда се, че наличието на нехарактерни стойности силно отдалечава модела (представен с пунктирна линия) от преобладаващата зависимост между фактора и изхода. С всяка следваща ите-

рация на RobLS моделът се доближава до този, който би се получил по LS, ако в данните нямаше нехарактерни стойности. ◇

Пример. *Кредитен риск: LS и RobLS за данни с тежка опашка*



Фигура 3.7. Нехарактерни стойности за фактора ‘GrAnInc’ – годишен доход

В този пример отново се използват данните, предоставени от „Експирин“ [87]. Независимите характеристики (фактори) са 17 и описват сегмент от $N = 3000$ кандидати за кредит. В този твърде малък набор (обикновено N е от порядъка на $10^5 - 10^6$) се съдържат нехарактерни стойности на фактора ‘GrAnInc’, както е показано на Фигура 3.7. Честотата на годишните доходи между 24 000 и 111 600 евро е 2990 и не е показана, за да може нехарактерните стойности да се видят.

Изследваната система е статична, с един изход и съответно моделът е с вектор на параметрите. Добавен е и постоянен член. В Таблица 3.9 са представени стойностите на VAF, изчислени за моделите от всяка итерация на RobLS. Първият модел, както в предния пример, съвпада с този, получен по LS. Вижда се, че от втора итерация, с въвеждане на теглата, моделите са значително по-достоверни. От таблицата се вижда, че нехарактерните наблюдения (лостови точки) водят до три пъти по-неточен модел ($\text{VAF} = 5.41\%$) в сравнение с този, получен от RobLS. Цената на това подобрение е итеративното оценяване на параметрите. ◇

3.3 Методи за модел с вектор на параметрите

В предишната точка са представени най-разпространените методи за оценяване на параметри, базирани на представяне на линейно параметризирания модел в общ вид с матрица на параметрите. В тази точка

Таблица 3.9. VAF за моделите от всяка итерация на RobLS. Моделът от нулевата итерация е формиран в инициализиращата стъпка с LS.

итерация	VAF %
0	5.41
1	15.82
2	15.87
3	15.89

същите методи са разгледани, но при извеждането им е заложено описание на модела с вектор на параметрите. Задачата на оценяване може да се декомпозира на множество подзадачи, всяка от които е определяне на параметрите на MISO подмодел. При такъв подход към MIMO системата във всяка подзадача изходът е един, което опростява вида на оценяваните модели. Въпреки това методите за оценяване на параметрите на MIMO системи ще се представят за многовходови и многоизходови модели. Предимствата (както и понякога необходимостта) на директното формиране на MIMO модела са изложени в точка 3.1.2.

3.3.1 Модел в общ вид с вектор на параметрите

Линейният регресионен модел може да се представи в следния общ вид:

$$y_k = \Phi_k \theta + e_k. \quad (3.50)$$

При този запис параметрите са подредени във вектора $\theta \in \mathcal{R}^p$, а матрицата $\Phi_k \in \mathcal{R}^{\ell \times p}$ съдържа регресорите, описващи изхода в k -тия такт.

Нека за представяне на системата се използва ARX модел:

$$A(q^{-1})y_k = B(q^{-1})u_k + e_k. \quad (3.51)$$

В точка 2.1.3.3 е описано преобразуването на (3.51) във вид (3.50). Има различни варианти на общия вид с вектор на параметрите, при които структурата на матрицата Φ_k се променя (с пренареждане на елементите на θ се изменя и позицията на регресорите във Φ_k). По-долу задачата за преобразуване е разгледана по-общо – без да се уточнява напълно подредбата на параметрите. След това е уточнен конкретният вид на θ и Φ_k , който е използван в монографията.

Нека полиномната матрица, отразяваща авторегресията в (3.51), се запише като:

$$A(q^{-1}) = I_\ell + A_1 q^{-1} + \dots + A_{na} q^{-na} = I_\ell + \tilde{A}(q^{-1}).$$

Тогава за ARX модела се получава

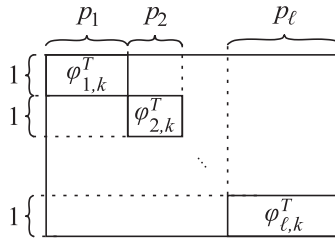
$$y_k = -\tilde{A}(q^{-1})y_k + B(q^{-1})u_k + e_k. \quad (3.52)$$

Нека матриците $Y_k \in \mathcal{R}^{\ell \times \ell^2}$ и $U_k \in \mathcal{R}^{\ell \times \ell m}$ са блокдиагонални, като всеки блок по диагонала им е съответно y_k^T и u_k^T . Нека също се въведат полиномните вектори $\underline{a}(q^{-1}) = \text{vec } \tilde{A}^T(q^{-1})$ и $\underline{b}(q^{-1}) = \text{vec } B^T(q^{-1})$. В такъв случай за (3.52) се получава:

$$y_k = -Y_k \underline{a}(q^{-1}) + U_k \underline{b}(q^{-1}) + e_k. \quad (3.53)$$

При извършване на умножението между Y_k и $\underline{a}(q^{-1})$ се получава вектор, чийто i -ти елемент е скаларното произведение на i -тия ред на $\tilde{A}(q^{-1})$ и вектора y_k . Този елемент е сумарното влияние на предисторията на изходите върху i -тия изход. Съответно резултатът от умножението между U_k и $\underline{b}(q^{-1})$ е вектор, чийто i -ти елемент е скаларното произведение на i -тия ред на $B(q^{-1})$ и u_k . Този елемент отразява сумарното влияние на входовете върху i -тия изход. Тъй като полиномите в $\tilde{A}(q^{-1})$ и $B(q^{-1})$ нямат свободни членове, то наистина всички регресори от дясната страна на (3.52) и съответно на (3.53), след умножението им по полиномите, са до $k-1$ -ия такт.

За да се стигне до общото представяне на модела с вектор на параметрите, е необходимо те да се отделят във вектора θ . Освен това всички регресори трябва да се подредят в матрицата на регресорите Φ_k по такъв начин, че ненулевите елементи от i -тия ѝ ред да са регресорите, описващи i -тия изход, и то подредени така, че да отговарят на подредбата на параметрите във вектора θ . В резултат на това се получава (3.50).



Фигура 3.8. Структура на матрица на регресорите Φ_k на ARX модел в k -тия такт

Параметрите може да се подредят по различен начин. Например, първите елементи на θ може да са всички параметри на полиноми от $A(q^{-1})$ (подредени последователно според редовете на $A(q^{-1})$), следвани от параметрите на полиномите от $B(q^{-1})$. За да отговаря Φ_k на тази

подредба, матрицата на регресорите се състои от две блокдиагонални матрици. Първата е AR частта, а другата съдържа предисторията на входовете. В монографията е прието подредбата на параметрите в θ да е следната. Първите $p_1 = \sum_{j=1}^{\ell} na_{1j} + \sum_{j=1}^m nb_{1j}$ параметри отговарят съответно на полиномите от първите редове на $\tilde{A}(q^{-1})$ и $B(q^{-1})$. Те участват в описанието на първия изход. Следват параметрите, описващи втория изход, и т.н. По този начин структурата на матрицата Φ_k е блокдиагонална, като всеки блок съдържа предисторията на изходите, следвана от тази на входните сигнали. Тази структура е показана на Фигура 3.8.

За определянето на параметрите на базата на наличните данни е нужно да се формира уравнение, в което в явен вид се обвързват данните, параметрите и остатъкът за интервала на наблюдение. За целта се въвеждат векторите:

$$y = [y_{n+1}^T \ y_{n+2}^T \ \dots \ y_N^T]^T \in \mathcal{R}^{(N-n)\ell},$$

$$e = [e_{n+1}^T \ e_{n+2}^T \ \dots \ e_N^T]^T \in \mathcal{R}^{(N-n)\ell},$$

както и матрицата:

$$\Phi = [\Phi_{n+1}^T \ \Phi_{n+2}^T \ \dots \ \Phi_N^T]^T \in \mathcal{R}^{(N-n)\ell \times p},$$

зависеща от вътрешната структура на (3.50). Тогава уравнението на модела за k -тия такт (3.50) може да се обобщи за интервала на наблюдение (т.е. за $k = \overline{n+1, N}$), в резултат на което се получава:

$$y = \Phi\theta + e.$$

Основният недостатък на общото представяне с матрица на параметрите се дължи именно на подредбата им в матрица. Когато степените на полиномите са различни, за конструиране на Θ се налага въвеждането на допълнителни елементи в полиномите от по-нисък ред, като целта е редът на всички полиноми в съответната полиномна матрица (или в неин стълб) да станат равни на най-голямата степен на полином в матрицата (в стълба).

За разлика от вида на модела с матрица на параметрите при представянето му с вектор θ има пълна свобода при избора на реда на полиномите в описанието [153]. Това води до повече възможности при определянето на структурата. За SISO модели, където полиномните матрици се свеждат до полиноми, няма разлика между реализациите на оценителите за модел с вектор и с матрица на параметрите. Но при MIMO моделите, където броят на параметрите може значително да нарасне,

тази допълнителна свобода позволява да се получи модел с подходяща структура, като се избягват някои проблеми, типични за многомерните описания, като ненужно усложняване на структурата, преоразмеряване на модела и лоша обусловеност на задачата.

Поради повечето възможности при избора на структурните параметри построяването на модел с вектор на параметрите може да отнеме значително повече време от изграждането на модел с матрица на параметрите. Затова е удачно оценителите, изведени за общия вид, разгледан по-долу, да се използват за прецизно уточняване на структурата на модела, а при много голям брой входове и изходи тази дейност да се предхожда от построяване на модел с матрица на параметрите (за бързата, но по-груба ориентация в данните).

От горното преобразуване се вижда, че $\Phi \in \mathcal{R}^{(N-n)\ell \times p}$ е със значително по-голяма размерност от варианта с матрица на параметрите, където $\Phi \in \mathcal{R}^{(N-n) \times z}$ (при динамични системи $z < p$, като разликата обикновено е в пъти). Има случаи, когато при голям брой наблюдения матрицата на данните може да нарасне толкова, че да не е възможно да се извърши оценяване на параметрите. Този проблем е преодолим, тъй като Φ е разрежена матрица – състои се от много нулеви елементи, чиито позиции са предварително известни. Затова при разработката на оценителите тя трябва да се представя именно като разрежена (sparse) матрица.

3.3.2 Метод на най-малките квадрати

По-долу е представена постановката на задачата на най-малките квадрати и са изведени оптималните оценки.

3.3.2.1 Показател на качеството

В LS се минимизира сумата от квадратите на остатъците, т.е. показателят на качеството (целевата функция) е

$$\mathcal{F}(\theta) = \sum_{k=n+1}^N \sum_{i=1}^{\ell} e_{i,k}^2.$$

Тъй като в (3.50) остатъците за интервала на наблюдение са групирани във вектор, $\mathcal{F}(\theta)$ може да се запише така:

$$\mathcal{F}(\theta) = e^T e.$$

За разлика от предишния подход тук $\mathcal{F}(\theta)$ е скаларна функция на векторния аргумент θ . Изключването на вектора e (неизвестен преди оценяването) от израза на целевата функция може да стане със следния запис:

$$\mathcal{F}(\theta) = \|e\|_2^2 = \|y - \Phi\theta\|_2^2. \quad (3.54)$$

Така $\mathcal{F}(\theta)$ зависи от търсените параметри и наличните данни. Този вид на целевата функция се използва при разглеждането на числените реализации на методите за оценяване на параметри. $\|\cdot\|_2$ е норма-2 на вектор и има смисъл на дължина на вектора. Нека x е произволен вектор с реални елементи. В такъв случай $\|x\|_2$ е квадратен корен от сумата от квадратите на елементите на x , т.е. $\|x\|_2 = \sqrt{x^T x}$. От (3.54) се вижда, че стойността на $\mathcal{F}(\theta)$ е всъщност квадратът на дължината на вектора e (този вектор е представен графично в точка 2.3.4.5).

В долните извеждания $\mathcal{F}(\theta)$ се представя с по-краткия запис $\mathcal{F}(\theta) = e^T e$. Тогава критерият, заложен в LS (за модел с вектор на параметрите), се записва по следния начин:

$$\min_{\theta} \mathcal{F}(\theta) = \min_{\theta} e^T e. \quad (3.55)$$

3.3.2.2 Оценки по LS

Отново се извършват трите стъпки, описани в точка 3.2.2.2, като разликата с предишното извеждане е, че тъй като θ е вектор, то и градиентът $g = \nabla \mathcal{F}(\theta)$ е вектор, а не матрица. При определянето на g се използват някои зависимости, валидни за диференцирането на скаларна функция на векторен аргумент. Нека са дадени матриците $A \in \mathcal{R}^{n \times n}$ и $B \in \mathcal{R}^{m \times n}$, както и векторите $x \in \mathcal{R}^n$ и $b \in \mathcal{R}^m$. Тогава са в сила съотношенията:

$$\nabla_x (x^T B^T b) = \nabla_x (b^T B x) = B^T b,$$

$$\nabla_x (x^T A x) = (A + A^T)x.$$

Когато A е симетрична, какъвто е случаят (в долните разглеждания $A = \Phi^T \Phi$), то $\nabla_x x^T A x = 2Ax$.

По-долу се определя оптималният вектор $\hat{\theta}$, изразен чрез наличните входно-изходни данни. Първата стъпка е $\mathcal{F}(\theta)$ да се представи като

функция на θ , т.е.

$$\begin{aligned}\mathcal{F}(\theta) &= e^T e \\ &= (y - \Phi\theta)^T (y - \Phi\theta) \\ &= y^T y - \theta^T \Phi^T y - y^T \Phi\theta + \theta^T \Phi^T \Phi\theta.\end{aligned}$$

След това последният израз се диференцира. В сила са зависимостите:

$$\begin{aligned}\nabla_{\theta}(y^T y) &= 0, \\ \nabla_{\theta}(\theta^T \Phi^T y) &= \nabla_{\theta}(y^T \Phi\theta) = \Phi^T y, \\ \nabla_{\theta}(\theta^T \Phi^T \Phi\theta) &= 2\Phi^T \Phi\theta,\end{aligned}$$

откъдето за градиента на $\mathcal{F}(\theta)$ се получава:

$$g = -2\Phi^T y + 2\Phi^T \Phi\theta.$$

Той се приравнява на нула и се определя $\hat{\theta}$, т.е.

$$g|_{\theta=\hat{\theta}} = 0_p \Leftrightarrow -\Phi^T y + \Phi^T \Phi\hat{\theta} = 0_p.$$

Векторът $\hat{\theta}$, при който g се нулира, е

$$\hat{\theta} = (\Phi^T \Phi)^{-1} \Phi^T y.$$

Анализът на решението, извършен за предишния подход, е в сила и за оценката $\hat{\theta}$, т.е. LS е BLUE оценител и съответно осигурява неизмествени и ефективни оценки, когато e_k е центриран бял Гаусов шум. Този анализ е много сходен за описанието с вектор на параметрите и по тази причина не е извършен.

3.3.3 Метод на претеглените най-малки квадрати

По-долу е представено обобщението WLS на стандартния LS, когато стойностите на остатъка се отчитат с различно тегло в целевата функция. В точка 3.2.3 е описан WLS за модел с матрица на параметрите. Там са разгледани различни приложения на метода. Затова по-долу вниманието е насочено единствено върху особеностите на WLS, когато описанието е с вектор на параметрите.

3.3.3.1 Показател на качеството

За модел с матрица на параметрите компактното представяне на оптималната оценка $\hat{\Theta}$, получена по WLS, е свързано с използването на тензори. За модел в общ вид с вектор на параметрите, извеждането на WLS значително се опростява, като не е нужно въвеждането на тензори.

Показателят на качеството, който се минимизира в WLS, е претеглена сума на квадратите на остатъците по всички изходи. За модел с вектор на параметрите тази сума се записва компактно по един от следните начини:

$$\mathcal{F}(\theta) = e^T W e = \|e\|_W^2 = \|y - \Phi\theta\|_W^2.$$

$W \in \mathcal{R}^{\ell(N-n) \times \ell(N-n)}$ е тегловната матрица, която е диагонална и се формира като:

$$W = \text{diag}[w_{n+1}^T \quad w_{n+2}^T \quad \dots \quad w_N^T].$$

Поради голямата размерност на W и факта, че се състои основно от нули, е подходящо тази матрица да се представи в паметта като разрежена. Векторите $w_k \in \mathcal{R}^\ell$ съдържат теглата (неотрицателни) по отношение на остатъците по всеки изход в k -тия такт.

3.3.3.2 Оценки по WLS

Отново, за да се намерят оптималните параметри, се определят производните на $\mathcal{F}(\theta)$. При диференциране на целевата функция по вектора $\theta \in \mathcal{R}^p$ се получава градиентът $g \in \mathcal{R}^p$. Извеждането на оптималното решение по WLS се извършва по същите стъпки като при LS. Първо $\mathcal{F}(\theta)$ се представя като функция на θ по следния начин:

$$\begin{aligned} \mathcal{F}(\theta) &= e^T W e \\ &= (y - \Phi\theta)^T W (y - \Phi\theta) \\ &= y^T W y - \theta^T \Phi^T W y - y^T W \Phi\theta + \theta^T \Phi^T W \Phi\theta. \end{aligned}$$

Тъй като W е симетрична, в горните равенства се използва, че $W = W^T$. След диференциране се получават следните зависимости:

$$\begin{aligned} \nabla_\theta(y^T W y) &= 0, \\ \nabla_\theta(\theta^T \Phi^T W y) = \nabla_\theta(y^T W \Phi\theta) &= \Phi^T W y, \\ \nabla_\theta(\theta^T \Phi^T W \Phi\theta) &= 2\Phi^T W \Phi\theta, \end{aligned}$$

откъдето за градиента се получава:

$$g = -2\Phi^T W y + 2\Phi^T W \Phi\theta.$$

Той се приравнява на нула и се определя $\hat{\theta}$, т.е.

$$g|_{\theta=\hat{\theta}} = 0 \Leftrightarrow -\Phi^T W y + \Phi^T W \Phi\hat{\theta} = 0.$$

Векторът $\hat{\theta}$, при който g се нулира, е

$$\hat{\theta} = (\Phi^T W \Phi)^{-1} \Phi^T W y.$$

3.3.4 Метод на обобщените най-малки квадрати

Тегловната матрица W в WLS е диагонална. Но това приемане се нарушава, когато остатъкът е цветен шум и $W = \Sigma_e^{-1}$. Отново, тъй като този метод е разгледан за модел с матрица на параметрите, акцентът в долното разглеждане е върху особеностите на GLS, за случая с вектор на параметрите.

3.3.4.1 Показател на качеството

Ако между елементите на вектора e има корелация (причините за това са споменати в точка 3.2.4), то матрицата $\Sigma_e \in \mathcal{R}^{(N-n)\ell \times (N-n)\ell}$ не е диагонална. Тогава, ако с цел получаване на ефективни и неизместени оценки, се зададе $W = \Sigma_e^{-1}$, то в целевата функция се появяват (претеглени) взаимни произведения от вида $e_{i,k_1} e_{i,k_2}$ за $k_1 \neq k_2$. Въпреки това не се променя начинът, по който се формира целевата функция, нито се изменя изразът на оптималното решение. Целевата функция е

$$\mathcal{F}(\theta) = e^T W e.$$

3.3.4.2 Оценки по GLS

Извеждането на оптималните оценки се извършва по същите стъпки, както при WLS, а оценките, минимизиращи $\mathcal{F}(\theta)$, са:

$$\hat{\theta} = (\Phi^T W \Phi)^{-1} \Phi^T W y.$$

Обикновено Σ_e не е известна и затова GLS (по подобие на модела с матрица на параметрите) се реализира като итеративна процедура. Тъй като остатъкът е корелиран с предикторите, е подходящо да се въведе формиращ филтър от тип AR.

3.3.5 Метод на разширените най-малки квадрати

Когато не е намерен подходящ ARX модел, част от поведението на системата се причислява към остатъка, който няма свойствата на бял шум (отново ще се бележи с $\tilde{e}_k$), а това води до изместеност на оценките. В тази точка е разгледан вариантът с въвеждане на формиращ филтър за избелване на остатъка. В точка 3.2.5.2 беше споменато, че един вариант за оценяване на параметрите на разширения с филтъра модел (когато целевата функция е сума от квадратите на остатъка) е да се приложи NLS. Това е свързано с прилагане на числен метод за нелинейна оптимизация. Другият вариант е да се използва итеративната

процедура на ELS. В тази точка е разгледана реализацията на ELS за модел, представен в общ вид с вектор на параметрите.

Идеята и особеностите на ELS за модела с вектор на параметрите са същите като при оценяването на ARMAX модел, представен с матрица на параметрите. Затова в долното изложение вниманието е насочено към описанието на системата с ARMAX модел и към съответния алгоритъм.

3.3.5.1 Описание на системата

Когато цветният шум $\tilde{e}_k$ се представи с помощта на формиращ филтър от тип МА:

$$\tilde{e}_k = C(q^{-1})e_k,$$

на входа на който постъпва (фиктивен) случаен сигнал e_k , който е бял шум, полученият разширен модел е ARMAX:

$$A(q^{-1})y_k = B(q^{-1})u_k + C(q^{-1})e_k.$$

Преобразуването на модела в общ вид с вектор на параметрите е извършено по подобие на преобразуването на ARX модела. Горният израз представлява система от ℓ уравнения (MISO модели), отговарящи на всеки от изходите на системата. i -тият ред (MISO модел) е

$$A_i(q^{-1})y_k = B_i(q^{-1})u_k + C_i(q^{-1})e_k.$$

Представен по-подробно, моделът е

$$\begin{aligned} a_{i1}(q^{-1})y_{1,k} + \dots + a_{i\ell}(q^{-1})y_{\ell,k} &= b_{i1}(q^{-1})u_{1,k} + \dots + b_{im}(q^{-1})u_{m,k} \\ &+ c_{i1}(q^{-1})e_{1,k} + \dots + c_{i\ell}(q^{-1})e_{\ell,k}. \end{aligned} \quad (3.56)$$

Целта на долните преобразувания е всички параметри да се отделят в един вектор, а регресорите – в матрица.

Нека за опростяване на записите индексът i се изпуска. Събираемите в горния израз може да се представят в следния векторен вид:

$$a_j(q^{-1})y_{j,k} = \begin{cases} y_{i,k} + y_{j,k}^T a_j, & \text{за } i = j, \\ y_{j,k}^T a_j, & \text{за } i \neq j, \end{cases}$$

$$b_j(q^{-1})u_{j,k} = u_{j,k}^T b_j$$

$$c_j(q^{-1})e_{j,k} = \begin{cases} e_{i,k} + e_{j,k}^T c_j, & \text{за } i = j, \\ e_{j,k}^T c_j, & \text{за } i \neq j. \end{cases}$$

Въведените вектори на регресорите са:

$$y_{j,k} = [y_{j,k-1} \ y_{j,k-2} \ \dots \ y_{j,k-na_j}]^T \in \mathcal{R}^{na_j},$$

$$u_{j,k} = [u_{j,k-1} \ u_{j,k-2} \ \dots \ u_{j,k-nb_j}]^T \in \mathcal{R}^{nb_j},$$

$$e_{j,k} = [e_{j,k-1} \ e_{j,k-2} \ \dots \ e_{j,k-nc_j}]^T \in \mathcal{R}^{nc_j},$$

а векторите на параметрите са:

$$a_{j,k} = [a_{j,1} \ a_{j,2} \ \dots \ a_{j,na_j}]^T \in \mathcal{R}^{na_j},$$

$$b_{j,k} = [b_{j,1} \ b_{j,2} \ \dots \ b_{j,nb_j}]^T \in \mathcal{R}^{nb_j},$$

$$c_{j,k} = [c_{j,1} \ c_{j,2} \ \dots \ c_{j,nc_j}]^T \in \mathcal{R}^{nc_j}.$$

Както се вижда, векторите $y_{j,k}$, $u_{j,k}$ и $e_{j,k}$ съдържат предисторията на съответните сигнали до $k-1$ -ия такт, а единствените стойности в k -тия такт са $y_{j,k}$ и $e_{j,k}$. Това позволява от (3.56) да се изрази $y_{i,k}$ като функция на предисторията на входовете и изходите, както и на остатъка в k -тия такт, т.е.

$$y_{i,k} = -y_{1,k}^T a_1 - \dots - y_{\ell,k}^T a_{\ell} + u_{i,k}^T b_i + \dots + u_{m,k}^T b_m + e_{1,k}^T c_1 + \dots + e_{\ell,k}^T c_{\ell} + e_{i,k}$$

или накратко (с въвеждане отново на индекса i)

$$y_{i,k} = \varphi_{\text{ARX}i,k}^T \theta_{\text{ARX}i} + \varphi_{\text{fi}i,k}^T \theta_{\text{fi}i} + e_{i,k}. \quad (3.57)$$

Параметрите на i -тия MISO модел се групират в двата вектора:

$$\theta_{\text{ARX}i} = [a_{i1}^T \ \dots \ a_{i\ell}^T \ b_{i1}^T \ \dots \ b_{im}^T]^T,$$

$$\theta_{\text{fi}i} = [c_{i1}^T \ \dots \ c_{i\ell}^T]^T,$$

а регресорите му са обединени в двата вектора:

$$\varphi_{\text{ARX}i,k} = [-y_{i1,k}^T \ \dots \ -y_{i\ell,k}^T \ -u_{i1,k}^T \ \dots \ -u_{im,k}^T]^T,$$

$$\varphi_{\text{fi}i,k} = [e_{i1,k}^T \ \dots \ e_{i\ell,k}^T]^T.$$

Моделите (3.57) за $i = \overline{1, \ell}$, записани заедно, са:

$$y_k = \Phi_k \theta + e_k,$$

където векторът $\theta \in \mathcal{R}^p$ съдържа параметрите на всички MISO модели. Той има размерност $p = \sum_i p_i$ и е със следната структура:

$$\theta = [\theta_{\text{ARX}}^T \ \theta_{\text{f}}^T]^T.$$

Векторът

$$\theta_{\text{ARX}} = [\theta_{\text{ARX1}}^T \ \theta_{\text{ARX2}}^T \ \dots \ \theta_{\text{ARX}\ell}^T]^T$$

съдържа всички параметри на полиномите в матриците $A(q^{-1})$ и $B(q^{-1})$ (ARX частта на разширения модел), а векторът

$$\theta_f = [\theta_{f1}^T \ \theta_{f2}^T \ \dots \ \theta_{f\ell}^T]^T$$

съдържа параметрите на полиномите в матрицата $C(q^{-1})$ (частта на филтъра от тип МА). При това разделяне на параметрите матрицата на регресорите на пълния МИМО ARMAX модел $\Phi_k \in \mathcal{R}^{\ell \times p}$ се състои от два блока, всеки от които е блокдиагонален. Нека двата блока са:

$$\Phi_{\text{ARX},k} = \text{diag}(\varphi_{\text{ARX1},k}^T, \varphi_{\text{ARX2},k}^T, \dots, \varphi_{\text{ARX}\ell,k}^T),$$

$$\Phi_{f,k} = \text{diag}(\varphi_{f1,k}^T, \varphi_{f2,k}^T, \dots, \varphi_{f\ell,k}^T),$$

като $\Phi_{\text{ARX},k}$ съдържа предисторията на (известните) входно-изходни сигнали, а $\Phi_{f,k}$ – стойностите на (неизвестния) остатък на ARMAX модела. Тогава пълната матрица на регресорите, участващи във формирането на изхода в k -тия такт, е

$$\Phi_k = [\Phi_{\text{ARX},k} \ \Phi_{f,k}]. \quad (3.58)$$

Описанието на модела за интервала на наблюдение е

$$y = \Phi\theta + e. \quad (3.59)$$

Векторите $y, e \in \mathcal{R}^{\ell(N-n)}$ и матрицата $\Phi \in \mathcal{R}^{\ell(N-n) \times p}$ в горното равенство са:

$$y = [y_{n+1}^T \ y_{n+2}^T \ \dots \ y_N^T]^T,$$

$$e = [e_{n+1}^T \ e_{n+2}^T \ \dots \ e_N^T]^T,$$

$$\Phi = [\Phi_{n+1}^T \ \Phi_{n+2}^T \ \dots \ \Phi_N^T]^T,$$

за $n = \max(na, nb, nc)$. С използване на разделното представяне на ARX и на МА частта, (3.59) може да се запише като:

$$y = [\Phi_{\text{ARX}} \ \Phi_f] \begin{bmatrix} \theta_{\text{ARX}} \\ \theta_f \end{bmatrix} + e.$$

3.3.5.2 Показател на качеството

Тъй като ELS е свързан с многократно прилагане на LS, показателят на качеството и критерият са същите като тези на стандартния LS. За модел с вектор на параметрите показателят може да се запише като:

$$\mathcal{F}(\theta) = e^T e$$

и съответно критерият е

$$\min_{\theta} \mathcal{F}(\theta) = \min_{\theta} e^T e.$$

3.3.5.3 Алгоритъм на метода

Методът ELS за оценка на MIMO модели с вектор на параметрите се състои от следните стъпки.

1. Инициализация

Задават се параметрите, които определят условията за спиране на итеративната процедура. Формира се матрицата Φ_{ARX} , отговаряща на избраната структура на ARX модела и се оценяват параметрите му:

$$\theta_{\text{ARX}}^{(0)} = (\Phi_{\text{ARX}}^T \Phi_{\text{ARX}})^{-1} \Phi_{\text{ARX}}^T y.$$

$\theta_{\text{ARX}}^{(0)}$ е изместена оценка поради цветния характер на остатъка $\tilde{e}_k$. Формира се векторът e , съдържащ оценените стойности на остатъка e_k , който трябва да е БШ. В инициализиращата стъпка e съвпада с вектора $\tilde{e}_k$, зависещ от цветните остатъци, т.е.

$$e^{(0)} = y - \Phi_{\text{ARX}} \theta_{\text{ARX}}^{(0)} \quad (= \tilde{e}^{(0)}). \quad (3.60)$$

2. LS за ARMAX модела

В i -тата итерация се формира вторият блок $\Phi_f^{(i)}$ в (3.58), както и пълната матрица:

$$\Phi^{(i)} = [\Phi_{\text{ARX}} \quad \Phi_f^{(i)}].$$

Векторът на параметрите на ARMAX модела в текущата итерация се формира като:

$$\theta^{(i)} = (\Phi^{(i)T} \Phi^{(i)})^{-1} \Phi^{(i)T} y.$$

3. Оценяване на остатъка e_k

Векторът e в текущата итерация се изчислява като:

$$e^{(i)} = y - \Phi^{(i)} \theta^{(i)}. \quad (3.61)$$

4. Проверка за спиране

Ако избраното условие за спиране на итеративната процедура е изпълнено, оценяването на параметрите на ARMAX модела се преустановява и се приема, че търсените оценки са $\hat{\theta} = \theta^{(i)}$. В противен случай $i \leftarrow i + 1$ и се продължава от стъпка 2.

3.3.6 Метод на инструменталните променливи

Методът на инструменталните променливи (IV) се използва, когато остатъкът не е бял шум, което води до изместени оценки. При него неизместеността се постига, като се изменя уравнението на модела. По-долу е представен IV за модел с вектор на параметрите (методът е разгледан по-подробно за модел с матрица на параметрите – точка 3.2.6).

3.3.6.1 Уравнение на модела

Отново с $\tilde{e}_k$ е означен оцветеният остатък, а с $\tilde{e}$ – векторът, съдържащ стойностите му по отношение на всички изходи за интервала на наблюдение. При IV моделът

$$y = \Phi\theta + \tilde{e} \quad (3.62)$$

се умножава с инструменталната матрица $Z \in \mathcal{R}^{(N-n)\ell \times p}$ и тогава

$$Z^T y = Z^T \Phi\theta + Z^T \tilde{e}.$$

Тази матрица се определя така, че стълбовете ѝ – инструменталните променливи, да са корелирани с данните, т.е. $Z^T y \neq 0$ и $Z^T \Phi \neq 0$, но да не са корелирани с остатъка. Когато остатъкът е центриран, $Z^T \tilde{e} = 0$ и тогава (3.62) се свежда до

$$Z^T y = Z^T \Phi \hat{\theta},$$

откъдето за оценките се получава:

$$\hat{\theta} = (Z^T \Phi)^{-1} Z^T y.$$

3.3.6.2 Алгоритъм на IV

Предисторията на входа не зависи от остатъка и затова е удачно входните сигнали за ARX модела да играят роля на инструментални променливи. Но тъй като y зависи от $\tilde{e}$, изходите не са подходящи за стълбове на Z . В долния алгоритъм за изграждане на Z се използва модел за формиране на сигнал, който е приближение на y , но не зависи от $\tilde{e}$. Тази идея е залегнала в следния алгоритъм.

Алгоритъм на IV

1. Инициализация

Задават се условията за спиране, формира се Φ с входно-изходните данни и се прилага LS:

$$\theta^{(0)} = (\Phi^T \Phi)^{-1} \Phi^T y.$$

2. Оценяване на y

В i -тата итерация се формира оценка на y , чийто елементи са:

$$y_k^{(i)} = \Phi'_k \theta^{(i)}.$$

В матрицата на регресорите Φ'_k , вместо изхода на системата, участва предисторията на неговата оценка.

3. Формиране на Z

$$Z^{(i)} = [\Phi_{n+1}'^T \quad \Phi_{n+2}'^T \quad \dots \quad \Phi_N'^T]^T.$$

4. Оценяване на θ по IV

$$\theta^{(i)} = (Z^{(i)T} \Phi)^{-1} Z^{(i)T} y.$$

5. Проверка за спиране

Ако избраното условие за спиране е изпълнено, итеративното оценяване се прекратява и се приема, че търсените параметри са $\hat{\theta} = \theta^{(i)}$. В противен случай $i \leftarrow i + 1$ и се продължава от стъпка 2.

При описанието на IV за модел с матрица на параметрите е засегнат въпросът за избора на Z , както и начини за подобрене на описания алгоритъм.

3.3.7 Робастен метод на най-малките квадрати

Идеята на метода и приложението му са разгледани в точка 3.2.7. Оценителите, използващи тази идея, са подходящи, когато разпределението на остатъка води до изместени оценки по LS или по разгледаните до момента негови модификации. Робастните оценители се изграждат така, че оценките да са нечувствителни към асиметриите и особено към евентуалните тежки опашки в разпределението на e .

3.3.7.1 Показател на качеството

Най-общо целевата функция на робастните оценители е

$$\mathcal{F}(\theta) = \sum_{i=1}^{\ell} \sum_{k=n+1}^N \zeta(e_{i,k}), \quad (3.63)$$

а критерият е минимум на $\mathcal{F}(\theta)$ по отношение на θ . Функцията $\zeta_{i,k} = \zeta(e_{i,k})$ е приносът на остатъка $e_{i,k}$ към стойността на $\mathcal{F}(\theta)$. $\zeta(\cdot)$ трябва да е неотрицателна (по същата причина като теглата в WLS), да е четна и унимодална с минимум $\zeta(0) = 0$.

За удобство при извеждането на RobLS е прието, че приносите в k -тия такт и техните производни по съответните остатъци са обединени във вектора:

$$\zeta_k = [\zeta_{1,k} \ \zeta_{2,k} \ \dots \ \zeta_{\ell,k}]^T, \quad \gamma_k = [\gamma_{1,k} \ \gamma_{2,k} \ \dots \ \gamma_{\ell,k}]^T.$$

3.3.7.2 Оценки по RobLS

Нека $\theta_i \in \mathcal{R}^{p_i}$ е векторът на параметрите на i -тия MISO модел, а j -тият елемент от този вектор е $\theta_h = [\theta_i]_j$. На него в общото представяне на модела отговаря $\varphi_h = [\varphi_{i,k}]_j$, където $\varphi_{i,k} \in \mathcal{R}^{p_i}$ е съответстващият на θ_i вектор на регресорите. От всички остатъци в k -тия такт само $e_{i,k}$ (и съответно $\zeta_{i,k}$) зависи от елементите на вектора θ_i (т.е. и от θ_h). Тогава ненулевите производни ζ_k по отношение на θ_h са единствено $\frac{\partial \zeta_{i,k}}{\partial \theta_h}$, които са:

$$\frac{\partial \zeta_{i,k}}{\partial \theta_h} = \frac{\partial \zeta_{i,k}}{\partial e_{i,k}} \frac{\partial e_{i,k}}{\partial \theta_h} = -\gamma_{i,k} \varphi_{h,k}. \quad (3.64)$$

Обобщавайки за $h = \overline{1, p_i}$ (т.е. за i -тия MISO модел), може да се запише:

$$\nabla_{\theta_i} \zeta_{i,k} = -\varphi_{i,k} \gamma_{i,k}.$$

За пълния градиент на приноса ζ_k (обобщавайки за $i = \overline{1, \ell}$) се получава:

$$\nabla_{\theta} \zeta_k = \begin{bmatrix} \nabla_{\theta_1} \zeta_{1,k} \\ \nabla_{\theta_2} \zeta_{2,k} \\ \vdots \\ \nabla_{\theta_{\ell}} \zeta_{\ell,k} \end{bmatrix} = - \begin{bmatrix} \varphi_{1,k} & 0 & \dots & 0 \\ 0 & \varphi_{2,k} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \varphi_{\ell,k} \end{bmatrix} = -\Phi_k^T \gamma_k.$$

Тогава градиентът на целевата функция е

$$g = \nabla_{\theta} \mathcal{F}(\theta) = \sum_{k=n+1}^N g_k = - \sum_{k=n+1}^N \Phi_k^T \gamma_k. \quad (3.65)$$

Нека се въведат отношенията:

$$w_{i,k} = \frac{\gamma_{i,k}}{e_{i,k}}. \quad (3.66)$$

Тогава производната на вектора на приносите в k -тия такт може да се запише като:

$$\gamma_k = \text{diag}(w_k) e_k.$$

Тук отношенията $w_{i,k}$ са обединени във вектора:

$$w_k = [w_{1,k} \quad \dots \quad w_{\ell,k}]^T, \quad (3.67)$$

а (3.65) добива вида:

$$g = - \sum_{k=n+1}^N \Phi_k^T \text{diag}(w_k) e_k.$$

Последният израз в компактен вид е

$$g = -\Phi^T \underbrace{W e}_{\gamma}, \quad (3.68)$$

където

$$e = [e_{n+1}^T \quad \dots \quad e_N^T]^T, \quad \Phi = [\Phi_{n+1}^T \quad \dots \quad \Phi_N^T]^T,$$

$$W = \text{diag}[w_{n+1}^T \quad \dots \quad w_N^T]. \quad (3.69)$$

Видът (3.68) на градиента на целевата функция съвпада с градиента по WLS. Наистина изразът за g от точка 3.3.3.2 е

$$g = -2\Phi^T W y + 2\Phi^T W \Phi \theta = -2\Phi^T W e,$$

като в (3.68) множителят 2 се причислява към елементите на W . Така минимизацията на $\mathcal{F}(\theta)$, формирана по (3.63), води до

$$\hat{\theta} = (\Phi^T W \Phi)^{-1} \Phi^T W y.$$

Разсъжденията за полученото решение са същите, а именно: тегловната матрица зависи от остатъка, който зависи от оценките, а те от своя страна зависят от теглата. Затова робастните оценители се реализират

като итеративни процедури. На всяка итерация се формират теглата, зависещи от остатъците от предишната итерация, и се прилага стандартният WLS. Получените оценки се използват за изчисляване на остатъците, които в следващата итерация се използват при формирането на нови тегла. Такъв алгоритъм е описан по-долу.

3.3.7.3 Алгоритъм на робастния LS

1. Инициализация

Задават се условията за спиране, например толеранси на изменение на $\mathcal{F}(\theta)$ или на оценките в съседни итерации. Също така се формира Φ и се прилага LS:

$$\theta^{(0)} = (\Phi^T \Phi)^{-1} \Phi^T y.$$

Това е равносилно на WLS с тегловна матрица $W^{(0)} = I$.

2. Определяне на остатъка

В i -тата итерация се формира векторът $e^{(i)}$:

$$e^{(i)} = y - \Phi \theta^{(i-1)}.$$

3. Формиране на тегловната матрица

Изчисляват се новите тегла:

$$w_{j,k}^{(i)} = \min(1, \max(\delta, |e_{j,k}^{(i)}|^{-1})), \quad (3.70)$$

като δ е малко положително число, например $\delta = 10^{-8}$. Те са необходими за формиране на новата тегловна матрица $W^{(i)}$ според (3.67) и (3.69).

4. Оценяване на параметрите по WLS

Параметрите на модела са:

$$\theta^{(i)} = (\Phi^T W^{(i)} \Phi)^{-1} \Phi^T W^{(i)} y.$$

5. Проверка за спиране

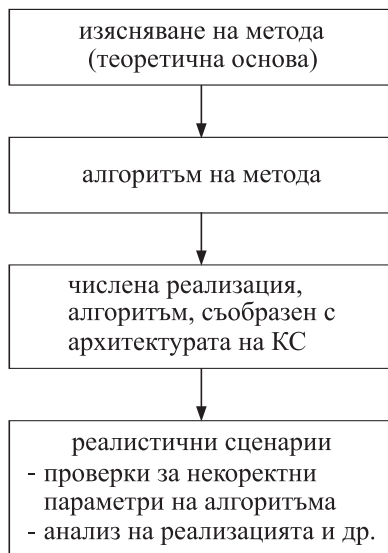
Ако условието за спиране е изпълнено, итеративното оценяване се преустановява и се приема, че търсените параметри са $\hat{\theta} = \theta^{(i)}$. В противен случай $i \leftarrow i + 1$ и се продължава от стъпка 2.

3.4 Числено устойчиви реализации на оценителите

При системи с голям брой входове и изходи често има стремеж идентификацията да се извършва, доколкото е възможно, без човешка намеса. Важен момент при проектирането на оценители на параметрите (Фигура 3.9), особено когато идентификацията е автоматизирана, е прилагането на числено устойчиви реализации на методите за оценяване.

По-долу са разгледани числените проблеми, техните източници, както и различни подходи за избягването им по време на работа на оценителите.

Описаните до момента методи се свеждат до еднократно или многократно прилагане на LS или на претегления му вариант WLS. Затова в тази тема ще бъде засегнато оценяването на параметри по LS.



Фигура 3.9. Етапи в проектирането на оценители на параметри (КС – компютърна система)

3.4.1 Числени проблеми при оценяването на параметри

Най-общо числените проблеми по време на идентификацията възникват по две причини. Едната е, че този подход за моделиране се реализира с помощта на компютри, което е свързано с грешки от изчисления [25]. Другата причина е използването на неподходящи данни.

При реализацията на методите за оценяване, независимо от данните, съществуват два вида грешки от изчисления. Това са грешки от:

- прекъсване;
- закръгляне.

3.4.1.1 Грешка от прекъсване

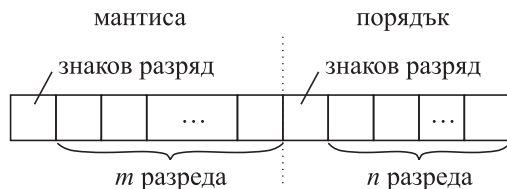
Тази грешка е характерна за операциите, където итеративно се търси оптимално решение. Най-общо, поради невъзможността числено да се сумират безкраен брой членове на (сходящи) редове, сумирането се извършва само между определен брой членове, пренебрегвайки останалите. Това води до т.нар. грешка от прекъсване. Тя се контролира с различни критерии за спиране на процеса на оптимизация.

3.4.1.2 Грешка от закръгляне

Всяко число x може да се представи във вида:

$$x = a_1 b^n + a_2 b^{n-1} + \dots + a_m b^{n-m+1},$$

като b е основата на бройната (позиционна) система, a_i са значещите цифри, а цялото число m определя броя на разрядите (значещите цифри), с които е записано числото. Например дробното число 12.35 може да се представи като $12.35 = 1 \times 10^1 + 2 \times 10^0 + 3 \times 10^{-1} + 5 \times 10^{-2}$.



Фигура 3.10. Представяне на числата с плаваща запетая в компютърните системи

В компютрите числата се представят с мантиса и порядък (Фигура 3.10). Мантисата съдържа значещите цифри, а порядъкът определя по-

зицията на плаващата запетая. Прието е в мантисата да е записана само дробната част, т.е. плаващата запетая да се намира пред най-старшия разряд на мантисата, който е различен от нула ($a_1 \neq 0$) (в случая 12.35 се представя като 0.1235×10^2).

В резултат от аритметични действия с числа, записани с m разряда може да се получат числа, които са с повече от m разряда. Така, ако числата в компютъра се представят с m цифри, то вместо с точния резултат се работи с негово приближение. Грешката, която се допуска в този случай, е грешка от закръгляне.

Пример. *Загуба на точност при аритметично действие*

Въпреки че числата в компютърните системи се представят в двоичен вид, за представяне на загубата на точност, се използват числа в десетичната бройна система. Числата 3 и 4 се записват само с една цифра. Въпреки това частното им е

$$x = \frac{4}{3} = 1.333 \dots = 1.3(3).$$

Ако x се представи с четири значещи цифри и порядък, то

$$x = 0.1333 \times 10^1.$$

С увеличаване на значещите цифри се получава все по-близка стойност до истинската (при пет значещи цифри $x = 0.13333 \times 10^1$ и т.н.), но никога истинската стойност не може да се запише с краен брой цифри. С други думи при записване на x винаги ще съществува грешка от закръгляне. $\diamond$

Ако данните, използвани от оценителя, са подходящи за обработка, грешките от изчисления не влияят чувствително на крайния модел. Но ако са неподходящи за задачата на оценяването (и не са коректно обработени), числените грешки може да доведат до значително отклонение на поведението на модела от това на реалната система.

Основният източник на числени грешки при реализацията на LS (както и по-общите случаи, разгледани до момента) е обръщането на матрицата на Фишер $\Phi^T \Phi$. Типични случаи, при които е възможно да възникне числен проблем, са:

- мултиколинеарност между факторите (източникът е системата и по-точно нейните свойства по отношение на величините, използвани като фактори);
- недостатъчно възбуден обект (източникът е експериментът);

- неподходящо подбрана структура на модела (източникът е моделът, например статична връзка между величини, която се описва с динамичен модел).

При едномерните системи, за да се гарантира числено устойчиво оценяване на параметрите, обикновено е достатъчно експериментът да е проведен коректно (най-вече тактът на дискретизация да е подходящ и обектът да е възбуден), а също и структурата на модела да е подходяща. При многомерните системи обаче е възможно тези условия да са изпълнени и въпреки това да възникнат числени проблеми. Тъй като при ММО системите u_k и y_k са векторни сигнали, е възможно някои от компонентите им да са (близки до) линейно зависими.

3.4.2 Линейно зависими фактори

Със следващия пример е показан проблемът, който може да възникне, когато по една или друга причина факторите в модела са линейно зависими.

Пример. *Модел с линейно зависими фактори*

Нека е даден MISO модел, зависещ от два фактора, и нека търсените параметри са $\theta^* = [\theta_1^* \ \theta_2^*]^T$. Оптималният модел е

$$\hat{y}_k = \theta_1^* \varphi_{1,k} + \theta_2^* \varphi_{2,k}. \quad (3.71)$$

Ако $\varphi_{1,k}$ и $\varphi_{2,k}$ са линейно зависими, като връзката е

$$\varphi_{1,k} = 2\varphi_{2,k},$$

тогава модел (3.71) е еквивалентен на:

$$\hat{y}_k = \theta_1^* \varphi_{1,k} + \theta_2^* \varphi_{2,k} + \underbrace{\alpha \varphi_{1,k} - 2\alpha \varphi_{2,k}}_{=0},$$

за произволно α . Или с други думи, съществуват безброй много модели от вида:

$$\hat{y}_k = \underbrace{(\theta_1^* + \alpha)}_{\tilde{\theta}_1} \varphi_{1,k} + \underbrace{(\theta_2^* - 2\alpha)}_{\tilde{\theta}_2} \varphi_{2,k} = \varphi_k^T \tilde{\theta}, \quad (3.72)$$

които теоретично съвпадат с (3.71). Например за $\alpha = 10^5$ и $\theta^* = [0.6 \ 0.4]^T$, моделите:

$$\hat{y}_k = 0.6\varphi_{1,k} + 0.4\varphi_{2,k}, \quad (3.73)$$

$$\hat{y}_k = 100000.6\varphi_{1,k} - 199999.6\varphi_{2,k}, \quad (3.74)$$

са идентични. Това показва, че α може да нараства неограничено, което води до нарастване на оценките по абсолютна стойност. Така при практическата реализация на оценителя в среда с крайна точност изходите може да се различават, което се дължи на грешките от изчисления (по-точно това са грешките от закръгляне). $\diamond$

Ако част от факторите, а те са стълбове на Φ , са (близки до) линейно зависими, матрицата на Фишер $\Phi^T \Phi$ е лошо обусловена. В такъв случай определянето на $(\Phi^T \Phi)^{-1}$ е свързано със значителни грешки в изчисленията, които обикновено водят до модел, който е напълно неизползваем.

От гледна точка на геометричната интерпретация на LS (Фигура 2.40) линейната зависимост между факторите означава, че направленията на част от стълбовете на Φ съвпадат, а ако факторите не са напълно, но близки до линейно зависими, направленията също са много близки. Това показва, че част от факторите в модела, освен че влошават оценките, са ненужни. Изборът на значими за описанието на изхода фактори е разгледан в точка 3.5, а целта на долните изменения на решението по LS е да се гарантират числено устойчиви оценки на параметрите, независимо от вида на данните. Причина за това е, че в практиката, при недостиг на фактори, може да се наложи да се използват и величини, между които съществува известна мултиколинеарност.

За разширяване на случая от горния пример нека $\text{rank } \Phi = r < p$, като стълбовете Φ са пренаредени така, че $\Phi = [\Phi_1 \ \Phi_2]$, където $\Phi_1 \in \mathcal{R}^{(N-n)\ell \times r}$ е от пълен ранг (т.е. не съдържа линейно зависими фактори), а $\Phi_2 \in \mathcal{R}^{(N-n)\ell \times p-r}$ съдържа останалите стълбове, които са линейно зависими със стълбове във Φ_1 . Аналогично на Φ , векторът на параметрите се размества така, че първите r елемента да съответстват на Φ_1 , а останалите $p - r$ реда – на Φ_2 , т.е. $\theta = [\theta_1^T \ \theta_2^T]^T$, $\theta_1 \in \mathcal{R}^r$ и $\theta_2 \in \mathcal{R}^{p-r}$.

Тъй като $\text{rank } \Phi = \text{rank } \Phi_1 = r$, то стълбовете на Φ_2 може да се представят като линейна комбинация от стълбовете на Φ_1 , т.е.

$$\Phi_2 = \Phi_1 \Theta,$$

като $\Theta \in \mathcal{R}^{p-r}$ е ненулева матрица. Тогава целевата функция може да се развие по следния начин:

$$\begin{aligned} \mathcal{F}(\theta) &= \|y - \Phi\theta\|_2^2 \\ &= \|y - \Phi_1\theta_1 - \Phi_2\theta_2\|_2^2 \\ &= \|y - \Phi_1(\theta_1 + \Theta\theta_2)\|_2^2 \\ &= \|Y - \Phi_1\theta_3\|_2^2. \end{aligned}$$

Понеже Φ_1 е от пълен ранг, то за оптималното решение се получава:

$$\hat{\theta}_3 = (\Phi_1^T \Phi_1)^{-1} \Phi_1^T y.$$

При тази формулировка на задачата за конкретен избор на $\hat{\theta}_1$ може да се намери такъв вектор $\hat{\theta}_2$, че да е изпълнено равенството $\hat{\theta}_3 = \hat{\theta}_1 + \Theta \hat{\theta}_2$. Това означава, че част от параметрите се оценяват еднозначно (елементите на θ_1), а останалите (θ_2) са свободни параметри. С други думи задачата за минимизация на $\mathcal{F}(\theta)$ има безброй много решения.

Нека Ω е множеството от оценки $\tilde{\theta}$, за които целевата функция има минимална стойност. В такъв случай, за да се избере конкретно решение, може да се наложат допълнителни изисквания. Например подходящо е 2-нормата на $\tilde{\theta}$ да е минимална [102, 148]. Тогава оптималното решение е

$$\hat{\theta} = \arg \min_{\tilde{\theta} \in \Omega} \|\tilde{\theta}\|_2^2. \quad (3.75)$$

Така се избягва опасността от неограничено нарастване на оценките по абсолютна стойност. За подобряване на резултата от оценяването се прилагат числено устойчиви реализации на стандартните методи, с извършване на регуляризации [35] или декомпозиции [103, 151]. Термините факторизация и декомпозиция, които се срещат в литературата, са взаимозаменяеми в разглеждания по-долу. Тук те имат смисъл на преобразуване на матрица в произведение от матрици (фактори) [103]. Въпреки това терминът декомпозиция е по-общ и се използва също за преобразувания на матрици, при които не е задължително факторите да се умножават. В литературата също така се среща и терминът триангулизация [59], използван за означаване на преобразуване, при което даден фактор е триъгълна матрица. За преобразуванията на матрици в долните разглеждания се използва терминът декомпозиция, а не по-точният термин факторизация, както и резултантните матрици не се наричат фактори, за да не се смесват понятията с използваните до момента.

По-долу реализациите на LS са представени подробно за модел с вектор на параметрите, след което е представен вариантът за модел с матрица на параметрите.

3.4.3 LS с регуляризация на Тихонов

3.4.3.1 Идея на реализацията

В оптимизационните процедури е възможно промяната в $\hat{y}_k$, дължаща се на числени грешки, да води до намаляване на целевата функция. При това оценителят, минимизирайки $\mathcal{F}(\theta)$, погрешно би изменял оценките на параметрите в посока, водеща до нарастване на числените грешки. Така оценките, свързани с линейно зависимите фактори, често нарастват с всяка итерация, при това неограничено. Такъв процес на оценяване е разходящ.

Ако се очаква поява на числени проблеми по време на оценяването на параметри, един начин за получаване на модел, нечувствителен към споменатите грешки, е да се наложи изискването 2-нормата на вектора на параметрите $\|\theta\|_2$, т.е. дължината на θ , да е възможно по-малка [126]. Това е идеята на регуляризацията на Тихонов (по-точно в нейната стандартна форма), използвана и в метода на Левенберг-Маркуард (Levenberg-Marquardt). Например за модел (3.73) $\|\theta\|_2 \approx 0.72$, а за модел (3.74) нормата е $\|\theta\|_2 \approx 2.24 \times 10^5$. Така измежду моделите (3.72) се избира такова описание на системата, при което стойността на α осигурява минимална 2-норма на параметрите. Подробно сравнение между регуляризацията на Тихонов и метода на Левенберг-Маркуард е извършено в [89].

3.4.3.2 Показател на качеството

За модел с вектор на параметрите, нормата е

$$\|\theta\|_2 = \sqrt{\theta^T \theta}.$$

Тогава ограничаването на нормата на θ може да се постигне, ако в целевата функция се въведе наказателен член, който нараства с увеличаване на $\|\theta\|_2^2$. Така $\mathcal{F}(\theta)$ добива вида:

$$\mathcal{F}(\theta) = e^T e + \lambda \theta^T \theta, \quad (3.76)$$

като $\lambda > 0$ е тегловен коефициент, който определя с каква тежест се предпочита оценките да минимизират нормата на вектора на параметрите спрямо минимизацията на квадрата на остатъка. Минимизацията на (3.76) води до регуляризацията на Тихонов.

3.4.3.3 Оценки по LS с регуляризация на Тихонов

Градиентът на целевата функция (3.76) е

$$\begin{aligned} g &= \nabla_{\theta} ((y - \Phi\theta)^T(y - \Phi\theta)) + \nabla_{\theta}(\lambda\theta^T\theta) \\ &= -2\Phi^T y + 2\Phi^T \Phi\theta + 2\lambda\theta. \end{aligned}$$

Приравнявайки g на нула за оптималните регуляризирани оценки по Тихонов, се получава:

$$\hat{\theta} = (\Phi^T \Phi + \lambda I)^{-1} \Phi^T y.$$

Въвеждането на λ води до подобряване на обусловеността на матрицата, която се обръща, но също и до изместеност на $\hat{\theta}$. Колкото по-голяма е стойността на λ , толкова по-изместени са оценките, но $\hat{\theta}$ е с по-малка норма и числените грешки намаляват. От друга страна, когато $\lambda \rightarrow 0$, изместеността намалява, но се усилват числените грешки. По тази причина изборът на λ е от основно значение за коректното получаване на оценките.

Един ефект от използването на регуляризацията е, че с нарастване на λ елементите на $\hat{\theta}$ намаляват и изходът на модела се изглажда.

3.4.3.4 Регуляризация на Тихонов и метод на максималната апостериорна плътност

Описаната връзка няма отношение към числените проблеми, но показва още едно приложение на описаната техника.

В точка 2.3.4.2 се споменава за връзката на регуляризацията с метода на максималната апостериорна плътност (МАР). В МАР и метода на Бейс се отчитат, освен апостериорната информация, съдържаща се в данните, и априорните знания за параметрите. Колкото по-голяма е неопределеността в предварителната информация за даден параметър, толкова оценките са по-чувствителни към данните, а когато параметърът е известен с по-голяма точност, оценките са по-чувствителни към априорната информация. Тази особеност на МАР (и на метода на Бейс) може да се представи като регуляризация на решението. За да се отчете априорната информация за параметрите θ_a , целевата функция (3.76) се формира така:

$$\mathcal{F}(\theta) = e^T e + \lambda(\theta - \theta_a)^T(\theta - \theta_a), \quad (3.77)$$

което е еквивалентно на допълнително намаляване на дължината на вектора $\theta - \theta_a$ (т.е. на неговата 2-норма $\|\theta - \theta_a\|_2^2$). Така, вместо оценките да се „придърпват“ към нула с допълнителния член, те се прибли-

жават към априорните им стойности. Минимизацията на първото събираемо в (3.77) е свързана с отчитането на апостериорната информация, а второто събираемо отразява предварителната информация за модела. Колкото по-достоверни са елементите на θ_a , толкова параметърът λ е по-голям и второто събираемо в (3.77) доминира.

След диференциране на $\mathcal{F}(\theta)$ за градиента ѝ се получава:

$$g = -2\Phi^T y + 2\Phi^T \Phi \theta + 2\lambda \theta - 2\lambda \theta_a,$$

откъдето оптималните оценки на параметрите, с отчитане на априорната информация, са:

$$\hat{\theta} = (\Phi^T \Phi + \lambda I)^{-1} (\Phi^T y + \lambda \theta_a).$$

Ако неопределеността във всеки от елементите на θ_a е различна, целевата функция може да се зададе като:

$$\mathcal{F}(\theta) = e^T e + (\theta - \theta_a)^T \Lambda (\theta - \theta_a).$$

Матрицата Λ е диагонална с тегла по диагонала, отразяващи степента на достоверност на предварителните стойности на отделните параметри. В този случай оптималните оценки на параметрите са:

$$\hat{\theta} = (\Phi^T \Phi + \Lambda)^{-1} (\Phi^T y + \Lambda \theta_a).$$

3.4.3.5 Оценки за модел с матрица на параметрите

Аналогът на 2-нормата, когато параметрите са в матрица, е Фробениус нормата:

$$\|\Theta\|_F = \sqrt{\text{tr}(\Theta^T \Theta)}.$$

Тогава ограничаването на нормата на Θ и по този начин намаляването на ефекта от числените грешки може да се постигне, ако целевата функция е

$$\mathcal{F}(\Theta) = \text{tr}(E^T E + \lambda \Theta^T \Theta).$$

Градиентната матрица на $\mathcal{F}(\Theta)$ е

$$G = -2\Phi^T Y + 2\Phi^T \Phi \Theta + 2\lambda \Theta.$$

Приравнявайки G на нула, за регуляризираните оценки по Тихонов се получава:

$$\hat{\Theta} = (\Phi^T \Phi + \lambda I)^{-1} \Phi^T Y.$$

Разгледаните регуляризации на LS водят до числено устойчиви оценки на параметрите на модела. Друг подход за избягване на числените проблеми, представен в следващите подточки, е свързан с използване на декомпозиции на матрици.

В общия случай декомпозираната матрица се представя като произведение от матрици, всяка от които има специфична структура (диагонална, триъгълна и др.) и свойства (като ортогоналност). С използването на особеностите на тези матрици е възможно да се избегне директното обръщане на матрицата на Фишер.

3.4.4 LS с QR декомпозиция

При QR декомпозицията дадена матрица $A \in \mathcal{R}^{m \times n}$ се представя като произведение на две матрици:

$$A = QR.$$

$Q \in \mathcal{R}^{m \times m}$ е ортогонална, а $R \in \mathcal{R}^{m \times n}$ е горна триъгълна. Свойството на Q , което се използва по-долу, е, че за ортогоналните матрици е в сила $Q^{-1} = Q^T$. От друга страна, структурата на R е удобна за откриване на числени проблеми.

Стандартният алгоритъм за QR декомпозиция представлява последователност от трансформации (елиминации) на Хаусхолдер (Householder) [131]. При тях се формира матрица на Хаусхолдер, която, умножена по дадена матрица, нулира елементите ѝ в определен стълб, които са под избран елемент. Така последователно матрицата A се преобразува в триъгълна. Матрицата Q от своя страна е произведение от матриците на Хаусхолдер, в съответния ред на прилагане към A . Този алгоритъм има различни варианти, част от които зависят от свойствата на A . Но изобщо характерно за методите за декомпозиране на матрица е, че са итеративни [26]. Това води до трудности при „разпаралеляването“ на изчислителния процес, което има отношение към бързодействието на оценителя особено когато броят на факторите е голям. Например има случаи, когато се търси модел с хиляди фактори. Тогава матрицата, която се декомпозира, се състои от милиони елементи. Поради итеративната работа на алгоритмите тази дейност не може да се ускори чувствително с помощта на паралелизъм в изчисленията. От друга страна, има бързи (непаралелни) реализации на методите и рядко декомпозицията е дейност, която забавя изграждането на модела.

Едни от най-ефективните (с по-малко изчисления и едновременно с това числено устойчиви) реализации на LS са тези, използващи QR декомпозиция. По-долу са дадени два такива варианта.

3.4.4.1 QR декомпозиция на матрицата на Фишер

Има различни модификации на стандартното решение на LS. В един вариант QR се прилага директно към информационната матрица, т.е.

$$\Phi^T \Phi = QR,$$

при което декомпозираната матрица се представя като произведение от ортогоналната Q и горната триъгълна матрица R . От това, че Q е ортогонална, следва, че $Q^T Q = I$. В такъв случай системата от уравнения:

$$\Phi^T \Phi \hat{\theta} = \Phi^T y \quad (3.78)$$

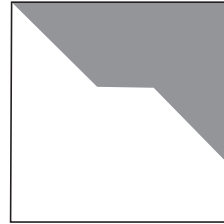
се изменя така:

$$QR\hat{\theta} = \Phi^T y \Leftrightarrow Q^T QR\hat{\theta} = Q^T \Phi^T y$$

или окончателно:

$$R\hat{\theta} = Q^T \Phi^T y.$$

Тъй като R е триъгълна матрица, то последният израз е триъгълна система, която може да се реши чрез обратна Гаусова елиминация [27], като по този начин не се изчислява инверсна матрица. Ако част от факторите са линейно зависими ($\text{rank}(\Phi^T \Phi) = r < p$), тогава структурата на R е показана на Фигура 3.11, което позволява откриването на числени проблеми и съответно тяхното избягване.



Фигура 3.11. Структура на R при съседни линейно зависими фактори (сивата област съдържа ненулевите елементи)

3.4.4.2 Оценки за модел с матрица на параметрите

Няма принципа разлика между горното изложение и случая с матрица на параметрите. Уравнение (3.78) съответства на:

$$\Phi^T \Phi \hat{\theta} = \Phi^T Y.$$

Отново, умножавайки отляво с Q , това равенство се изменя по следния начин:

$$QR\hat{\Theta} = \Phi^T Y \Leftrightarrow Q^T QR\hat{\Theta} = Q^T \Phi^T Y$$

или окончателно:

$$R\hat{\Theta} = Q^T \Phi^T Y. \quad (3.79)$$

Получената система е триъгълна и отново може да се реши чрез обратна Гаусова елиминация. Въпреки че $\hat{\Theta}$ е матрица, решаването на (3.79) не се усложнява, понеже всеки стълб на матрицата на параметрите се определя независимо от останалите.

3.4.4.3 QR декомпозиция на разширена матрица на данните

При друга реализация на LS се формира разширената матрица $[\Phi \ y]$, към която се прилага QR. Получава се:

$$[\Phi \ y] = QR = Q \begin{bmatrix} R_{11} & R_{12} \\ 0 & R_{22} \end{bmatrix}.$$

От горното равенство може да се запише, че

$$Q^T \Phi = \begin{bmatrix} R_{11} \\ 0 \end{bmatrix} \quad \text{и} \quad Q^T y = \begin{bmatrix} R_{12} \\ R_{22} \end{bmatrix}.$$

За $\mathcal{F}(\theta)$ във вид на норма е в сила:

$$\begin{aligned} \mathcal{F}(\theta) &= \|y - \Phi\theta\|_2^2 = \|Q^T y - Q^T \Phi\theta\|_2^2 = \left\| \begin{bmatrix} R_{12} - R_{11}\theta \\ R_{22} \end{bmatrix} \right\|_2^2 \\ &= \|R_{12} - R_{11}\theta\|_2^2 + \|R_{22}\|_2^2. \end{aligned}$$

В последното равенство е отчетено, че умножението на $e = y - \Phi\theta$ с ортогонална матрица не променя стойността на нормата. Причина за това е, че ефектът от умножението с Q^T е единствено завъртане на вектора в пространството, без това да е свързано с промяна на неговата дължина. Последната норма е $\|R_{22}\|_2^2 = R_{22}$, тъй като при модел с вектор на параметрите R_{22} е скалар. От равенството се вижда, че R_{22} не

зависи от θ , а отразява неопределеността в системата, дължаща се на случайни фактори. Тогава

$$\mathcal{F}(\theta) \propto \|R_{12} - R_{11}\theta\|_2^2$$

и за оптималната оценка се получава:

$$\hat{\theta} = R_{11}^{-1} R_{12}.$$

R_{11} е триъгълна и затова за намиране на $\hat{\theta}$ отново може да се използва обратна Гаусова елиминация, приложена към системата $R_{11}\hat{\theta} = R_{12}$. Както се вижда, оценките зависят от стойностите на част от матрицата R . Поради това при тази модификация не е нужно да се определя пълната матрица R , както и Q , с което се намалява броят на изчисленията. Тази числено устойчива реализация на LS е използвана в Matlab.

3.4.4.4 Оценки за модел с матрица на параметрите

При този вариант разширената матрица е $[\Phi \ Y]$, а нейната QR декомпозиция е

$$[\Phi \ y] = QR = Q \begin{bmatrix} R_{11} & R_{12} \\ 0 & R_{22} \end{bmatrix}.$$

За $\mathcal{F}(\Theta)$, записана във вид на норма, е в сила:

$$\mathcal{F}(\Theta) = \|Y - \Phi\Theta\|_F^2 = \|R_{12} - R_{11}\Theta\|_F^2 + \|R_{22}\|_F^2.$$

В случая R_{22} е $\ell \times \ell$ мерна матрица, която отново не зависи от модела, а от случайните фактори в поведението на системата. Тогава

$$\mathcal{F}(\Theta) \propto \|R_{12} - R_{11}\Theta\|_2^2$$

и за оптималната оценка се получава:

$$\hat{\Theta} = R_{11}^{-1} R_{12}.$$

3.4.5 LS с SVD декомпозиция

Декомпозицията по сингулярни стойности (SVD – Singular Value Decomposition) на матрицата $A \in \mathcal{R}^{m \times n}$ е

$$A = USV^T,$$

където $U \in \mathcal{R}^{m \times m}$ и $V \in \mathcal{R}^{n \times n}$ са ортогонални, а $S \in \mathcal{R}^{m \times n}$ е диагонална, с елементи по диагонала, равни на сингулярните числа на A , т.е.

$\text{diag } S = [s_1 \ s_2 \ \dots \ s_h]$, за $h = \min(m, n)$, при това – подредени в намаляваща последователност. По-долу се използва икономичната SVD, при която, ако $m > n$, а такъв е разглежданият случай, матриците са с размерност $U \in \mathcal{R}^{m \times n}$ и $V \in \mathcal{R}^{n \times n}$ и $S \in \mathcal{R}^{n \times n}$.

Матрицата на Фишер, която се обръща в LS, е реална и симетрична. За нейната диагонализация може да се използва методът на Якоби (Jacobi), при който се изпълнява последователност от т.нар. равнинни ротации [131]. Всяка от тези ортогонални трансформации нулира недиагонален елемент, като следващата ротация води до ненулеви стойности на нулираните до момента елементи. Въпреки това при многократно прилагане на ротациите недиагоналните елементи стават все по-малки. Както се вижда, този метод също е итеративен, а броят на итерациите много зависи от матрицата, която се декомпозира. Ефективни алгоритми за SVD се получават, ако матрицата (реална и симетрична) първо се приведе в тридиагонална форма (с ненулеви елементи по главния диагонал, над и под него), а след това се редуцира до диагонална. За първата част може да се използват споменатите в предишната точка трансформации на Хаусхолдер или ротациите на Гивънс (Givens). В основата на вторите е методът на Якоби, но тъй като при Гивънс началната матрица е в тридиагонална форма, при подходяща последователност от ротации, след всяка нова ротация, нулираните вече елементи не се променят. Въпреки това редукцията на Хаусхолдер е предпочитана поради по-малкия брой итерации (за матрицата на Фишер $F \in \mathcal{R}^{p \times p}$, броят им е $p - 1$).

Споменатите ефективни алгоритми отново зависят от итеративни действия, а това ограничава възможностите за паралелно изпълнение.

3.4.5.1 SVD декомпозиция на матрицата на данните

Когато част от факторите са линейно зависими, съществуват безброй модели, минимизиращи целевата функция. Тогава, за да се избере подходящ модел, се налагат допълнителни изисквания като (3.75). Тази оценка, която минимизира $\mathcal{F}(\theta)$ и същевременно е с минимална 2-норма, може да се определи с използване на SVD декомпозицията, приложена към матрицата на данните [61, 85, 146]. Тъй като в задачата за оценяване $(N - n)\ell \gg p$, то може да се формира икономичната SVD:

$$\Phi = USV^T. \quad (3.80)$$

Матриците $U \in \mathcal{R}^{(N-n)\ell \times p}$ и $V \in \mathcal{R}^{p \times p}$ са ортогонални, а $S \in \mathcal{R}^{p \times p}$ е квадратна и диагонална, с неотрицателни елементи, равни на сингулярните числа на Φ , подредени така, че $s_1 \geq s_2 \geq \dots \geq s_p$. Когато

$s_1 \gg s_p$, числото на обусловеност е $\text{cond } \Phi = \text{cond } S = \frac{s_1}{s_p} \gg 1$, а това е индикация за възникване на числени проблеми. При пълна линейна зависимост между стълбове на Φ , част от сингулярните числа са 0, т.е. $s_p = 0$ и $\text{cond } S = \infty$).

Нека s_l е долна граница на сингулярните числа на Φ , под която се приема, че стойностите им са незначими. По-долу се показва, че именно близките или равни на нула сингулярни числа водят до числени проблеми [52, 85, 86]. Нека с помощта на s_l се определи $S_1 \in \mathcal{R}^{r \times r}$ – диагонална подматрица на S , която съдържа доминиращите, а също и подматрицата S_2 от незначимите сингулярни числа. Тогава $\text{cond } S_1 = \frac{s_1}{s_r}$ е ограничено с помощта на праговата стойност s_l (понеже $s_r \geq s_l$), което осигурява числено устойчиви оценки. На практика S_1 отразява наличието на полезната (търсена) съставка в данните, а S_2 зависи от неопределеността. Колкото по-високо е нивото на смущенията в данните, толкова диагоналните елементи на S_2 са по-големи. С използване на това разделяне на S (3.80) може да се запише като

$$\Phi = [U_1 \ U_2] \begin{bmatrix} S_1 & 0 \\ 0 & S_2 \end{bmatrix} \begin{bmatrix} V_1^T \\ V_2^T \end{bmatrix}.$$

За да се реши задачата на LS по числено устойчив начин, се въвежда диагоналната матрица $\tilde{S}$. Нейните диагонални елементи се задават по следния начин:

$$\tilde{s}_i \begin{cases} s_i, & \text{за } s_i \geq s_l, \\ 0, & \text{за } s_i < s_l, \end{cases}$$

за $i = \overline{1, p}$. В такъв случай целевата функция, записана като 2-норма на остатъка, е

$$\mathcal{F}(\theta) = \|y - \Phi\theta\|_2^2 = \|y - USV^T\theta\|_2^2.$$

Понеже елементите на S_2 са незначими, то

$$\begin{aligned} \mathcal{F}(\theta) &\approx \|y - U\tilde{S}V^T\theta\|_2^2 \\ &= \left\| y - [U_1 \ U_2] \begin{bmatrix} S_1 & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} V_1^T \\ V_2^T \end{bmatrix} \theta \right\|_2^2 \\ &= \|y - U_1 S_1 \theta_{v1} - 0\theta_{v2}\|_2^2. \end{aligned} \tag{3.81}$$

В горното преобразуване е положено:

$$\begin{bmatrix} \theta_{v1} \\ \theta_{v2} \end{bmatrix} = \begin{bmatrix} V_1^T \\ V_2^T \end{bmatrix} \theta. \quad (3.82)$$

Приемането, че $S_2 = 0$, е равносилно на допускането, че $p - r$ стълба на матрицата Φ са линейно зависими, а останалите r стълба предоставят значимо (според прага s_l) количество информация за описанието на изходите. Това означава, че в S_1 няма сингулярни стойности, които са близки до 0.

Минималната стойност на $\mathcal{F}(\theta)$ се достига, когато $\tilde{\theta}_{v1} = S_1^{-1} U_1^T y$, а векторът $\tilde{\theta}_{v2}$ е произволен (тъй като не влияе на нормата (виж (3.81))). До този извод се стига и в точка 3.4.2. Така за оценките на параметрите на множеството Ω от безброй много модели, минимизиращи $\mathcal{F}(\theta)$, се получава:

$$\tilde{\theta} = \begin{bmatrix} V_1 & V_2 \end{bmatrix} \begin{bmatrix} \tilde{\theta}_{v1} \\ \tilde{\theta}_{v2} \end{bmatrix} = V_1 S_1^{-1} U_1^T y + V_2 \tilde{\theta}_{v2}. \quad (3.83)$$

Тъй като $V_2^T V_1 = 0$, нормата на $\tilde{\theta}$ е

$$\|\tilde{\theta}\|_2^2 = \|V_1 S_1^{-1} U_1^T y\|_2^2 + \|V_2 \tilde{\theta}_{v2}\|_2^2.$$

Очевидно нормата $\|\tilde{\theta}\|_2^2$ е минимална, когато векторът на свободните параметри е $\tilde{\theta}_{v2} = 0$ и така за оптималното решение на (3.75) се получава:

$$\hat{\theta} = V_1 S_1^{-1} U_1^T y.$$

Последният израз отговаря на числено устойчива реализация на LS с използване на икономичната SVD.

В действителност почти винаги $S_2 \neq 0$. Тъй като $\text{cond } S_1 = \frac{s_1}{s_r}$ се контролира с прага s_l , той се избира така, че обръщането на S_1 (практически изчисляването на s_i^{-1}) да не води до числени проблеми при формирането на $\hat{\theta}$.

Предимство на тази реализация е, че изолирането на подпространството, което съдържа значимата информация за обекта, се постига с просто сравнение на сингулярните числа с въведения праг s_l . Така търсенето на оптималните оценки се извършва в подпространство на това

на параметрите, в което няма опасност от възникване на числени проблеми [36, 97, 130].

3.4.5.2 Оценки за модел с матрица на параметрите

При това представяне Φ има различна структура, но това не променя идеята на реализацията. Нека отново в икономичната SVD на матрицата на данните се извърши разделянето:

$$\Phi = [U_1 \ U_2] \begin{bmatrix} S_1 & 0 \\ 0 & S_2 \end{bmatrix} \begin{bmatrix} V_1^T \\ V_2^T \end{bmatrix},$$

като е използван предварително зададеният праг на значимост s_l . Целевата функция, записана като Фробениус норма, е

$$\mathcal{F}(\Theta) = \|Y - \Phi\Theta\|_F^2.$$

След аналогични разсъждения като при модела с вектор на параметрите, се достига до

$$\mathcal{F}(\Theta) \approx \|Y - U_1 S_1 \Theta_{v1} - 0 \Theta_{v2}\|_2^2$$

като е положено $\Theta_{v1} = V_1^T \Theta$ и $\Theta_{v2} = V_2^T \Theta$. Минимумът на $\mathcal{F}(\Theta)$ се получава за $\tilde{\Theta}_{v1} = S_1^{-1} U_1^T Y$ и произволна матрица $\tilde{\Theta}_{v2}$, т.е. съществуват безброй много модели, минимизиращи $\mathcal{F}(\Theta)$. От тях за матрицата $\hat{\Theta}$ с минимална Фробениус норма се получава:

$$\hat{\Theta} = V_1 S_1^{-1} U_1^T Y.$$

И други декомпозиции, освен разгледаните QR и SVD, се използват при изграждането на числено устойчиви реализации на методите за оценяване на параметри. Например декомпозицията по собствени стойности (EVD – Eigenvalue Decomposition) е приложима към някои видове квадратни матрици, към които спада и матрицата на Фишер. Въпреки че SVD може да се прилага към произволни $m \times n$ мерни матрици, тя е много сходна с EVD, когато декомпозираната матрица е симетрична и с реални елементи. Затова EVD, за която има бързи алгоритми, често се използва за числена реализация на LS. Декомпозицията на Холески (Cholesky) [98, 118, 45], UD декомпозицията (известна още като декомпозиция на Бирман (Bierman)) [51] и др. също намират приложение при реализацията на оценителите.

3.5 Избор на структурата на модела

Изходите на системата отразяват нейната реакция. Те се определят още в началото на идентификацията, когато се конкретизира системата, и зависят от целта, за която ще се използва моделът. Но тъй като източникът на информация в идентификацията основно са данните, в случай че има изходи, които не са възбудени (и не може да се проведат допълнителни експерименти), те се премахват в третия етап (събирането и подготовката на данните за същинското моделиране). Така в набора, подготвен за изграждането на модела, се съдържат само изходи, чиито стойности се изменят.

Характерно за техническата област е наличието на априорна информация за входните величини. В тези случаи и въздействията може да се конкретизират преди същинското моделиране. От друга страна, в пазарния сектор, социологията, медицината и други области често входовете не са предварително известни. В този случай по време на първия етап от идентификацията се избира множество от потенциални въздействия, от което впоследствие се конкретизира подмножеството от входове на модела, чието комбинирано влияние върху изходите е доминиращо. А пренебрегнатите величини, които не участват в крайния модел, се разглеждат като смущения от околната среда. Уточняването на подмножеството от значими въздействия е част от дейностите по избора на структурата на модела, описани в тази точка.

Броят на факторите в една динамична система е по-голям от броя на входовете. Причината е, че в описанието на динамиката участват предишни стойности на сигналите. Ако например дадена система е с 5 входа и 5 изхода и се описва с ARX модел с матрица на параметрите, като всички полиноми в $A(q^{-1})$ и $B(q^{-1})$ са от първи ред, то броят на регресорите е $z = 10$, а параметрите са $p = 50$, но ако полиномите са от втори ред, то факторите стават 20, а параметрите 100.

Целта на разгледаната по-долу част от изграждането на модела е да се определи множество от подходящи въздействия, а ако системата е динамична, да се изберат също степени на полиномите в полиномните матрици, както и да се определят евентуалните закъснения в системата. Тъй като изходите се уточняват в предишни етапи от идентификацията, определянето им не се разглежда по-долу.

3.5.1 Оптимизиране на структурата

Обикновено изборът на структурата се реализира като оптимизационна задача. Освен това, когато е голям броят на потенциалните фактори, между които може да има неинформативни или мултиколинеарни (което е характерно за много ММО системи), е удачно определянето на структурата да се автоматизира, като се търси оптимум на определен показател. В някои случаи уточняването на структурните параметри е продиктувано и от бизнес нуждите, физическата реализация на системите и др. Тези изисквания играят роля на ограничения или на допълнителни критерии в задачата за избор на структура.

Често причината за големия брой потенциални фактори е липсата на априорна информация за това, кои величини са значими. Факторите допълнително нарастват след прилагането на някои техники по време на предварителната обработка на данните. Примери за това са категоризацията, където факторът се замества с множество фиктивни променливи, играещи ролята на нови фактори; трансформация, като се запазват отделните нелинейни трансформации на факторите, с цел впоследствие автоматично да се избере най-подходящата, и др. В резултат от тези дейности броят на факторите може да нарасне десетки пъти. Такъв е случаят при изграждането на модел за оценка на кредитния риск, където факторите след споменатата обработка може да станат хиляди.

Намаляването на броя фактори още по време на предварителната обработка може значително да повиши ефективността на цялостния процес на идентификация. Това е една от причините за анализа на данните, който се извършва преди същинското моделиране.

Търсенето на оптимална, в определен смисъл, структура е свързано както с максимизиране на достоверността на модела, така и с опростяване на неговата структура. Използването на модел с полиноми от по-нисък ред и/или с по-малко входове опростява структурата на други елементи от по-машабни системи. Например елементи на системите за управление, които зависят от модела, са наблюдатели на състоянието, компенсатори, регулатори и др.

Отчитането на точно определени въздействия в някои случаи се налага по икономически причини, като целта е намаляване на вложените средства. Това може да доведе до усложняване на модела, което показва, че изборът на структурата е многокритериална задача. Например, възможно е някои технически средства за автоматизация (изпълнителни механизми, регулиращи органи, датчици и др.) да са предпочита-

ни пред други. Друг пример е ограничаването на характеристиките на клиентите, които банката купува от кредитни бюра. Тук има стремеж в модела да участват възможно по-малко фактори, които са свързани с отделянето на значителни средства за сметка на използването на други (често повече) фактори, чието ползване не изисква много средства.

Не на последно място, с пренебрегването на величини, които не допринасят за точността на модела, може да се избегне преоформяването на описанието. Допълнителен ефект е намаляването на вероятността от числени проблеми.

3.5.2 Структурни параметри на многомерни модели

Структурата на модела се дефинира с помощта на набор от т.нар. структурни параметри, чиито оптимални (в определен смисъл) стойности се търсят. В долните разглеждания се приема, че по време на предварителната обработка и анализ на данните всички изходи, които не са подходящи за моделиране, са премахнати.

3.5.2.1 Входни величини

Една група структурни параметри са индексите на въздействията, които участват в модела. По отношение на входовете целта е да се подберат само тези, комбинацията от които описва максимално точно изходите. Нека преди конкретизирането на входните сигнали първоначално да са налични данни за $\tilde{m}$ потенциални входове (по-долу те се означават с вектора $\tilde{u}_k \in \mathcal{R}^{\tilde{m}}$). Също така нека след избора на входните величини, индексите на участващите в описанието да са:

$$h_i \text{ за } i = \overline{1, m} \text{ и } m \leq \tilde{m}.$$

Така, ако измежду всички въздействия (компоненти на $\tilde{u}_k$), за входове на модела са избрани 2-и, 5-и и 7-и сигнал, то $m = 3$, а $h_1 = 2$, $h_2 = 5$ и $h_3 = 7$ и съответно $u_k = [\tilde{u}_{2,k} \ \tilde{u}_{5,k} \ \tilde{u}_{7,k}]^T$. В общия случай

$$\begin{aligned} u_k &= [u_{1,k} \ \dots \ u_{m,k}]^T \\ &= [\tilde{u}_{h_1,k} \ \dots \ \tilde{u}_{h_m,k}]^T. \end{aligned}$$

По този начин търсените структурни параметри, с които се описват входовете на модела, са индексите h_i . Понякога като параметър, който характеризира структурата на ММО модела, се разглежда и броят на входовете m .

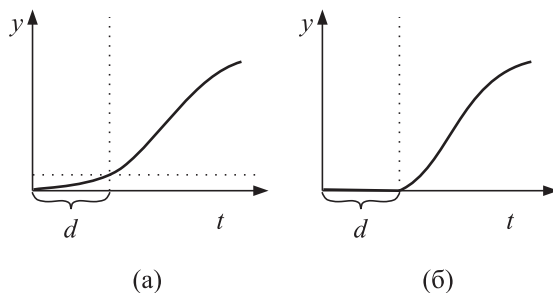
Едно приложение на представянето на модела с матрица на параметрите е, когато броят на входовете трябва да се минимизира (например по икономически причини). В този случай може първо да се потърси модел, описан с матрица на параметрите. Идеята е, че ако за описанието на един изход е добавено определено въздействие в модела, то е удачно то да се използва, доколкото е възможно, за описанието и на другите изходи.

3.5.2.2 Степени на полиномите

Когато броят на входовете и изходите е голям и няма първоначална информация за максималните степени на полиномите, може първо да се потърси модел с матрица на параметрите, при което се уточняват двата параметъра na и nb (или na_i и nb_j , за $i = \overline{1, \ell}$ и $j = \overline{1, m}$). След тази, първоначално груба, ориентация по отношение на структурата трябва да се пристъпи към точното определяне на степените na_{ij} и nb_{ij} на отделните полиноми, като се използва описание на модела с вектор на параметрите.

3.5.2.3 Закъснение

Интервалът от време от изменението на даден фактор до наблюдаването на реакцията от това изменение е закъснение в системата. Тази особеност на динамичните системи може да се отчете, като в модела се въведат структурни параметри, отговарящи на закъсненията. Поради инерционността на някои системи реакцията от изменението на даден фактор може да е плавна, както е показано на Фигура 3.12 (а).



Фигура 3.12. Закъснение в система, дължащо се на инерционност (а) и чисто закъснение (б)

С цел опростяване на структурата на модела първоначалното бавно изменение на изхода, докато е незначително, може да се разглежда като чисто закъснение. Друг тип закъснение, дадено на Фигура 3.12 (б), е транспортното (например дължащо се на поточна линия за пренос на материал); закъснение, дължащо се на забавяне в комуникации и т.н. Икономически обекти със закъснение са примерно свързани с възвръщаемостта на инвестиции, в медицината – описващи ефекта от лекарство и т.н.

Тъй като закъснението е свързано с параметри в модела, чиито стойности са нула (известно време не се наблюдава реакция) или са пренебрежими, определянето му е важно за получаването на структура, в която не се съдържат излишни параметри. Например, ако реакцията на даден изход при изменението на вход се наблюдава след $d + 1$ такта, то освен свободния параметър $b_{ij,0}$ следващите d параметъра на съответния полином на $V(q^{-1})$ са нули и не трябва да се оценяват. Така в модела не се добавят изкуствено степени на свобода по отношение на параметри, които трябва да са нула, и по този начин се намалява вероятността от преоразмеряване.

3.5.3 Стъпкови методи

Ако са известни граничните стойности на структурните параметри на ММО модела, теоретично може да се формират всички възможни модели и от тях да се избере този с най-добра, в определен смисъл, структура. Но на практика, когато системата е с много структурни параметри поради огромния брой комбинации между стойностите на тези параметри, подходът с формирането на всички възможни модели е неприложим.

Съществуват различни подходи за избор на структура [11, 147, 126, 62], при които не се оценяват всички потенциални модели. Обикновено методите водят до различен резултат. Затова често се говори не за оптимална, а за подходяща структура на модела, като в практиката се избира такъв подход, който определя задоволително структурата на модела с възможно по-малко оценени модели.

3.5.3.1 Идея на методите

В повечето случаи методите за избор на структура са стъпкови, като на всяка стъпка се сравнява подобрението/влошаването на достоверността на модела при промяна на някой от структурните параметри.

На практика промяната на тези параметри (независимо дали са закъснения, степени на полиноми или брой входни въздействия) е равностойна на добавяне или премахване на фактори в модела.

Ако към факторите, описващи даден изход, се добави/премахне фактор, в зависимост от промяната в достоверността на модела може да се определи доколко факторът е значим за описанието на даден изход. Този стъпков процес продължава, докато не се изпълнят определени условия, свързани със значимостта на факторите или съответно със степента на изменение на достоверността на модела.

При ММО системите, когато моделът е представен с матрица на параметрите, всяка промяна във вектора на регресорите води до изменение на модела по всички изходи. Но когато представянето е с вектор на параметрите, добавянето/премахването на фактор по отношение на един изход не влияе на качеството на модела по отношение на останалите изходи. Затова в този случай на всяка стъпка е удачно да се извършва едновременно по една промяна във всеки MISO подмодел (т.е. общо ℓ изменения в определени части на ММО модела). Това може значително да намали времето за уточняване на структурата особено ако наборът от входно-изходни данни е голям и на всяка стъпка е необходимо той да се прочита (четенето на данни е значително по-бавна операция в сравнение с изчислителните операции). Както ще бъде показано, при линейно параметризираните модели е възможно стъпковият метод да се изпълни с едно прочитане на данните. Но ако моделът е нелинеен по параметри, обикновено на всяка итерация на метода се налага данните многократно да се прочитат.

3.5.3.2 Видове стъпкови методи

Тези методи може да се класифицират според начина, по който се изменят търсените структурни параметри (дали нарастват, намаляват или промяната е в двете посоки). Също методите се отличават и според класа на модела, както и по критерия за избор на най-подходяща промяна на структурата. Тези класификации са разгледани по-долу.

Изменение на структурните параметри

Според тази класификация методите се делят на:

- права стъпкова регресия;
- обратна стъпкова регресия;
- комбинирана стъпкова регресия;
- регресия с пълно множество от модели.

Последният вариант не е стъпков в истинския смисъл на думата, тъй като няма значение каква е последователността на изменение на структурата на модела. Избраната структура винаги е една и съща – оптимална, според зададения критерий и фиксираните граници на изменение на параметрите.

За разлика от формирането на пълно множество от модели първите три метода са итеративни, в смисъл че текущият резултат от дадена итерация оказва влияние на следващите итерации и съответно на крайната структура на модела. Обикновено тези методи водят до формирането на различни модели. Само при регресията с пълно множество от модели се гарантира, че измежду възможните структури е избрана най-добрата. Затова, когато става дума за стъпков метод (където не се обхождат всички възможни структури), под оптимална се има предвид структурата, която е избрана от конкретния метод.

Ако структурата се изменя така, че факторите се добавят и/или премахват на групи, то регресията е групова (съответно права, обратна или комбинирана).

Клас на модела

В зависимост от това, дали моделът е линеен по отношение на параметрите, методите най-общо биват:

- стъпкова линейна регресия;
- стъпкова нелинейна регресия.

За да се определи най-подходящата промяна в текущата структура, при стъпковата линейна регресия се формират множество от съревноваващи се модели, без да е необходимо наборът от данни да се прочита за всеки модел, както е описано в точка 3.5.5.3. От получените модели се избира текущият най-подходящ и процедурата продължава със следващата стъпка на изменение на структурата. Но когато описанието е нелинейно по параметри, в общия случай за формирането на модел, на дадена стъпка се изпълнява (итеративен) метод за оптимизация. Това означава, че данните се прочитат на всяка итерация за формиране на модел, като тази дейност продължава, докато оценките не сходят. Ако на всяка стъпка се изграждат съревноваващи се модели, т.е. оптимизацията се стартира многократно, времето за избор на структура би нараснало драстично. По тази причина принципът на работа на методите за нелинейни модели е различен от този за линейни (по параметри). Тук не се формират съревноваващи се модели, а на базата на текущия модел се извършва бърз анализ на повърхнината на целевата функция, в смисъл – как тя се изменя, ако структурата се промени. След това се

избира тази нова структура, за която се очаква, че подобрението ще е най-голямо (ако се добавя нов фактор) или влошаването ще е най-малко (при премахване на фактор). Накрая еднократно се прилага метод за оптимизация, с който се намира модел с подобрената структура.

Критерии за оценка на текущата структура

За да може стъпковите алгоритми да функционират коректно, е необходимо предварително да се уточнят границите на изменение на структурните параметри. Така предварително ще е известно множеството от потенциални фактори, които могат да участват в модела и измежду които ще се избира подходящо подмножество. В стъпковите методи се използват критерии за избор на текущ най-значим фактор за добавяне в модела или най-незначим фактор за премахване от модела (т.е. за промяна на структурата). За тази цел може да се използват тестове за статистическа проверка на хипотези. За линейни модели и при нормално разпределение на изхода (съответно и на остатъка) намират приложение F и t -тестът. Може би най-разпространеният тест в практиката, използван за оценка на значимостта на фактор, е F -тестът [80, 11, 5] и затова по-долу са изложени стъпкови методи, използващи този тест.

Когато описанието е нелинейно, в зависимост от функцията на правдоподобие се използват различни тестове. Например при биномно разпределение (изходът приема две възможни стойности) тестовата статистика е с χ^2 разпределение.

Решение за промяна в структурата може да се вземе с помощта на коефициента на корелация между остатъка на текущия модел и всички неучастващи в модела фактори. Според този вариант (удачен за линейни по параметри модели) най-подходящ за добавяне е този фактор, който е най-силно корелиран с остатъка, т.е. описва в най-голяма степен поведението на системата, което все още не е отчетено от модела.

Освен определянето на най-подходящата промяна в структурата е необходимо и да се следи доколко тази промяна подобрява модела, и съответно да се вземе решение дали да се продължи с изменението на структурата, или стъпковата процедура да се прекрати. От една страна, за това се използват доверителните вероятности от споменатите тестове за значимост (използването им е показано по-долу). Също така е възможно да се използват критерии, с които се балансира между достоверност и сложност на модела, т.е. най-добрата структура според даден критерий може да не отговаря на най-достоверния модел, а на такъв,

който е едновременно удобен (с проста структура) и достатъчно точен за приложението, за което се изгражда. Подходящи критерии са AIC, BIC, а също често използвана е и Малоу статистиката (Cp) (описани в точка 2.4.6).

В допълнение обикновено се следят различни спомагателни статистики. Част от тях характеризират достоверността на текущия модел като цяло, например коефициентът на определеност, средноквадратичната грешка и др., както и статистики, свързани с оценките на параметрите поотделно като стандартна грешка, частичната F-статистика и нейната доверителна вероятност и др.

Добре е също така да се следи доколко всеки от неучастващите в модела фактори се описва с избраните до момента. За целта се използва статистиката VIF (Variance Inflation Factor), описана в точка 2.4.4.2. Колкото по-малка е стойността на VIF, толкова повече уникална информация се съдържа в изследвания фактор.

Не на последно място, при избора на структурата е важно да се отчете и предварителното знание за системата (ако такова съществува), свързано с нестатистически изисквания към структурата на модела. Понякога (статистически) значим според избрания критерий фактор може да се премахне от модела, ако участието му не е логично от гледна точка на априорното знание.

3.5.4 Регресия с пълно множество от модели

Този метод е приложим при разумен брой на възможните структури на модела. Следният пример е даден, за да се добие представа за броя на моделите, които трябва да се формират според този метод.

Пример. *Формиране на пълно множество от модели на многомерна система*

Нека е дадена статична MISO система с $\tilde{m} = 3$ потенциални входа, измежду които трябва да се определи подмножество, което е подходящо за описание на изхода. Тогава възможните модели имат следните входове: $\{u_{1,k}\}$, $\{u_{2,k}\}$, $\{u_{3,k}\}$, $\{u_{1,k}, u_{2,k}\}$, $\{u_{1,k}, u_{3,k}\}$, $\{u_{2,k}, u_{3,k}\}$ и $\{u_{1,k}, u_{2,k}, u_{3,k}\}$. С други думи при 3 потенциални входа моделите са $n_3 = 2^3 - 1 = 7$ на брой.

Нека системата е с $\tilde{m} = 20$ потенциални входа. Тогава броят на възможните модели е

$$n_{20} = 2^{20} - 1 \approx 1.0485 \times 10^6.$$

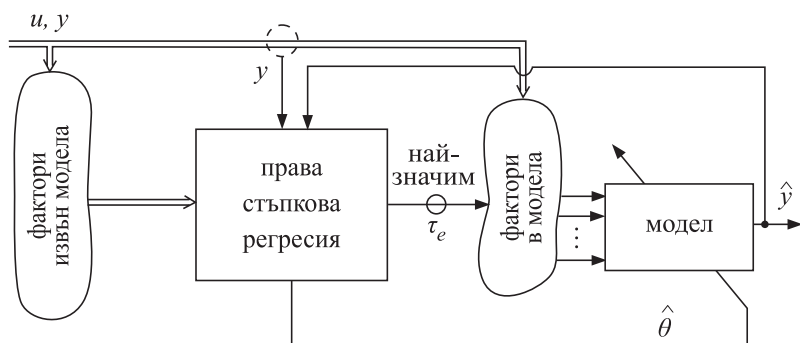
Ако броят на потенциалните фактори е $\tilde{m} = 50$, тогава $n_{50} \approx 1.1259 \times 10^{15}$. Както неведнъж беше споменато, има области, където потенциалните входове са стотици или хиляди. Също така, ако се търси динамичен модел, броят на потенциалните структури нараства значително. ◇

Има бързи реализации на този подход, например като се формират последователности от модели, които се различават с по един фактор. Повече информация за такава реализация е дадена в [124]. В точка 3.5.5.3 е представен начин за имплементация на алгоритъма на правата стъпкова регресия. Описаната идея може да се използва за ефективна реализация и на метода с пълно множество от модели.

3.5.5 Права стъпкова регресия

При голям брой фактори за предпочитане е да се използва правата пред обратната (точка 3.5.6) стъпкова регресия. При правата структурата постепенно се усложнява, докато подобрението е значимо, и така не се налага да се построяват модели с ненужно голяма размерност.

3.5.5.1 Идея на метода



Фигура 3.13. Права стъпкова регресия

Алгоритъмът на правата стъпкова регресия стартира с празен модел (т.е. модел само със свободни членове), като на всяка стъпка се добавя нов фактор или фактори в зависимост от представянето на модела (Фигура 3.13). По този начин, ако не съществува мултиколинеарност в данните, се добавят фактори, които описват изхода във все по-малка степен, и така подобрението в модела намалява. Изпълнението на

стъпковия алгоритъм се прекратява при достигане до определен праг на значимост или до изчерпване на факторите. Прагът на значимост за добавяне на фактори в модела τ_e в долния алгоритъм има смисъл на доверителната вероятност, над която подобрението в описанието на y_k при добавяне на нов фактор се приема, че се дължи на случайността (т.е. е незначително). Смисълът на τ_e се изяснява в представения по-долу алгоритъм.

При динамични системи даден фактор, който постъпва в модела, може да е нов вход за текущото описание, но може да бъде и допълнително изместен във времето сигнал, който вече участва в модела. Така с помощта на стъпковия алгоритъм се определят значимите входове на модела, а също и степените на полиномите и чистите закъснения.

3.5.5.2 Алгоритъм на метода

За добиване на по-ясна представа как се изпълнява правата стъпкова регресия, е описан алгоритъм за представяне с вектор на параметрите. Освен това моделите са MISO, защото при повече изходи всяка стъпка, без инициализиращата, се изпълнява ℓ пъти. Алгоритъмът на правата стъпкова регресия, използващ F-тест за определяне на значимостта на факторите, е следният.

1. Инициализация

Задава се прагът на значимост за добавяне на фактори τ_e (често се означава като SLE – Significance Level to Enter [62]). Ако няма специфични изисквания, подходяща стойност на този параметър е $\tau_e = 0.05$. Ако има стремеж моделът да е с по-проста структура (по-малък брой параметри), се избира по-малка стойност, но това води до по-неточно описание на изхода. Ако основната цел е изходът да се опише по-точно или ако изследователят предпочита да избере структурата на модела, се задава по-голям праг, например $\tau_e = 0.15$. Така след прекратяването на итеративната процедура изследователят може да избере някой от междинните най-добри модели, които са с по-проста структура, като при избора отчита статистиките, свързани с модела и известната му априорна информация.

Също по време на инициализацията се задават максималните степени на полиномите $na_{max} = \max(na_j)$ и $nb_{max} = \max(nb_j)$, ако е избран ARX модел, а също и чистите закъснения $d_{y,j}$ и $d_{u,j}$ (ако са известни). Тези величини са всъщност границите на изменение на отместването на сигналите във времето и от тях зависи

броят на потенциалните фактори. В допълнение може да се инициализират и други параметри, имащи отношение към числената реализация на оценителя (например минималната значима сингулярна/собствена стойност, ако се използва SVD или EVD декомпозиция и т.н.).

2. Множество от съревноваващи се модели

Нека в i -тата итерация има n_{out} фактора извън модела. Тогава се формират n_{out} съревноваващи се модела. Те съдържат включените до момента фактори и по един фактор извън модела. За i -тата итерация те са

$$y_{+j}^{(i)} = \Phi_{+j}^{(i)} \theta_{+j}^{(i)}, \text{ за } j = \overline{1, n_{out}},$$

където матрицата на факторите за j -тия модел е $\Phi_{+j}^{(i)} = [\Phi^{(i-1)} \varphi_j]$. Стълбовете на $\Phi^{(i-1)}$ се състоят от стойностите на участващите фактори в модела от предишната стъпка, а векторът φ_j съдържа стойностите на j -тия невключен в модела фактор (кандидат за включване), т.е.

$$\varphi_j = [\varphi_{j,n+1} \ \varphi_{j,n+2} \ \dots \ \varphi_{j,N}]^T. \quad (3.84)$$

Оценките на параметрите на j -тия съревноваващ се модел са

$$\theta_{+j}^{(i)} = (\Phi_{+j}^{(i)T} \Phi_{+j}^{(i)})^{-1} \Phi_{+j}^{(i)T} y.$$

3. Значимост на потенциалните фактори

Критерият за значимост е

$$F_{+j}^{(i)} = \frac{v_2}{v_1 - v_2} \frac{\text{SSR}_{+j}^{(i)} - \text{SSR}_{+j}^{(i-1)}}{\text{SSE}_{+j}^{(i)}}, \quad (3.85)$$

където

$$\text{SSR}_{+j}^{(i)} = y_{+j}^{(i)T} y_{+j}^{(i)}, \quad \text{SSE}_{+j}^{(i)} = e_{+j}^{(i)T} e_{+j}^{(i)}.$$

Величината $F_{+j}^{(i)}$ е частичната F статистика за модела и отразява подобрението му при добавянето на фактор. Степените на свобода на текущия най-добър модел (с p параметъра), избран в $i - 1$ -та итерация, са $v_1 = N - p$, а на съревноваващите се модели (с $p + 1$ параметри), определени в i -та итерация, са $v_2 = N - (p + 1)$. Знаменателят $v_1 - v_2$ е 1, понеже на всяка стъпка се добавя по един нов фактор. Ако факторите се добавят на групи (виж точка 3.5.8), $v_1 - v_2$ е броят на факторите в добавената група. От (3.85) се вижда, че с нарастване (намаляване) на значимостта на фактора φ_j

разликата между $SSR^{(i-1)}$ и $SSR_{+j}^{(i)}$ нараства (намалява), откъдето и $F^{(i,j)}$ също нараства (намалява).

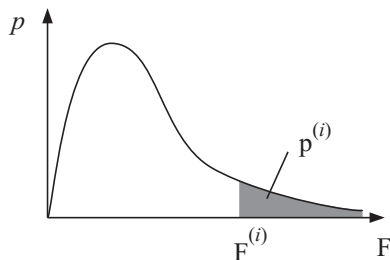
4. Избор на най-значимия фактор

Тук се определя новият най-добър (според F критерия) модел от всички потенциални модели, или с други думи на тази стъпка от неучастващите в модела фактори се избира този, който максимизира стойността на частичната F статистика. Нека търсеният фактор е с индекс

$$j^* = \arg \max_j F_{+j}^{(i)}.$$

5. Проверка за край

Статистиката $F^{(i)} = F_{+j^*}^{(i)}$ има F -разпределение и се използва за вземане на решение дали подобрението на модела с добавяне на новия j -ти фактор е значимо, или се дължи на случайността. За тази цел се прилага F -тестът, при който се задава нулевата хипотеза, която гласи, че при добавяне на новия j^* -ти фактор, подобрението в модела се дължи на случайността. Съответната алтернативна хипотеза е, че подобрението не се дължи на случайността (факторът е значим).



Фигура 3.14. Вероятност подобрението в модела да се дължи на случайността

Изчислява се доверителната вероятност $p^{(i)}$, показана на Фигура 3.14. Тя отразява риска да се допусне грешка, като се приеме, че алтернативната хипотеза е вярна, и се нарича още равнище на значимост. В [131] е описан алгоритъм за изчисляване на $p^{(i)}$, както и свързаната с него теория.

На фигурата се вижда, че колкото стойността на $F^{(i)}$ е по-голяма (факторът е по-значим), толкова площта $p^{(i)}$ на опашката на разпределението, определена от $F^{(i)}$, е по-малка, т.е. по-трудно може

да се приеме нулевата хипотеза. На практика в началото на процедурата влизат най-значимите фактори, а с нарастване на стъпките (съответно и на броя на факторите в модела) тенденцията е извън модела да останат все по-незначими фактори. При това положение с нарастване на i $F^{(i)}$ намалява, а $p^{(i)}$ расте. В граничния случай, ако новият фактор е напълно незначим, $F^{(i)} = 0$, а $p^{(i)} = 1$ (вероятността подобрението да се дължи на случайността е 1, т.е. факторът е напълно незначим). При избора на структура до такъв случай не се достига, тъй като условието за край на стъпковия алгоритъм е

$$p^{(i)} > \tau_e.$$

Ако това условие е изпълнено, то най-значимото подобрение в модела не е достатъчно значимо и процедурата се прекратява. Приема се, че

$$\hat{\theta} \leftarrow \theta^{(i-1)}, \Phi \leftarrow \Phi^{(i-1)} \text{ и } p \leftarrow i - 1.$$

В противен случай $i \leftarrow i + 1$ и $n_{out} \leftarrow n_{out} - 1$ и се продължава от стъпка 2.

Понякога е желателно някои фактори, макар и не много значими, да участват в модела (поради изисквания, които нямат общо със статистическите характеристики на факторите). Тогава началният модел, от който стартира алгоритъмът, съдържа именно тези фактори.

Част от изискванията към модела се дефинират с помощта на статистически величини, характеризиращи неговото поведение. По тази причина същинският избор на това – кой модел (структура) е най-подходящ, се извършва след анализа на достоверността на модела в смисъл на тези характеристики. Въпреки нуждата от тази допълнителна дейност, като цяло процесът на избор на структурата е удобен за автоматизация.

3.5.5.3 Ефективна реализация на алгоритъма

Нека потенциалните фактори са 100, а системата е MISO. Това означава, че в първата итерация на стъпковия метод трябва да се формират 100 модела със свободен член и един фактор (от вида $\hat{y}_k = \theta_0 + \theta_1 \varphi_{1,k}$). Ако се стартират 100 оценители на параметрите, то данните ще бъдат прочетени 100 пъти. Съответно в следващата итерация моделите ще са 99 от вида $\hat{y}_k = \theta_0 + \theta_1 \varphi_{1,k} + \theta_2 \varphi_{2,k}$, а данните ще бъдат прочетени 99 пъти и т.н., докато стъпковият алгоритъм не се прекрати.

В общия случай, когато се прилага LS, оценките на модела може да се запишат като:

$$\hat{\theta} = (\Phi^T \Phi)^{-1} \Phi^T y = F^{-1} b. \quad (3.86)$$

Възможно е множествата модели от всички стъпки да се формират само с един прочит на данните, в началото на стъпковия алгоритъм. За целта се формира векторът $b_{full} = \Phi_{full}^T y \in \mathcal{R}^{p_{full}}$ и пълната матрица на Фишер $F_{full} = \Phi_{full}^T \Phi_{full} \in \mathcal{R}^{p_{full} \times p_{full}}$, зависещи от всички възможни, p_{full} на брой, фактори. Както се вижда, размерностите на F_{full} и b_{full} не зависят от броя на наблюденията, т.е. те са със значително по-малък брой елементи от $\Phi_{full} \in \mathcal{R}^{(N-n)\ell \times p_{full}}$. Това означава, че веднъж изчислени, може да останат заредени в оперативната памет или в бързата локална памет на видеокартата. След това, когато трябва да се определят оценките на модел с конкретни фактори, от F_{full} и b_{full} се извличат съответните редове и стълбове и се формират съответните подматрици F и b . След това оценките се определят от (3.86).

Например нека в началото на трета стъпка факторите, участващи в модела (от предходната стъпка), са φ_2 и φ_5 , а текущият съревноваващ се модел е с допълнителен фактор φ_1 . Тогава $F \in \mathcal{R}^{3 \times 3}$ съдържа елементите на 2-и, 5-и и 1-и ред и 2-и, 5-и и 1-и стълб на матрицата F_{full} , а $b \in \mathcal{R}^3$ съдържа 2-и, 5-и и 1-и елемент на вектора b_{full} . По този начин за формиране на F и b за всеки модел, а оттам и за определяне на съответните оценки $\hat{\theta}$, не е нужно да се прочитат входно-изходните данни.

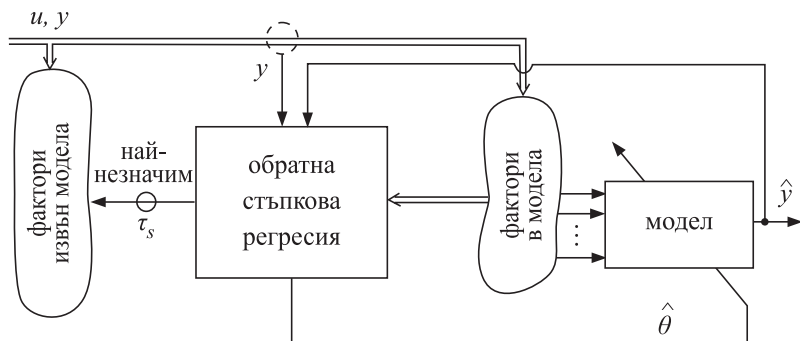
Възможни са допълнителни подобрения с избягване на явното обръщане на матриците на Фишер за съревноваващите се модели, при положение че е известна F^{-1} за модела от предишната итерация. Също така изчисляването на статистиките (например на $SSR_{+j}^{(i)}$ и $SSE_{+j}^{(i)}$) може да се извърши без обхождане на данните.

3.5.6 Обратна стъпкова регресия

Тази стъпкова регресия е подходяща, когато потенциалният брой фактори е сравнително малък (от порядъка на десетки фактори), тъй като тя стартира с пълния модел, след което структурата се опростява. Ако все пак тя е избрана, а факторите са много, то подходящо е да се приложат техники за редуциране на броя им преди стартиране на стъпковия алгоритъм. За целта в практиката се използват методи, като стъпков дискриминантен анализ [137, 136], клъстеризация и др. или

някои от описаните дейности, по време на предварителната обработка на данните.

3.5.6.1 Идея на метода



Фигура 3.15. Обратна стъпкова регресия

При обратната стъпкова регресия се стартира с пълния модел, а след това на всяка стъпка се премахват фактори (Фигура 3.15). Ако не съществува мултиколинеарност в данните, то първо се премахват най-малко информативните фактори, като на всяка следваща стъпка в модела остават все по-значими фактори. Изпълнението на стъпковия алгоритъм се прекратява, когато влошаването на модела от премахването на поредния фактор е по-голямо от допустимото. Прагът на значимост на факторите, оставащи в модела, е означен с τ_s и по подобие на τ_e в долния алгоритъм има смисъл на доверителната вероятност, над която влошаването в описанието на y_k при премахване на фактор се приема, че се дължи на случайността. Смисълът на τ_s се изяснява в представения по-долу алгоритъм.

3.5.6.2 Алгоритъм на метода

За добиване на по-ясна представа как се изпълнява обратната стъпкова регресия, отново е описан алгоритъм за представяне с вектор на параметрите. Също така моделите са MISO, защото при повече изходи всяка стъпка освен инициализиращата се изпълнява ℓ пъти.

1. Инициализация

Задава се прагът на значимост на факторите, оставащи в модела, означен с τ_s (често се означава като SLS – Significance Level to

Stay). Подходяща стойност за него е $\tau_s = 0.05$. Също се задават границите на изменение на изместването на сигналите във времето, тъй като от тях зависи структурата на пълния модел (с тяхна помощ се формират F_{full} и b_{full}).

2. Множество от съревноваващи се модели

Нека в i -тата итерация има n_{in} фактора в модела. Тогава се формират n_{in} съревноваващи се модели, които се получават от текущия модел (от предишната итерация) с премахване на един фактор. За i -тата итерация те са:

$$y_{-j}^{(i)} = \Phi_{-j}^{(i)} \theta_{-j}^{(i)}, \text{ за } j = \overline{1, n_{in}},$$

където матрицата на факторите за j -тия модел е

$$\Phi_{-j}^{(i)} = [1_{N-n} \quad \varphi_1^{(i-1)} \quad \dots \quad \varphi_{j-1}^{(i-1)} \quad \varphi_{j+1}^{(i-1)} \quad \dots \quad \varphi_{n_{in}}^{(i-1)}]. \quad (3.87)$$

Първият стълб на $\Phi_{-j}^{(i)}$ се състои от единици поради наличието на свободен член в модела. Така, ако в набора от данни няма значими фактори, алгоритъмът ще определи като най-добър празния модел, т.е. този, който се състои само от свободен член. Останалите стълбове в (3.87) са дефинирани в (3.84). Оценките на параметрите на j -тия съревноваващ се модел са:

$$\theta_{-j}^{(i)} = (\Phi_{-j}^{(i)T} \Phi_{-j}^{(i)})^{-1} \Phi_{-j}^{(i)T} y.$$

3. Значимост на потенциалните фактори

Критерият за значимост е

$$F_{-j}^{(i)} = \frac{v}{v_2 - v_1} \frac{SSR^{(i-1)} - SSR_{-j}^{(i)}}{SSE^{(i-1)}}, \quad (3.88)$$

където

$$SSR_{-j}^{(i)} = y_{-j}^{(i)T} y_{-j}^{(i)}, \quad SSE^{(i-1)} = e^{(i-1)T} e^{(i-1)}.$$

Величината $F_{-j}^{(i)}$ е частичната F статистика за модела и отразява влошаването му при премахването на фактор. Степените на свобода на текущия най-добър модел (с p параметъра), избран в $i-1$ -ата итерация, са $v_1 = N - p$, а на съревноваващите се модели (с $p-1$ параметъра), определени в i -та итерация, са $v_2 = N - (p-1)$. Знаменателят $v_2 - v_1$ отново е 1, понеже на всяка стъпка се премахва по един фактор. Ако факторите се премахват на групи, знаменателят е броят на факторите в премахнатата група.

Статистиката $F_{-j}^{(i)}$ се тълкува по същия начин като $F_{+j}^{(i)}$, опреде-

лена от (3.85), в смисъл че, колкото по-голяма е стойността $\hat{y}$, толкова по-значим е факторът (който в случая се изключва от модела).

4. Избор на най-незначимия фактор

Тук се определя най-ненужният (според F-критерия) фактор в модела от предишната итерация. Това е факторът, при който стойността на частичната F статистика е минимална. Нека търсеният фактор е с индекс

$$j^* = \arg \min_j F_{-j}^{(i)}.$$

5. Проверка за край

Статистиката $F^{(i)} = F_{-j^*}^{(i)}$ има F-разпределение и се използва за вземане на решение дали влошаването на модела с изключване на j^* -ти фактор е достатъчно незначимо (дължи се на случайността). За тази цел се прилага F-тестът (същият като при правата стъпкова регресия). Тук условието за край на стъпковия алгоритъм е

$$p^{(i)} < \tau_s.$$

Когато това условие се изпълни, то най-незначимият фактор се очаква, че е по-значим от допустимото (премахването му би влошило чувствително модела) и в такъв случай процедурата се прекратява. Приема се, че

$$\hat{\theta} \leftarrow \theta^{(i-1)}, \Phi \leftarrow \Phi^{(i-1)} \text{ и } p \leftarrow p_{full} - (i - 1).$$

В противен случай $i \leftarrow i + 1$ и $n_{in} \leftarrow n_{in} - 1$ и се продължава от стъпка 2.

3.5.7 Комбинирана стъпкова регресия

Правата и обратната стъпкови регресии са подходящи, когато потенциалните фактори зависят слабо един от друг (ако те са напълно независими, двата метода водят до един и същи резултат – това ще бъде изяснено по-долу).

С помощта на комбинираната стъпкова регресия се отчита евентуалната мултиколинеарност между факторите в модела. Темата за мултиколинеарността и линейната зависимост е разгледана в точка 2.2.3.3, от гледна точка на предварителната обработка на данните. По-долу ак-

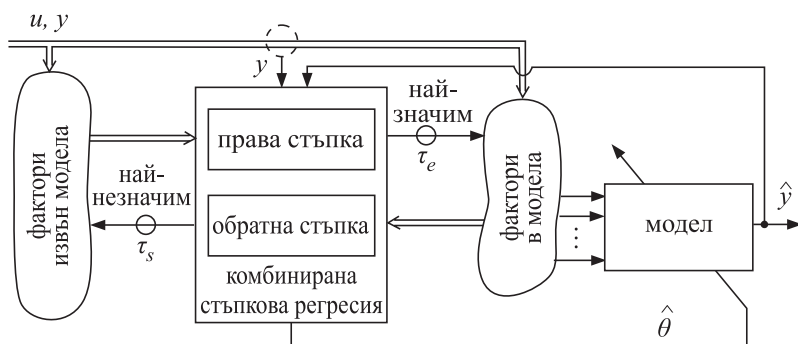
центът е върху отчитането на тази особеност в данните, но по време на уточняването на структурата на модела.

3.5.7.1 Мултиколинеарни фактори

В много реални ситуации потенциалните фактори съдържат обща информация за изхода на системата. При подготовката на данните е желателно да се премахнат линейно зависимите фактори. От друга страна, често се предпочита в набора да се запазят фактори, които не са напълно линейно зависими. Причината е, че на този предварителен етап, където акцентът са свойствата на данните, невинаги е ясно кои от близките фактори е най-добре да участват в крайния модел. Отговорът на този въпрос е естествено да бъде получен, когато се пристъпи към изграждане на описанието.

На Фигура 3.16 е представена идеята на комбинираната стъпкова регресия и най-вече как от множество фактори, сред които има мултиколинеарни, се извлича подмножество, в което всеки фактор допринася съществено за описанието на изхода на системата. При този стъпков метод формирането на модел с подходяща структура отново се постига с многократно оценяване на модели с различна структура, като се следи промяната в тяхното качество, но едновременно с това се отчита и евентуалната мултиколинеарност.

Тъй като мултиколинеарността е характерна за ММО системите, комбинираната стъпкова регресия се е утвърдила в различни приложни области като основен метод за определяне на структурата на модела.



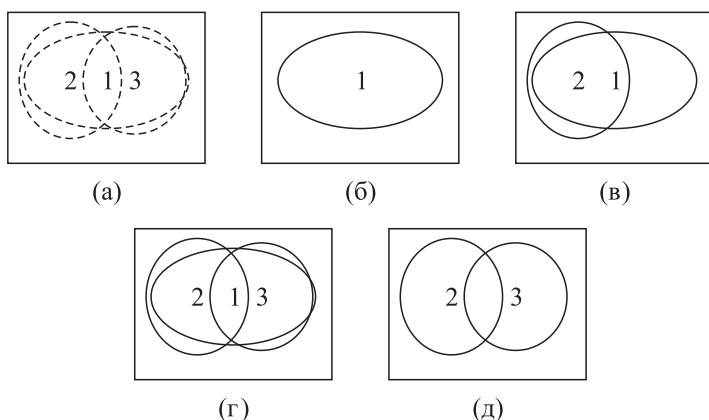
Фигура 3.16. Комбинирана стъпкова регресия

3.5.7.2 Идея на метода

При комбинираната стъпкова регресия се редуват двете стъпки, разгледани до момента, а именно:

- права стъпка (избор на фактор);
- обратна стъпка (елиминация на фактор),

както е показано на Фигура 3.16. Начинът, по който методът отчита евентуалната зависимост между факторите, е описан по-долу и представен графично на Фигура 3.17 с помощта на диаграмите на Вен.



Фигура 3.17. Нужда от права и обратна стъпка при зависими фактори

Нека са налице три фактора, които описват определени аспекти от поведението на даден изход. Площта на правоъгълника на Фигура 3.17 представлява информацията, съдържаща се в y , а площите на трите елипси (диаграма (а)) отразяват полезната информация във всеки фактор. При тази постановка, ако два фактора не са корелирани, те няма да имат обща площ (т.е. няма да съдържат обща информация, с която да обясняват зависимата променлива). От друга страна, ако два фактора са линейно зависими, то площите на съответните фигури напълно ще се припокриват. (Възможно е факторите да съдържат информация, която няма отношение към изхода, и тогава част от площите им ще са извън правоъгълника, но този случай не е от интерес за разглежданията.)

Нека стъпковата регресия стартира с празен модел. В процеса на работа на алгоритъма се генерират следните модели. Първо се изпълнява

правата стъпка, в резултат на което се избира фактор ‘1’ (случай (б)), тъй като той е най-информативен (покрива най-голяма площ). Понеже в модела няма незначими фактори, изпълнението на обратната стъпка не променя структурата му. При извършване на следващата права стъпка се добавя вторият фактор (случай (в)). Причината е, че при комбинирането му с първия фактор общата повърхност (информация) е по-голяма в сравнение с общата повърхност между първия и третия фактор. След това отново се прави опит за елиминация, който е неуспешен поради значимостта и на двата участващи в модела фактора. На третата стъпка, която е права, тъй като третият фактор съдържа значима информация, която не се припокрива с тази на влезлите до момента, в структурата на модела се включва и този фактор. Както се вижда от случай (г), площта на най-значимия фактор (добавен на първа стъпка) почти се припокрива от двата не толкова значими фактора. Това означава, че в текущия модел този фактор вече не допринася чувствително към описанието на y . Така, след добавянето на третия фактор, най-информативният (поотделно) фактор става незначим. Затова на следващата стъпка (обратна) първият фактор се премахва от модела (случай (д)).

3.5.7.3 Алгоритъм на метода

Представеният алгоритъм е комбинация от описаните алгоритми за права и обратна стъпкова регресия.

1. Инициализация

Задават се праговете τ_e и τ_s , както и границите на изменение на степените на полиномите (за динамични системи). С тях се дефинира множеството от n_{out} потенциални фактори. Ако има априорна информация за очаквани значими фактори, то те се добавят в началния модел и алгоритъмът продължава от стъпка 6, като се проверява дали някои от тези начални фактори са ненужни. Ако значимостта на факторите предварително не е известна, представеният алгоритъм стартира с празен модел и се продължава от стъпка 2.

2. Права регресия: съревноваващи се модели

Тази и следващите три точки се изпълняват, ако предишната стъпка е инициализираща (и в модела няма фактори, а само свободен член) или ако последната дейност е обратна регресия, която е неуспешна, т.е. няма фактори в модела, които да са незначими. В i -тата итерация се формират n_{out} съревноваващи се модели:

$$y_{+j}^{(i)} = \Phi_{+j}^{(i)} \theta_{+j}^{(i)}, \text{ за } j = \overline{1, n_{out}},$$

съдържащи включените до момента фактори и по един фактор (j -тият), който е извън модела. Оценките на параметрите на j -тия модел са

$$\theta_{+j}^{(i)} = (\Phi_{+j}^{(i)T} \Phi_{+j}^{(i)})^{-1} \Phi_{+j}^{(i)T} y.$$

3. *Права регресия: значимост на потенциалните фактори*

Изчислява се критерият за значимост:

$$F_{+j}^{(i)} = \frac{v}{v_1 - v_2} \frac{\text{SSR}_{+j}^{(i)} - \text{SSR}^{(i-1)}}{\text{SSE}_{+j}^{(i)}}.$$

4. *Права регресия: най-значим фактор*

От всички потенциални модели се избира този, при който се максимизира стойността на частичната F -статистика. Нека търсеният най-значим фактор измежду тези, които все още не са в модела, е с индекс

$$j^* = \arg \max_j F_{+j}^{(i)}.$$

5. *Права регресия: проверка за край*

Нека $F^{(i)} = F_{+j^*}^{(i)}$ и се изчисли съответната доверителна вероятност $p^{(i)}$. Прилага се F -тестът. Ако условието

$$p^{(i)} \leq \tau_e \quad (3.89)$$

е изпълнено и модел със същата структура като новоизбраната не е формиран до момента, за текущ модел се избира j^* -тият, а $i \leftarrow i + 1$ и $n_{out} \leftarrow n_{out} - 1$ и се продължава от стъпка 6 с проверка дали мултиколинеарността между факторите в новия модел е недопустимо висока. Ако новоизбраната структура е формирана и в предишна итерация, то моделът не се актуализира и алгоритъмът се прекратява. Така се избягва цикличното добавяне и премахване на едни и същи фактори.

Ако (3.89) не е изпълнено, то най-значимият фактор от тези извън модела не допринася чувствително за достоверността на описанието. Освен това в модела няма незначими фактори (защото последната обратна регресия е била неуспешна). В такъв случай моделът не се променя, а стъпковият алгоритъм се прекратява и се приема, че

$$\hat{\theta} \leftarrow \theta^{(i-1)}, \quad \Phi \leftarrow \Phi^{(i-1)}.$$

6. *Обратна регресия: съревноваващи се модели*

Тази и следващите три точки се изпълняват, ако последната (правата или обратната) регресия е успешна или алгоритъмът стартира с начален модел, съдържащ известен брой фактори. Ако последно е добавен фактор, то трябва да се провери дали някой от предишните фактори в модела е станал незначим. От друга страна, ако предишната регресия е обратна, то фактор е премахнат от модела и алгоритъмът продължава от тази точка с проверка дали има още фактори, които са незначими.

Формират се n_{in} съревноваващи се модели:

$$y_{-j}^{(i)} = \Phi_{-j}^{(i)} \theta_{-j}^{(i)}, \text{ за } j = \overline{1, n_{in}},$$

които се получават от текущия с премахване на един фактор. Оценките на параметрите са:

$$\theta_{-j}^{(i)} = (\Phi_{-j}^{(i)T} \Phi_{-j}^{(i)})^{-1} \Phi_{-j}^{(i)T} y.$$

7. *Обратна регресия: значимост на потенциалните фактори*

Изчислява се критерият за значимост:

$$F_{-j}^{(i)} = \frac{v}{v_2 - v_1} \frac{\text{SSR}^{(i-1)} - \text{SSR}_{-j}^{(i)}}{\text{SSE}^{(i-1)}}.$$

8. *Обратна регресия: най-незначим фактор*

Тук се определя най-ненужният (според F-критерия) фактор, измежду всички, които участват в текущия модел. Това е факторът, при който стойността на частичната F-статистика е минимална. Нека търсеният фактор е с индекс

$$j^* = \arg \min_j F_{-j}^{(i)}.$$

9. *Обратна регресия: проверка за край*

Нека $F^{(i)} = F_{-j^*}^{(i)}$ и се изчисли съответната доверителна вероятност $p^{(i)}$. Прилага се F-тестът. Ако условието

$$p^{(i)} \geq \tau_s. \quad (3.90)$$

е изпълнено и ако модел със същата структура като новоизбраната не е формиран в предишни стъпки, то за текущ модел се избира редуцираният с премахнат j^* -ти фактор, а $i \leftarrow i + 1$ и $n_{out} \leftarrow n_{out} + 1$ и се продължава от стъпка 6 с проверка дали мултиколинеарността между факторите в новия модел все още е недопустимо висока. Ако новоизбраната структура е формирана в

предишна итерация, то моделът не се актуализира и алгоритъмът се прекратява (избягвайки цикличното добавяне и премахване на едни и същи фактори).

Ако (3.90) не изпълнено, то моделът не се променя, $i \leftarrow i + 1$ и се продължава от стъпка 2 с проверка дали има фактори извън модела, които е удачно да се добавят.

Поведението на алгоритъма зависи от двата прага на значимост τ_e и τ_s . Колкото по-малки са те, толкова по-ниско е нивото на мултиколинеарност между факторите в крайния модел.

Когато моделът е ММО, при определянето на степените на свобода във формулата за частичната F-статистика е нужно за всеки изход да се следи текущият брой параметри. Причината е, че може по едни изходи да се добавят фактори, а в същата итерация, по други – да се премахват. Също е важно моделът да се проверява по всеки изход за преоразмеряване. При ММО модели, алгоритъмът спира, когато няма изход, чийто MISO модел да може да се подобри.

3.5.8 Групова стъпкова регресия

Понякога наборите от данни може да съдържат групи от фактори, които по време на моделирането се разглеждат като неделими множества. Например може да бъде наложено изискването, ако някои фиктивни променливи се включат в модела, то задължително да се добавят и останалите фиктивни променливи, описващи съответната променлива. Смисълът на това е, че наличието дори на една фиктивна променлива в модела предполага набавянето на данни за първоначалната променлива. Тогава (особено ако това е свързано с отделени средства) би трябвало цялата осигурена информация да участва в модела (при условие че не възникват числени проблеми). Важен ефект от работата с неделими групи от фактори е, че така намалява броят на първоначалните променливи в модела. Приложението на груповата стъпкова регресия в пазарния сектор е подробно обсъдено в [19, 18].

Заклучение

Монографията засяга идентификацията на системи, като акцентът е върху многомерността и динамиката. Разглеждането на тези две свойства на системите е естествен резултат от изискванията, наложени от много практически задачи, решаването на които е свързано с използването на модел. В основата на изискванията е стремежът да се построи възможно по-точно описание на изследвания аспект от реалността. Със съвременните технически средства използването на многомерен модел е напълно достъпно и не затруднява неговата употреба. От друга страна, отчитането на факта, че много величини са взаимноsvъrzани, води до качествено подобрене на крайния резултат от идентификацията (в сравнение с едномерния случай).

Въпреки че статичните системи са частен случай на динамичните, като цяло дейностите в идентификацията на динамични системи не могат да се разглеждат като обобщение на тези, използвани за статични системи. В монографията са описани и техники, характерни за изграждането на статични модели, които не са директно приложими при отчитането на динамика.

В области, където експерименталният подход е доминиращ, а моделите – по правило статични, вниманието все по-често се насочва към използването на динамични описания. Примери за това са различни сектори на икономиката като търговия, туризъм, финанси и т.н.

В монографията се разглеждат и решения за обработка на големи набори от данни. Въпреки че статистически от един момент, с увеличаване на наблюденията, не се постига подобрене в модела, има случаи когато, по една или друга причина не е желателно данните да се редуцират – например, ако нивото на неопределеността е високо или се наблюдават сегменти в данните, отразяващи различно поведение на системата (като режими на работа, различни групи от населението, продукти с характерно поведение на продажбите и др.). Като се добави голямата размерност на някои системи, както и малкото априорна информация, възниква нуждата идентификацията да се автоматизира. Така намеса-

та на експерт се измества от анализа на конкретни сигнали и вземане на рутинни решения, към формулиране на подходяща методология за извличане на полезната информация от данните, както и глобална настройка на дейностите, обхващани от идентификацията. Това позволява многократно да се премине през целия процес на моделиране, като последователно методологията се подобрява, а параметрите (например прагови стойности при автоматизираното вземане на решения) се прецизират. В резултат на това конструираната „поточна линия“ за изграждане на модел може да се използва за производството на бъдещи модели.

Друг аспект от работата с големи набори от данни, който е засегнат в труда, е ефективното използване на възможностите на изчислителните системи, например използването на графичните процесори на видеокартите. Също така тензорното представяне и подходящото съхраняване на разреждени матрици в паметта допълнително подобряват изпълнението на някои дейности в идентификацията.

Характерно за многомерните системи е възникването на числени проблеми. То се дължи на наличието на мултиколинеарност между факторите. Затова е отделено внимание и на числено устойчивите реализации на методите за оценяване на параметри.

В теоретично отношение в монографията е извършено обобщение на възможните представяния на регресионните модели в общ вид, когато те са или се разглеждат като линейно параметризирани. Едното представяне е параметрите да се групират в матрица ($\Theta \in \mathcal{R}^{z \times \ell}$), а регресорите във вектор. Другото е параметрите да са подредени във вектор ($\theta \in \mathcal{R}^p$), а регресорите в матрица. В зависимост от това, как са разпределени параметрите и регресорите в съответните матрици и вектори, е възможно да се получат различни варианти на общите представяния. Това води до промяна на свойствата на представянията или до по-удобна структура на модела за конкретно негово приложение.

При използването на описание с матрица на параметрите обикновено по-бързо се уточнява структурата му. Основният недостатък е, че има вероятност тази структура да е ненужно усложнена. Причината за това е, че подредбата на параметрите в Θ ограничава свободата при избора на структурата на модела. Това се дължи на изискването всички полиноми в коя да е полиномна матрица (или в неин стълб) да са от един и същи ред. Този недостатък се избягва, когато моделът се описва с вектор на параметрите и редът на всеки полином се задава независимо от останалите. Това представяне е подходящо за прецизно определяне на

структурата на модела.

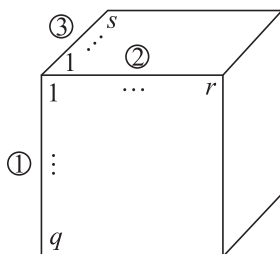
Когато структурата не е известна, не само в смисъл на степени на полиноми, а и на това – кои входни въздействия да участват в описанието, подходящо е първо да се използва представянето с матрица на параметрите. След първоначалното бързо, но по-грубо определяне на структурата е добре да се използва второто представяне, за да се доуточнят редовете на полиномите в модела. Също така понякога недостатъкът на общия вид с матрица на параметрите е желан ефект, например ако участието на всяка величина в модела е свързано с влагане на средства. Тогава, ако се вземе решение дадена величина да е фактор за описание на някой изход, то би трябвало от нея да се извлече максимална полза, като тя се добави в модела за описание и на другите изходи.

В монографията е разгледано получаването на модел, който е линейно параметризиран, а също и на такъв, който е нелинеен по параметри, но за целите на оценяването се приема за линеен. В отделен труд ще бъде разширена задачата на идентификацията, в т.ч. и оценяването на параметри и изборът на структура на модели, които са нелинейни по параметри.

Приложение

Умножения на тензор и матрица

Тензор е обобщено понятие, което включва в себе си понятията скалар, вектор и матрица. По-точно тензорът от нулев ред е скалар, от първи ред е вектор, а от втори ред – матрица. Когато размерностите на тензора са три, той е от трети ред и т.н.



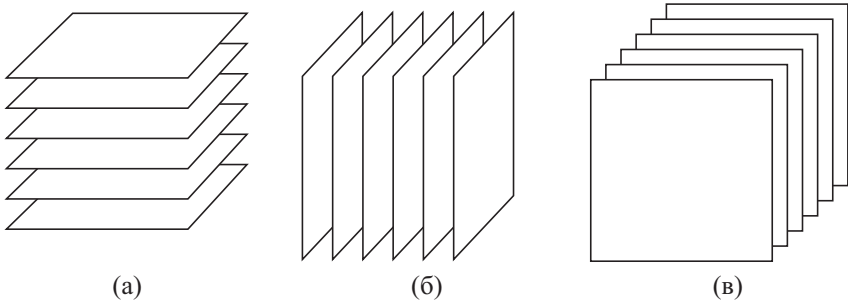
Фигура П.1. Тензор от трети ред $T \in \mathcal{R}^{q \times r \times s}$ и размерности

Съществуват различни умножения, свързани с тензори. В монографията се извършват два типа умножение между тензор и матрица, като едното е стандартното, а другото е въведено за опростяване на записа на някои изрази. Затова в приложението е обърнато внимание на тези две действия.

Стандартно умножение на тензор и матрица

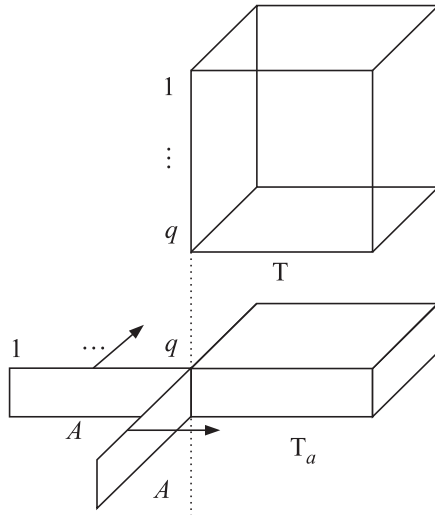
При матриците начинът, по който се извършва умножението, се определя еднозначно: елементите на резултантната матрица се получават в резултат от скаларното умножение на редовете на първата матрица и стълбовете на втората. От друга страна, при умножението между тензор с повече от две размерности и матрица е нужно да се уточни спрямо коя размерност на тензора се извършва умножението.

Нека T е тензор от трети ред с размерност $q \times r \times s$ (Фигура П.1), с реални елементи, което се записва накратко като $T \in \mathcal{R}^{q \times r \times s}$. Той се състои от q на брой $r \times s$ мерни матрици, а може да се разглежда и като



Фигура П.2. Тензорът $T \in \mathcal{R}^{q \times r \times s}$ може да се разглежда като обект, състоящ се от матриците $T_{i_1..} \in \mathcal{R}^{r \times s}$, за $i_1 = \overline{1, q}$ (а), от матриците $T_{.i_2.} \in \mathcal{R}^{q \times s}$, за $i_2 = \overline{1, r}$ (б) или от матриците $T_{..i_3} \in \mathcal{R}^{q \times r}$, за $i_3 = \overline{1, s}$ (в)

обект, състоящ се от r на брой $q \times s$ мерни матрици или от s на брой $q \times r$ мерни матрици, както е показано на Фигура П.2.



Фигура П.3. Умножението $T_a = T \circ_1 A$ на тензор от трети ред и матрица. Представяне с обхождане по втората и по третата размерност на T .

В монографията стандартното умножение на тензора T с матрицата M [47] спрямо неговата i -та размерност се означава като

$$T \circ_i M. \quad (3.91)$$

Прието е винаги тензорът да се записва пръв. При умножение на две матрици втората размерност на първата матрица съвпада с първата размерност на втората, например при умножението XY между матриците $X \in \mathcal{R}^{m \times n}$ и $Y \in \mathcal{R}^{p \times q}$ е нужно $n = p$. По-същия начин при записа (3.91) задължително втората размерност на матрицата трябва да съвпада с i -тата размерност на T .

Нека $A \in \mathcal{R}^{h \times q}$, $B \in \mathcal{R}^{h \times r}$, $C \in \mathcal{R}^{h \times s}$. Тензорът $T_a \in \mathcal{R}^{h \times r \times s}$, получен от умножението на T с A по неговата първа размерност, се записва така:

$$T_a = T \circ_1 A,$$

(ако $D \in \mathcal{R}^{q \times h}$, то $T_d = T \circ_1 D^T$, като D е транспонирана, понеже втората размерност на D^T е q). Резултатът от умножението отново е тензор, за който са в сила равенствата:

$$T_{a, \dots i_2} = AT_{\dots i_2}, \text{ за } i_2 = \overline{1, r}, \quad (3.92)$$

($T_{\dots i_2}$ е матрица с размерност $q \times s$). Те са еквивалентни на

$$T_{a, \dots i_3} = AT_{\dots i_3}, \text{ за } i_3 = \overline{1, s}. \quad (3.93)$$

Тези действия графично са показани на Фигура П.3. По подобен начин, тензорът $T_b \in \mathcal{R}^{q \times h \times s}$, получен от умножението на T с B по неговата втора размерност, се записва така:

$$T_b = T \circ_2 B.$$

Това действие също е възможно, защото втората размерност на T е същата като втората на B . Резултатът отново е тензор, за който е в сила:

$$T_{b, i_1 \dots} = BT_{i_1 \dots} \Leftrightarrow T_{b, \dots i_3} = T_{\dots i_3} B^T. \quad (3.94)$$

Съответно тензорът $T_c \in \mathcal{R}^{q \times r \times h}$, получен от умножението на T с C по неговата трета размерност, е

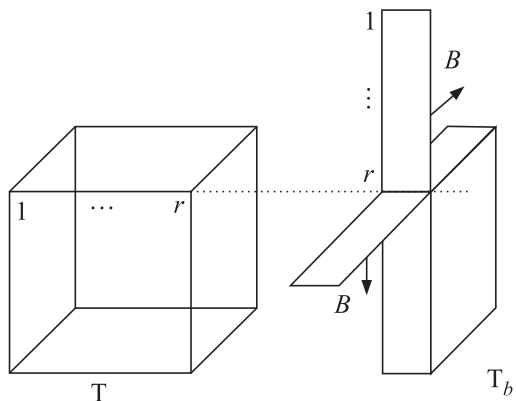
$$T_c = T \circ_3 C,$$

като е в сила:

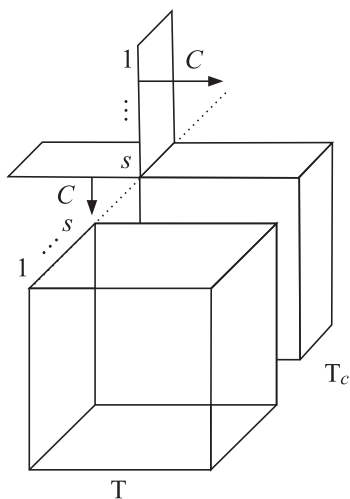
$$T_{c, i_1 \dots} = T_{i_1 \dots} C^T \Leftrightarrow T_{c, \dots i_2} = T_{\dots i_2} C^T. \quad (3.95)$$

Тук отново е изпълнено необходимото съответствие между размерностите, в смисъл че този път третата размерност на T съвпада с втората на C .

На Фигура П.4 и П.5. графично са показани съответно действията (3.94) и (3.95).



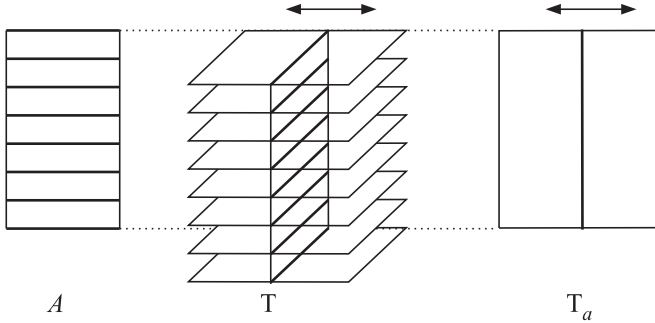
Фигура П.4. Умножението $T_b = T \circ_2 B$ на тензор от трети ред и матрица. Представяне с обхождане по първата и по третата размерност на T .



Фигура П.5. Умножението $T_c = T \circ_3 C$ на тензор от трети ред и матрица. Представяне с обхождане по първата и по втората размерност на T .

В някои случаи описаното умножение се извършва между тензора и няколко матрици по различни негови размерности. Например, ако T се умножи с A по първата размерност и с B по втората, резултатът е $T_{ab} \in \mathcal{R}^{h \times h \times s}$. Действието се записва така:

$$T_{ab} = T \circ_1 A \circ_2 B.$$



Фигура П.6. Свиващото повекторно умножение $T_a = T \circ_3^1 A$ на тензор от трети ред и матрица

Свиващо повекторно умножение на тензор и матрица

Другият тип умножение е въведено в монографията за опростяване на изложението. То е т.нар. свиващо повекторно умножение на тензор и матрица и е показано на Фигура П.6. Наречено е „свиващо“, понеже резултатът от него е тензор, на който една от размерностите се редуцира до 1. Повекторно е, тъй като елементите на резултантния тензор са получени от повекторно (скаларно) умножение на стълбовете на матрицата и определени вектори на тензора.

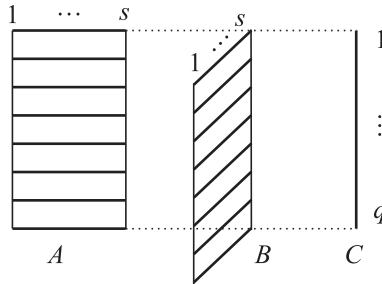
Преди да се дефинира това умножение, нека със символа ‘ $\bullet$ ’ се означава повекторното умножение на матрици. Така, ако $A \in \mathcal{R}^{s \times q}$ и $B \in \mathcal{R}^{s \times q}$, то $c = A \bullet B$ е вектор с елементи, които са скаларните произведения на стълбовете на A и B . С други думи

$$c_j = A_{\cdot j}^T B_{\cdot j}, \text{ за } j = \overline{1, q},$$

както е показано на Фигура П.7. По този начин може да се запише действието:

$$T_{a, i_2 1} = A \bullet T_{\cdot i_2 \cdot}, \text{ за } i_2 = \overline{1, r}, \quad (3.96)$$

в резултат на което се получава векторът $T_{a, i_2 1} \in \mathcal{R}^q$.



Фигура П.7. Повекторно скалярно умножение на матрици.

Размерностите на матрицата трябва да съвпадат с две от размерностите на тензора. Тъй като $T \in \mathcal{R}^{q \times r \times s}$, а оттам $T_{i_2..} \in \mathcal{R}^{q \times s}$ и понеже умножението се извършва спрямо третата размерност, която в (3.96) се редуцира до 1, то е необходимо A да е $A \in \mathcal{R}^{q \times s}$. Отново тензорът се записва пръв, а матрицата така, че втората ѝ размерност да е равна на размерността на тензора, отговаряща на броя на векторите му, които се умножават по тези на A . Описаното действие, което е обобщение на (3.96), се записва като:

$$T_a = T \circ_3^1 A. \quad (3.97)$$

Долният индекс е размерността, по която тензорът се свива (до 1), а горният индекс е размерността, определяща спрямо коя размерност се извършва свиването. Както е показано на Фигура П.6, тази размерност е първата (i_1 -тият стълб на A се умножава по всички вектори на матрицата $T_{i_1..} \in \mathcal{R}^{r \times s}$).

Самостоятелно тензорът T_a може да се разглежда като матрица, но когато участва в израз, понякога е важна неговата пълна размерност.

Означения

Малки латински букви

$a_{ij}(q^{-1})$	полиномът, отразяващ влиянието на j -тия изход върху i -тия (ij -ти елемент на $A(q^{-1})$); за SISO системи е $a(q^{-1})$
a_{ij}	векторът на параметрите на ij -тия полином в матрицата $A(q^{-1})$
$a_{i,j}$	векторът на параметрите пред степента q^{-j} на полиномите от i -тия стълб на матрицата $A(q^{-1})$
$a_{ij,h}$	h -тият параметър на ij -тия полином на $A(q^{-1})$
b	вектор в общ вид
$b_{ij}(q^{-1})$	полиномът, отразяващ влиянието на j -тия вход върху i -тия изход (ij -ти елемент на $B(q^{-1})$); за SISO системи е $b(q^{-1})$
b_{ij}	векторът на параметрите на ij -тия полином в $B(q^{-1})$
$b_{i,j}$	векторът на параметрите пред степента q^{-j} на полиномите от i -тия стълб на матрицата $B(q^{-1})$
$b_{ij,h}$	h -тият параметър на ij -тия полином на $B(q^{-1})$
c	скалар в общ вид (обикновено константа)
$c_{ij}(q^{-1})$	полиномът, отразяващ влиянието на j -тия остатък върху i -тия изход (ij -ти елемент на $C(q^{-1})$); за SISO системи е $c(q^{-1})$
c_{ij}	векторът на параметрите на ij -тия полином в $C(q^{-1})$
$c_{ij,h}$	h -ти параметър на ij -тия полином на $C(q^{-1})$
d	диференциал; закъснение

$d_{i,k}$	фиктивна променлива, описваща (не)съвпадението на стойността на неметрична променлива в k -тия такт с нейната i -та възможна стойност
e	$\ell(N - n)$ -мерният вектор от стойностите на остатъка за интервала на наблюдение (за статични системи е ℓN -мерен); остатъкът в общ вид; експонента
e_k	ℓ -мерният вектор на остатъка в k -тия такт
$e_{ij,k}$	nc_{ij} -мерният вектор на регресорите в k -тия такт, свързани с ij -тия полином в $C(q^{-1})$
$\tilde{e}_k$	остатъкът, когато е цветен шум в k -тия такт
f	функция в общ вид
f_s	честотата на дискретизация в Hz
g	градиентът на целевата функция $\nabla \mathcal{F}(\theta)$
h	помощен индекс
i	индекс (обикновено съответства на текущия изход); индекс на итерация (когато е горен индекс в скоби)
j	индекс (обикновено съответства на пореден фактор в рамките на полиномна матрица)
k	дискретен момент от времето (такт)
$l(\theta, \tilde{\theta})$	функцията на загубите в Бейсовия метод
ℓ	броят изходи на модела
m	броят входове на модела
n	максималният ред на полином в модела
na, nb, nc	максималните степени на полиномите съответно в $A(q^{-1})$, $B(q^{-1})$ и $C(q^{-1})$
na_i	максималната степен на полиномите в i -тия стълб на $A(q^{-1})$ (аналогично nb_i за $B(q^{-1})$ и т.н.)
na_{ij}	степената на ij -тия полином на $A(q^{-1})$ (аналогично nb_{ij} за $B(q^{-1})$ и т.н.)

n_{in}	брой фактори в модела (при стъпковите методи за избор на структура)
n_{out}	брой фактори извън модела (при стъпковите методи за избор на структура)
p	броят на параметрите на модела (p_i – брой на параметрите на i -тия MISO модел)
p	доверителната вероятност в тестовете за значимост на фактор в стъпковите алгоритми
$p(\theta)$	плътността на разпределение на параметрите
$p(\theta y)$	условната (апостериорна) плътност на разпределение на параметрите
$p(y \theta)$	функцията на правдоподобие (плътността на вероятността за наблюдаване на данните при фиксирани параметри на модела)
q	оператор за отместване на сигнал във времето ($x_k q^{-i} = x_{k-i}$)
r	ранг на матрица; помощен индекс
s	помощен индекс
s_x^2	оценка на дисперсията на x
s_i	i -тата сингулярна стойност
t	непрекъснато време
u	входът на обекта в общ вид
u_k	m -мерният вектор от входовете на обекта в k -тия такт
$u_{ij,k}$	nb_{ij} -мерният вектор на регресорите, свързани с ij -тия полином в $B(q^{-1})$
w_k	ℓ -мерният тегловен вектор в WLS в k -тия такт; системната неопределеност
x	сигнал в общ вид (скаларен или векторен)

y	$\ell(N - n)$ -мерният вектор от стойностите на изхода на обекта за интервала на наблюдение (за статични системи е ℓN -мерен); изходът на обекта в общ вид
y_k	ℓ -мерният вектор от изходите на обекта в k -тия такт
$y_{ij,k}$	na_{ij} -мерният вектор на регресорите, свързани с ij -тия полином в $A(q^{-1})$
$\hat{y}_k$	изходът на модела
z	броят на регресорите в модела

Главни латински букви

A	матрица в общ вид
$A(q^{-1})$	$\ell \times \ell$ мерната полиномна матрица, отразяваща авто-регресията по изхода y_k
A_i	$\ell \times \ell$ мерната матрица от коефициентите на полиномите в $A(q^{-1})$ пред степента q^{-i}
AIC	информационният критерий на Акайке (Akaike Information Criterion)
$B(q^{-1})$	$\ell \times m$ мерната полиномна матрица, отразяваща влиянието на предисторията на входа u_k
B_i	$\ell \times m$ мерната матрица от коефициентите на полиномите в $B(q^{-1})$ пред степента q^{-i}
BIC	Бейсовият информационен критерий (Bayesian Information Criterion)
$C(q^{-1})$	$\ell \times \ell$ мерната полиномна матрица, отразяваща влиянието на предисторията на остатъка e_k
C_i	$\ell \times \ell$ мерната матрица от коефициентите на полиномите в $C(q^{-1})$ пред степента q^{-i}
Cp	статистиката на Малоу (Malloy Statistics или Malloy Cp)
D	индексът на различие (Diversity Index)

E	$N - n \times \ell$ -мерната матрица от стойностите на остатъка за интервала на наблюдение (за статични системи е $N \times \ell$ -мерна матрица)
E_{-i}	$N - n \times \ell$ -мерната матрица от стойностите на остатъка от $n + 1 - i$ -тия до $N - i$ -тия такт за интервала на наблюдение
F	матрицата на Фишер
F	тензорът от матриците на Фишер; F-разпределение
F_{-j}	частичната F-статистика, отразяваща влошаването на модела при премахване на j -тия фактор
F_{+j}	частичната F статистика, отразяваща подобрението на модела при добавяне на j -тия фактор
$\mathcal{F}$	целевата функция ($\mathcal{F}(\Theta)$ или $\mathcal{F}(\theta)$)
G	градиентната матрица $\nabla \mathcal{F}(\Theta)$
H	Хесианът $\nabla^2 \mathcal{F}(\theta)$
I	единична матрица (I и I_n)
J	Якобианът на грешката
L_θ	функцията на правдоподобие
$L_{\theta,k}$	приносът на k -тото наблюдение към функцията на правдоподобие
$L_y(\omega)$	логаритмичната амплитудночестотна характеристика
M	матрица в общ вид
$M(q^{-1})$	полиномна матрица в общ вид
$M\{\}$	операторът за математическо очакване
N	броят наблюдения
$\mathcal{N}$	нормално разпределение
O_n	членове от n -ти и по-висок ред
P	ковариационната матрица на параметрите

P	тензорът от ковариационните матрици
Q	ортогоналната матрица в QR декомпозицията
R	горнотриъгълната матрица в QR декомпозицията
R_x	автокорелационната функция на скаларния сигнал x
R_{xy}	вазимокорелационната функция на сигналите x и y
R^2	коефициентът на определеност/детерминация (coefficient of determination)
R_a^2	коригираният коефициент на определеност (adjusted R^2)
$\mathcal{R}$	множеството на реалните числа
$\mathcal{R}^n$	множеството на реалните n -мерни вектори
$\mathcal{R}^{m \times n}$	множеството на $m \times n$ -мерните матрици
$\mathcal{R}_s^{m \times n}(q^{-1})$	множеството на $m \times n$ -мерните полиномни матрици с полиноми по степените на q^{-1} от ред максимум s
S	диагоналната матрица, съдържаща сингулярните стойности (в SVD декомпозицията)
S_x	диагоналната матрица, съдържаща стандартните отклонения на компонентите на вектора x
SC	Шварц критерият (Schwarz Criterion), наричан още Бейсов информационен критерий (BIC)
SSE	сумата от квадратите на разликата между изхода на обекта и на модела (Error Sum of Squares)
SSR	сумата от квадратите на центрирания изход на модела (Regression Sum of Squares)
SST	сумата от квадратите на центрирания изход на обекта (Total Sum of Squares)
T	тензор в общ вид
T_s	период на дискретизация
U	ортогонална матрица в SVD декомпозицията

U_{-i}	$N - n \times m$ -мерната матрица от стойностите на входа от $n + 1 - i$ -тия до $N - i$ -тия такт за интервала на наблюдение
V	ортогонална матрица в SVD декомпозицията
VAF	показателят на отчетената вариация (Variance Accounted For)
VIF	показател, наречен фактор на изкуствено нарастване на вариацията (Variance Inflation Factor)
VoI	показател, наречен степен на информативност (Value of Information)
W	тегловната матрица в WLS
W	тегловният тензор в WLS
WoE	теглото на достоверност (Weighted of Evidence)
X	матрица в общ вид
Y	матрица в общ вид
Y	$N - n \times \ell$ -мерната матрица от стойностите на изхода за интервала на наблюдение (за статични системи е $N \times \ell$ -мерна матрица)
Y_{-i}	$N - n \times \ell$ -мерната матрица от стойностите на изхода от $n + 1 - i$ -тия до $N - i$ -тия такт за интервала на наблюдение
Z	инструменталната матрица в IV
$\mathcal{Z}$	множеството на целите числа

Малки гръцки букви

α (alpha)	скалар в общ вид (обикновено приемащ големи стойности); тъгъл
β (beta)	тъгъл
γ (gamma)	градиентът на приноса към целевата функция в робастен оценител

ε (epsilon)	случайна величина
ζ (zeta)	приносът към целевата функция в робастен оценител
η (eta)	входната неопределеност ($\eta \in \mathcal{R}^m$ е разликата между управлението, изчислено от регулатора, и въздействието, формирано от регулиращия орган; смущения от околната среда, навлизащи по канала на управлението; неопределеност, приведена към входа)
θ (theta)	p -мерният вектор на параметрите на модела
$\hat{\theta}$	векторът на оптималните оценки на параметрите
$\tilde{\theta}$	разликата между оптималните параметри и техните оценки
θ^*	оптималните (неизвестни) параметри
θ_i	i -тият елемент на вектора θ
θ_a	априорният вектор на параметрите
θ_i	p_i -мерният вектор от параметрите на i -тия MISO модел
λ (lambda)	параметърът в регуляризацията на Тихонов
ν (nu)	изходната неопределеност ($\nu \in \mathcal{R}^\ell$ е шум от измерване на изхода и смущения, влияещи директно на изхода)
ξ (ksi)	неопределеността в системата, приведена към изхода ($\xi \in \mathcal{R}^\ell$ не зависи от конкретен модел за разлика от e)
π (pi)	константата пи
ρ (rho)	корелационен коефициент (коефициент на Пийърсън) в общ вид
ρ_x	корелационният коефициент на векторния сигнал $x \in \mathcal{R}^m$ ($\rho_x \in \mathcal{R}^{m \times m}$ – симетрична матрица)
ρ_{xy}	корелационният коефициент между x и y (ако $x \in \mathcal{R}^m$ и $y \in \mathcal{R}^n$, то $\rho_{xy} \in \mathcal{R}^{m \times n}$)

$\rho_{y x}$	множественият корелационен коефициент, отразяващ връзката между векторния сигнал $x \in \mathcal{R}^m$ и скаларния сигнал y (друг запис е $\rho_{y x_1, \dots, x_m}$)
σ (sigma)	стандартно отклонение в общ вид
σ_x	стандартното отклонение на сигнала x
σ_x^2	дисперсията на x (когато x е скалар, вектор или матрица и σ_x^2 е съответно скалар, вектор или матрица)
τ (tau)	толеранс (например $\tau_{\mathcal{F}}$, τ_{θ} , τ_e , τ_s и др.) в условието за спиране на итеративните алгоритми и стъпковата регресия
v (upsilon)	степени на свобода
φ_k (phi)	векторът на регресорите (факторите) в k -тия такт
$\varphi_{i,k}$	векторът на регресорите на i -тия MISO модел в k -тия такт
χ^2	хи-квадрат разпределение
ω (omega)	честотата в rad/s на хармоник
ω_s	честотата на дискретизация

Главни гръцки букви

Θ	$z \times \ell$ -мерната матрица на параметрите
$\hat{\Theta}$	матрицата на оптималните оценки на параметрите
$\tilde{\Theta}$	разликата между оптималните параметри и техните оценки
Θ^*	оптималните (неизвестни) параметри
Σ	ковариационна матрица в общ вид
Σ_x	$m \times m$ -мерната ковариационна матрица на векторния сигнал $x \in \mathcal{R}^m$
Σ_{xy}	$m \times n$ -мерната ковариационна матрица на векторните сигнали $x \in \mathcal{R}^m$ и $y \in \mathcal{R}^n$

Σ_x	$N-n \times N-n$ или $\ell(N-n) \times \ell(N-n)$ -мерната ковариационна матрица на векторния сигнал x за интервала на наблюдение
Φ	$N-n \times z$ -мерната матрица на данните (факторите, регресорите)
Φ_k	$\ell \times p$ -мерната матрица на регресорите в k -тия такт
Ω	множеството на оптималните оценки на параметрите при наличие на линейно зависими фактори

Други символи

$\diamond$	край на пример
∂	частен диференциал
$\approx$	е приблизително равно на
$\propto$	е пропорционално на
$\equiv$	съвпада с
$\neq$	е различно от
$\in$	принадлежи на
$\subseteq$	е подмножество на или съвпада с
$\subset$	е подмножество на
$\exists$	съществува
$\Leftrightarrow$	е еквивалентно на
$\leftarrow$	се замества с
$\times$	умножение, задаване на размерност на матрици и тензори
∞	безкрайност
$\otimes$	умножение на Кронекер

$*$	поелементно умножение на вектори и матрици (например, ако $A, B \in \mathcal{R}^{m \times n}$, то $A * B = C \in \mathcal{R}^{m \times n}$, като $c_{ij} = a_{ij}b_{ij}$)
$\parallel$	поелементно деление на вектори и матрици (например, ако $A, B \in \mathcal{R}^{m \times n}$, то $A \parallel B = C \in \mathcal{R}^{m \times n}$, като $c_{ij} = \frac{a_{ij}}{b_{ij}}$)
$\circ_i$	умножение между тензор и матрица, спрямо неговата i -та размерност (виж Приложението)
$\odot_i^j$	свиващо повекторно тензорно умножение (виж Приложението)
$\bullet$	повекторно умножение на матрици (виж приложението)
∇	операторът (набла) $\nabla_x = [\frac{\partial}{\partial x_1} \quad \frac{\partial}{\partial x_2} \quad \dots \quad \frac{\partial}{\partial x_n}]$ за дефиниране на градиент (ако аргументът е матрица, то $[\nabla_x]_{ij} = \frac{\partial}{\partial x_{ij}}$)
Σ	сума
Π	произведение
1_n	вектор от единици с размерност n
$0_{m \times n}$	нулева матрица с размерност $m \times n$
A^T	транспонираната матрица на матрицата A
$A_{i.}$	i -тият ред на матрицата A ($A_{.j}$ е j -ти стълб на A). Ако A е тензор, то $A_{i1.}$, $A_{.i2.}$, $A_{..i3}$ са матрици, а $A_{i1i2.}$, $A_{i1.i3}$, $A_{.i2i3}$ са вектори.
$[A]_{ij}$	ij -тият елемент на матрицата A . Ако A е тензор, то $[A]_{i1i2i3}$ е $i1i2i3$ -тият елемент на тензора.
$\ x\ _2$	2-нормата на вектора x (например $\ x\ _2 = \sqrt{x^T x} = \sqrt{x_1^2 + x_2^2 + \dots + x_n^2}$)
$\ A\ _F$	Фробениус нормата на матрица A (например $\ A\ _F = \sqrt{a_{11}^2 + a_{12}^2 + \dots + a_{nn}^2}$)
$\ A\ _W$	претеглената 2-норма на вектора x (например $\ x\ _W = \sqrt{x^T W x}$)

$\lfloor a \rfloor$	най-голямото цяло число, по-малко или равно на a
$ a $	абсолютната стойност на a
$\text{cond } A$	числото на обусловеност на A
const	константа
$\text{cov}(x, y)$	взаимната ковариация между сигналите x и y (ако $x \in \mathcal{R}^m$ и $y \in \mathcal{R}^n$, то $\text{cov}(x, y) \in \mathcal{R}^{m \times n}$)
$\text{cov } x$	автоковариацията на векторния сигнал x
$\det A$	детерминантата на A
$\text{diag } A$	векторът (стълб) с елементи, равни на елементите по диагонала на A
$\text{diag } a$	диагоналната матрица с елементи по диагонала, равни на елементите на вектора a
$\text{diag}(a_1, \dots, a_n)$	диагоналната матрица с елементи по диагонала, равни на скаларите a_i
$\text{diag}(A_1, \dots, A_n)$	блокдиагоналната матрица с блокове по диагонала, равни на матриците A_i
$\text{mean } x, \bar{x}$	средната стойност на x (когато x е скаларен, векторен или матричен сигнал и $\bar{x}$ е съответно скалар, вектор или матрица)
$\lim$	граница на функция
$\text{median } x$	медианата на x
$\text{mode } x$	модата на x
$\text{rank } A$	рангът на A
$\text{range } A$	образът на A (пространството, в което стълбовете на A формират базис)
$\text{sum } A$	сумата от всички елементи на вектора/матрицата A
$\text{tr } A$	следата на A (сума от диагоналните елементи)
$\text{vec } A$	векторът със структура $\text{vec } A = [A_{:,1}^T \ A_{:,2}^T \ \dots \ A_{:,n}^T]^T$, за $A \in \mathcal{R}^{m \times n}$

Съкращения

БШ	бял шум
ДБШ	дискретен бял шум
ДУ	диференциално уравнение
ИИД	извличане на информация от данните (Data mining)
ЛАЧХ	логаритмична амплитудночестотна характеристика
ЦШ	цветен шум
AR	авторегресионен модел (AutoRegressive model)
ARARMAX	авторегресионен модел с външен входен сигнал и формиращ филтър от тип авторегресия пълзящо средно (AutoRegressive model with eXogenous input with AutoRegressive Moving Average noise model)
ARARX	авторегресионен модел с външен входен сигнал и формиращ филтър от тип авторегресия (AutoRegressive model with eXogenous input with AutoRegressive noise model)
ARMAX	авторегресионен модел с външен входен сигнал и формиращ филтър от тип пълзящо средно (AutoRegressive model with eXogenous input with Moving Average noise model)
ARX	авторегресионен модел с външен входен сигнал (AutoRegressive model with eXogenous input)
BLUE	най-добър линеен неизместен оценител (Best Linear Unbiased Estimator)
CPU	централен процесор (Central Processing Unit)
ELS	метод на разширените най-малки квадрати (Extended Least Squares)
EVD	декомпозиция по собствени стойности (Eigenvalue Decomposition)

GLS	метод на обобщените най-малки квадрати (Generalized Least Squares)
GPU	графичен процесор (Graphical Processing Unit)
IV	метод на инструменталните променливи (Instrumental Variable)
LOO	K -кратна кросвалидация – граничен случай (Leave One Out)
LS	метод на най-малките квадрати (Least Squares)
MA	модел от тип пълзящо средно (Moving Average)
MAP	метод на максималната апостериорна плътност (Maximum A-Posteriori)
MDS	софтуер за експериментално моделиране на компанията „Експиритън“ (Model Development Studio)
ML	метод на максималното правдоподобие (Maximum Likelihood)
NGLS	метод на нелинейните обобщени най-малки квадрати (Nonlinear Generalized Least Squares)
NLS	нелинеен метод на най-малките квадрати (Nonlinear Least Squares)
NWLS	метод на нелинейните претеглени най-малки квадрати (Nonlinear Weighted Least Squares)
QR	QR декомпозиция
R	статистически софтуер (от първата буква на създателите му Robert Gentleman и Ross Ihaka)
RobLS	робастен метод на най-малки квадрати (Robust Least Squares)
SAS	статистически софтуер (Statistical Analysis System)
SLE	прагът на значимост за добавяне на фактори τ_e в стъпков алгоритъм за избор на структурата (Significance Level to Enter)
SLS	прагът на значимост за премахване на фактори τ_s в стъпков алгоритъм за избор на структурата (Significance Level to Stay)
SPSS	статистически софтуер (Statistical Product and Service Solutions)
SVD	декомпозиция по сингулярни стойности (Singular Value Decomposition)
WLS	метод на претеглените най-малки квадрати (Weighted Least Squares)

Библиография

- [1] Айдемирски, П. *Изследване на методи и алгоритми за адаптивна обработка на сигнали*. Дисертация. Българска академия на науките, София, 1991.
- [2] Божанов, Е. и И. Вучков. *Статистически методи за моделиране и оптимизиране на многофакторни обекти*. Техника, София, 1973.
- [3] Бокс, Дж. и Г. Дженкинс. *Анализ временных рядов. Прогноз и управление*. Мир, 1974. Част 1 и част 2, превод от английски.
- [4] Вентцель, Е. С. и Л. А. Овчаров. *Теория вероятностей и ее инженерные приложения*. Высшая школа, Москва, второ издание, 2000.
- [5] Вучков, И. *Идентификация*. ИК Юрапел, София, 1996.
- [6] Вучков, И., Н. Козарев, С. Стоянов и В. Цочев. *Ръководство за лабораторни упражнения по статистически методи*. Нови знания, първо издание, 2002.
- [7] Въндев, Д. *Записки по теорията на вероятностите*, януари, 2002.
- [8] Въндев, Д. *Записки по приложна статистика 1*, Софийски университет „Св. Климент Охридски“, ФМИ, юни, 2003.
- [9] Въндев, Д. *Записки по приложна статистика 2*, Софийски университет „Св. Климент Охридски“, ФМИ, юни, 2003.
- [10] Гарилов, Е. *Идентификация на системи. Част I. Въведение*. Технически университет – София, второ преработено издание, 2004.
- [11] Гарилов, Е. *Идентификация на системи. Част II. Идентификация чрез дискретни стохастични регресионни модели*. Технически университет – София, второ преработено издание, 2004.
- [12] Гатев, Г. *Анализ и синтез на автоматични системи*. Техника, София, 1978.
- [13] Генчев, Т. *Обикновени диференциални уравнения*. Университетско издателство „Св. Климент Охридски“, трето издание, София, 1999.
- [14] Георгиев, В. Ц. *Създаване на модели и софтуер, приложими при обработка на данни от селскостопански изследвания*. Докторантура. Пловдив, 2004.

- [15] Джиев, С. *Моделиране и оптимизация на процеси*. Технически университет – София, 2000 URL <http://anp.vmei.acad.bg/Djiev/MOP.htm>
- [16] Димитров, Б, Н. Янев. *Вероятности и статистика*. СОФТЕХ, София, 2000.
- [17] Ефремов, А. *Идентификация на многомерни нестационарни процеси при непълна информация*. Дисертация. 2008.
- [18] Ефремов, А., А. Атанасов, Ф. Томова и Д. Ефремова. Групова стъпкова регресия за моделиране на пазарни системи. – В: *Известия на съюза на учените – Сливен*, том 20, 2012.
- [19] Ефремов, А., А. Атанасов, Ф. Томова и Д. Ефремова. Комбинирана групова стъпкова регресия и приложението ѝ в пазарния сектор. – В: *Автоматика и информатика*, том 1, 2012.
- [20] Иванов, Р. *Цифрова обработка на едномерни сигнали*. Габрово: „Принт“ ЕООД, София, 1999.
- [21] Ищев, К. *Теория на автоматичното управление*. ИК Кинг, София, 2000.
- [22] Калинов, К. *Практическа статистика за археолози и антрополози*. Нов български университет, София, 2002.
- [23] Маджаров, Н. Рекурсивна оценка на параметри на дискретни линейни динамически модели. – В: *Автоматика и изчислителна техника*, том 1, стр. 10 – 18, 1976.
- [24] Маджаров, Н. *Стохастични процеси в системите за управление*. Технически университет – София, 1993.
- [25] Пашева, В. *Въведение в числените методи*. Технически университет – София, 2009.
- [26] Перев, К. *Анализ на линейни системи. Ръководство за лабораторни упражнения*. Технически университет – София, 2004.
- [27] Петков, П. *Многомерни системи за управление*. Технически университет – София, 1997.
- [28] Петков, П. и М. Константинов. *Робастни системи за управление. Анализ и синтез с MATLAB*. АВС Техника, София, първо издание, 2002.
- [29] Петков, Т. *Идентификация на обектите за автоматизация*. Техника, София, 1972.
- [30] Петков, Т. *Идентификация на обектите на управлението*. Техника, София, 1984.

- [31] Растригин, Л. А., Н. Е. Маджаров и С. И. Марков. *Оценяване на параметри и състояния на динамически обекти*. Техника, София, 1978.
- [32] Славова, С. К. Компютърно моделиране на фрактали. Ръководство за дистанционно обучение, 2010.
- [33] Соболев, И. М. *Численные методы Монте-Карло*. Наука, Москва, 1973.
- [34] Стоянов, С. *Оптимизация на технологични процеси*. Техника, София, 1993.
- [35] Тихонов, А. Н. и В. Я. Арсенин. *Методы решения некорректных задач*. Наука, 1974.
- [36] Форсайт, Д., М. Малкълм и К. Молър. *Компютърни методи за математически пресмятания*. Наука и изкуство, София, второ издание, 1986.
- [37] Цыпкин, Я. З. Оптимальная идентификация динамических объектов. *Измерения, контроль, автоматизация*, том 3, № 47, стр. 47 – 60, 1983.
- [38] Цыпкин, Я. З. *Основы информационной теории идентификации*. Наука, Москва, 1984.
- [39] Эйкхофф, П., А. Ванечек, Е. Савараги, Т. Созда, Т. Накамизо, Х. Акаике, Н. Райбман и В. Петерка. *Современные методы идентификации систем*. Мир, Москва, 1983. Под ред. П. Эйкхоффа; пер. с англ. под ред. Я. З. Цыпкина.
- [40] Adam-Medina, M., D. Theilliol and D. Sauter. Simultaneous fault diagnosis and robust model selection in multiple linear models framework. – In: *5th IFAC Symposium on Fault Detection, Supervision and Safety of Technical Processes*, pp. 513 – 518. Washington, D.C., USA, June 2003.
- [41] Agresti, A. *Categorical Data Analysis*. John Wiley & Sons, Inc., Hoboken, New Jersey, July 22, 2002.
- [42] Ahmadi-Kashani, K. and A. J. Bell. The analysis of cables subject to uniformly distributed loads. *Engineering Structures*, volume 10, № 3, pp. 174 – 184, July 1988.
- [43] Ali, A. B. M. S. and Y. Xiang. *Dynamic and Advanced Data Mining for Progressing Technological Development: Innovations and Systemic Approaches*. Information Science Reference, 701 E. Chocolate Avenue, Hershey PA 17033, USA, 2010.

- [44] Allen, D. The relationship between variable selection and data augmentation and a method for prediction. – *In: Technometrics*, volume 16, № 1, pp. 125 – 127, 1974.
- [45] Altman, M., J. Gill and M. P. McDonald. *Numerical Issues in Statistical Computing for the Social Scientist*. Wiley-Interscience, first edition, 2003.
- [46] Anderson, R. *The Credit Scoring Toolkit. Theory and Practice for Retail Credit Risk Management and Decision Automation*. Oxford University Press, Inc., New York, 2007.
- [47] Bader, B. W. and T. G. Kolda. Algorithm 862: MATLAB tensor classes for fast algorithm prototyping. Sandia National Laboratories, volume 32, № 4, pp. 635 – 653, December 2006.
- [48] Baltagi, B. H. and D. Levin. Cigarette taxation: raising revenues and reducing consumption. – *In: Structural Change and Economic Dynamics*, volume 3, № 2, pp. 321 – 335, 1992.
- [49] Bar-Shalom, Y. and W. D. Blair. *Multitarget-Multisensor Tracking: Applications and Advances*. Artech House, Inc., London, third edition, 2000.
- [50] Bernard, J. T., N. Idoudi, L. Khalaf and C. Yeloud. Finite sample multivariate structural change tests with application to energy demand models. – *In: Journal of Econometrics*, 2007.
- [51] Bierman, G. J. Measurement updating using the U-D factorization. – *In: Automatica*, volume 12, № 4, pp. 375 – 382, July 1976.
- [52] Bjorkberg, J. and G. Kristensson. Three-dimensional subterranean target identification by use of optimization techniques. – *In: Electromagnetics Research*, volume 15, pp. 141 – 164, 1997.
- [53] Blackman, S. S. *Multiple Target Tracking with Radar Applications*. Artech House, Inc., London, 1986.
- [54] Blackman, S. S. and R. Popoly. *Design and Analysis of Modern Tracking Systems*. Artech House, Inc., Boston-London, 1999.
- [55] Blattberg, R. C. and S. A. Neslin. *Sales Promotion Models*, volume 5, chapter 12, pp. 553 – 609. Eds. J. Eliasberg & G. Lilien, Elsevier Science Publishers B. V., Amsterdam, The Netherlands, second edition, 1993.
- [56] Boyd, S. and L. Vandenberghe. *Convex Optimization*. Cambridge University Press, UK, 2004.
- [57] Bozdogan, H. Akaike's information criterion and recent developments in information complexity. – *In: Journal of Mathematical Psychology*,

- volume 44, № 44, pp. 62–91, 2000. URL <http://www.sciencedirect.com>.
- [58] Brandimarte, P. *Numerical Methods in Finance. A MATLAB-Based Introduction*. John Wiley & Sons, Inc., New York, 2002.
- [59] Brohus, H., C. Frier and P. Heiselberg. Stochastic load models based on weather data. Technical Report DK-9000, Department of Building Technology and Structural Engineering, Aalborg University, Aalborg, Denmark, June 2002.
- [60] Bunday, B. D. *Basic Optimization Methods*. Hodder Arnold, London, U.K.: E. J. Arnold, 1984.
- [61] Burnik, U., J. Tasic1 and G. Cain. Two-dimensional least square SVD algorithm. – In: *International Workshop on Image Processing: Theory, Methodology, Systems and Applications*. Budapest, Hungary, 20 – 22 June 1994.
- [62] Casella, G., S. Fienberg and I. Olkin. *Applied Regression Analysis – A Research Tool*. Springer-Verlag, New York, second edition, 1998.
- [63] Castaneda, C. *The Active Side of Infinity*. Harper Perennial, Harper Collins Publishers 10 East 53rd Street, New York, NY 10022, December, 22, 1999.
- [64] Chattefuee, S. and A. S. Hadi. *Regression Analysis by Example*. John Wiley & Sons, Inc., Hoboken, New Jersey, fourth edition, 2006.
- [65] David, B. and G. Bastin. An estimator of the inverse covariance matrix and its application to ML parameter estimation in dynamical systems. – In: *Automatica*, volume 37, № 1, pp. 99 – 106, 2001.
- [66] Dayal, B. S. and J. F. MacGregor. Multi-output process identification. – In: *Journal of Process Control*, volume 7, № 4, pp. 269 – 282, 1997.
- [67] Den Hof, P. M. J. V. Model sets and parametrizations for identification of multivariable equation error models. – In: *Automatica*, volume 30, № 3, pp. 433 – 446, 1994.
- [68] Deschamps, P. J. Full maximum likelihood estimation of dynamic demand models. – In: *Journal of Econometrics*, volume 82, pp. 335 – 359, 1997.
- [69] Deschamps, P. J. Exact small-sample inference in stationary, fully regular, dynamic demand models. – In: *Journal of Econometrics*, volume 97, pp. 51 – 91, 2000.
- [70] Dominick's database, provided by the James M. Kilts Center, GSB, University of Chicago. URL <http://research.chicagogsb.edu/marketing/databases/dominicks/>.

- [71] Douma, S. G., X. Bombois and P. M. V. den Hof. Validity of the standard cross-correlation test for model structure validation. – *In: Proceedings of the 16th World IFAC Congress*, pp. 1 – 6. Prague, Czech Republic, 4 – 8 July 2005.
- [72] Duda, R. O., P. E. Hatr and D. G. Stork. *Pattern Classification*. John Wiley & Sons, Inc., New York, second edition, 2001.
- [73] Efremov, A. An automatic identification cycle for dynamic demand models generation. – *In: International Conference on Computer Systems and Technologies, CompSysTech*, pp. V.5-1 – V.5-6. Varna, Bulgaria, 16 – 17 June, 2005.
- [74] Efremov, A. General Forms of a Class of Multivariable Regression Models. – *In: Journal of Information Technologies and Control*, Accepted for publication. Sofia, Bulgaria, 2014.
- [75] Efremov, A. Generalized representations multivariable linear parameterized models. – *In: International Conference of Automatics and Informatics*, pp. I-233 – I-236. Sofia, Bulgaria, 3 – 7 October, 2013.
- [76] Efremov, A. Least squares and weighted least squares for multivariable models represented in a parameter matrix form. – *In: International Conference of Automatics and Informatics*, pp. I-237 – I-240. Sofia, Bulgaria, 3 – 7 October, 2013.
- [77] Efremov, A. Linear Approach for Parameters Estimation of Multivariable Models in Parameter Matrix Form. – *In: Journal of Information Technologies and Control*, Accepted for publication. Sofia, Bulgaria, 2014.
- [78] Efremov, A. Real-time estimation of multivariate dynamic time-varying market representations. – *In: International Conference on Computer Systems and Technologies, CompSysTech*, pp. VI.20-1 – VI.20-7. Rousse, Bulgaria, 14 – 15 June, 2007.
- [79] Efremov, A. Recursive estimation of dynamic time-varying demand models. – *In: International Conference on Computer Systems and Technologies CompSysTech*, pp. V.7-1 – V.7-6. Veliko Tarnovo, Bulgaria, 15 – 16 June, 2006.
- [80] Efremov, A. System identification based on stepwise regression for dynamic market representation. – *In: International Conference on Data Mining and Knowledge Engineering*, volume 64, 2, pp. 132 – 137. Rome, Italy, 28 – 30 April, 2010.
- [81] Efremov, A. and A. Todorov. Automatic multivariate system identification dealing with problematic datasets. – *In: Cybernetics and Information Technologies*, volume 7, № 2, pp. 3 – 21, 2007.

- [82] Efremov, A. and A. Todorov. Modifications of estimation procedures and validation criterion, applied for market systems identification. – *In: Cybernetics and Information Technologies*, volume 7, № 3, pp. 3 – 12, 2007.
- [83] Efremov, A., D. Mihajlova, F. Tomova. Development of a system for demand analysis and sales forecast. – *In: International Conference of Automatics and Informatics*, pp. III-33 – III-36. Sofia, Bulgaria, 29 September – 4 October, 2009.
- [84] Efremov, A., I. Dimitrov and T. Djamiykov. Adaptive Device for Data Evaluation from Optical Reflection Sensors. – *In: The Tenth International Scientific and Applied Science Conference, Electronics*, pp. 219 – 224. Sozopol, Bulgaria, 26 – 28 September 2001.
- [85] Elden, L. *Numerical linear algebra and applications in data mining*. Preliminary version. Lecture Notes, Department of Mathematics, Linköping University, Sweden, 20 December 2005.
- [86] Eriksson, J. and P.-A. Wedin. Truncated Gauss-Newton algorithms for ill-conditioned nonlinear least squares problems. – *In: Optimization Methods and Software*, volume 19, № 6, pp. 721 – 737, December 2004.
- [87] Experian Inc. URL <http://www.experian.com/>.
- [88] Eykhoff, P. *System Identification: Parameter and State Estimation*. Wiley, New York, 1974.
- [89] Fang, Q. Distinctions between Levenberg-Marquardt method and Tikhonov regularization. June 30, 2004.
- [90] Faraway, J. J. Practical regression and ANOVA using R, 2002.
- [91] Farber, R. *CUDA. Application Design and Development*. Elsevier – Morgan Kaufmann, 225 Wyman Street, Waltham, MA 02451, USA, 2011. NVIDIA Corporation.
- [92] Fassois, S. D. MIMO LMS – ARMAX identification of vibrating structures – part I: the method. – *In: Mechanical Systems and Signal Processing*, volume 15, № 4, pp. 723 – 735, 2001.
- [93] Foubert, I. *Modelling Isothermal Cocoa Butter Crystallization: Influence of Temperature and Chemical Composition*. Ph.D. thesis, Ghent University, Belgium, 2003.
- [94] Fox, J. *An R and S-PLUS Companion to Applied Regression. Appendix: Robust Regression*. SAGE Publications, Inc., first edition, June, 2002.
- [95] Friesen, J., S. Capalbo and M. Denny. Dynamic factor demand equations in U.S. and canadian agriculture. – *In: Agricultural Economics*, volume 6, pp. 251 – 266, 1992.

- [96] Gallus Ltd. URL <http://www.gallus-consulting.com/>.
- [97] Gaop, J. and J. Zhang. Clustered SVD strategies in latent semantic indexing. – In: *Information Processing and Management*, volume 41, pp. 1051 – 1063, September 2005.
- [98] Gill, J. and G. King. What to do when your hessian is not invertible – alternatives to model respecification in nonlinear estimation. – In: *Sociological Methods & Research*, volume 32, № 4, pp. 1 – 34, apr 2004.
- [99] Gonzalo, J. and O. Martinez. Threshold integrated moving average models (Does size matter? Maybe so). – In: *European Economic Association, Econometric Society, North American Winter Meetings 145 – EEA-ESEM*. Econometric Society, 2004.
- [100] Gonzalo, J. and O. Martinez. Large shocks vs. small shocks. (Or does size matter? May be so.). – In: *Journal of Econometrics*, volume 135, pp. 311 – 347, 2006.
- [101] Greene, W. H. *Econometric Analysis*. Prentice Hall, Upper Saddle River, New Jersey, 07458, 2003.
- [102] Gregorcic, G. *The singular value decomposition and the pseudoinverse*. Lecture Notes, Department of Electrical Engineering, University College Cork, Ireland, 28 November, 2001.
- [103] Grewal, M. S. and A. P. Andrews. *Kalman Filtering: Theory and Practice Using Matlab*. John Wiley & Sons, Inc., New York, second edition, 2001.
- [104] Hamilton, J. D. *State-Space Models, Part 10 – Theory and Methods for Dependent Processes*. Handbook of Econometrics. North-Holland, Amsterdam, 1994.
- [105] Hamister, J. W. and N. C. Suresh. The impact of pricing policy on sales variability in a supermarket retail context. *Int. J. Production Economics*, 2007.
- [106] Hannan, E. J. The estimation of the order of an ARMA process. – In: *The Annals of Statistics*, volume 8, pp. 1071 – 1081, 1981.
- [107] Hanssens, D. M. and L. J. Parsons. *Econometric and Time Series Market Response Models*, volume 5 of *Handbooks in Operations Research and Management Science*, chapter 9. Eds. J. Eliasberg & G. Lilien, Elsevier Science Publishers B.V., Amsterdam, The Netherlands, second edition, 1993.
- [108] Hanssens, D. M., L. J. Parsons and R. L. Schultz. *Market Response Models: Econometric and Time Series Analysis*. Kluwer Academic Publishers, Boston, second edition, 2001.

- [109] Hastie, T., R. Tibshirani and J. Friedman. *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. Springer Series in Statistics, 233 Spring Street, New York, NY 10013, USA, second edition, 2009.
- [110] Hoffmann, J. P. *Generalized Linear Models – an Applied Approach*. Pearson Education, Inc., Boston, 2004.
- [111] Huber, P. J. Robust estimation of a location parameter. – In: *Annals of Mathematical Statistics*, volume 35, pp. 73 – 101, 1964.
- [112] Isenman, A. J. *Modern Multivariate Statistical Techniques. Regression, Classification, and Manifold Learning*. Springer-Verlag, 2008.
- [113] Isermann, R. and M. Munchhof. *Identification of Dynamic Systems: An Introduction with Applications*. Advanced Textbooks in Control and Signal Processing. Springer, December 1, 2010.
- [114] Jorgensen, S. B. and J. H. Lee. Recent advances and challenges in process identification. – In: *AIChE Symposium Series, ISU 326*, pp. 55 – 74. Budapest, Hungary, 2002. Invited paper in Chemical Process Control – 6.
- [115] Kalman, R. E. A new approach to linear filtering and prediction problems. – In: *Transactions of the ASME–Journal of Basic Engineering*, volume 82, Series D, pp. 35 – 45, 1960.
- [116] Klapper, D. and H. Herwartz. Forecasting market share using predicted values of competitive behavior: further empirical results. – In: *International Journal of Forecasting*, volume 16, pp. 399 – 421, July – September 2000.
- [117] Klapper, D. and H. Herwartz. Relevance of functional flexibility for heterogeneous sales response models: A comparison of parametric and semi-nonparametric models. – In: *European Journal of Operational Research*, volume 174, pp. 1009 – 1020, October 16, 2006.
- [118] Kulhavy, R. Restricted exponential forgetting in real-time identification. – In: *Automatica*, volume 23, pp. 589 – 600, 1987.
- [119] Ljung, L. *System Identification: Theory for the User*. Prentice Hall, Englewood Cliffs, 1987.
- [120] Ljung, L. *System Identification Toolbox. For Use with Matlab*. User's Guide. The MathWorks, Inc, MA, USA, version 6, 2004.
- [121] Makridakis, S., S. Wheelwright and H. R. *Forecasting – Methods and Applications*. John Wiley & Sons, Inc., New York, 1998.
- [122] Maleshkov, S., S. Markov, I. Kalaikov and N. Madjarov. Lspst subroutine package for identification and parameter estimation of linear dynamical models. pp. 531 – 538, 1980.

- [123] McFadden, D. L. *Economics and statistics*. chapter 4: Instrumental variables. URL http://elsa.berkeley.edu/users/mcfadden/e240b_f01/e240b.html. Lecture notes, University of California, Berkeley, 1999.
- [124] Montgomery, D. C., E. A. Peck and G. G. Vining. *Introduction to Linear Regression Analysis*. Willey Series in Probability and Statistics, Hoboken, New Jersey, forth edition, 2006.
- [125] Moor, B. D. and P. V. Overschee. *Trends in Control: a European Perspective*. Springer-Verlag London Limited, 1995.
- [126] Nelles, O. *Nonlinear System Identification. From Classical Approaches to Neural Networks and Fuzzy Models*. Springer-Verlag, Berlin Heidelberg, 2001.
- [127] Nocedal, J. and S. J. Wright. *Numerical Optimization*. Springer Series in Operations Research and Financial Engineering. Springer Science + Business Media, LLC, USA, 2006.
- [128] Oppenheim, A. V. and A. S. Willsky with S. Hamid Nawab. *Signals and Systems*. Second Edition, Prentice Hall, Inc., 1997.
- [129] Paoletti, S. *Identification of Piecewise Affine Models*. Ph.D. thesis, Italy, 2004.
- [130] Pereyra, V. and G. Scherer. Least squares scattered data fitting by truncated SVDs. – *In: Applied Numerical Mathematics*, volume 40, pp. 73 – 86, 2002.
- [131] Press, W. H., S. A. Teukolsky, W. T. Vetterling and B. P. Flannery. *Numerical Recipes in C – The Art of Scientific Computing*. Cambridge University Press, 40 West 20th Street, New York, NY 10011-4211, USA, 1992.
- [132] Raorane, A. A., R. V. Kulkarni and B. D. Jitkar. Association rule – extracting knowledge using market basket analysis. – *In: Research Journal of Recent Sciences*, volume 1, № 2, pp. 19 – 27, February, 2012. International Science Congress Association.
- [133] Rencher, A. C. *Methods of Multivariate Analysis*. John Wiley & Sons, Inc., Hoboken, New Jersey, second edition, 2002.
- [134] Retail-Analytics Ltd. URL <http://retail-analytics.com/>.
- [135] Ridder, F., R. Pintelon, J. Schoukens and D. P. Gillikin. Modified AIC and MDL model selection criteria for short data records. – *In: IEEE Transactions on Instrumentation and Measurement*, volume 54, № 1, pp. 144 – 150, 2005.
- [136] Romisch, U., H. Jager and D. Vandev. Application of interactive regularized discriminant analysis to wine data. – *In: Biometric*

- Conference of the Region Austria-Switzerland of the International Biometric Society*. Graz, September 26 – 29, 2005.
- [137] Romisch, U., H. Jager and D. Vandev. Interactive regularized discriminant analysis for determining the origin of wines from five countries. – In: *DAGStat and Biometric Conf. of the German Region of the Intern. Biom. Society*. Bielefeld, March 27 – 30, 2007.
- [138] Shalizi, C. *Classification and regression trees*. Lecture Notes for the course 36-350: Data Mining, Center for the Neural Basis of Cognition, Carnegie Mellon University, 2011.
- [139] Soderstrom, T. and P. Stoica. *System Identification*. International Series in Systems and Control Engineering. Prentice Hall, March 1994.
- [140] Stone, M. Cross-validation choice and assessment of statistical predictions. – In: *Journal of the Royal Statistical Society*, volume B, № 36, pp. 111 – 147, 1974.
- [141] Sugiyama, M. and H. Ogawa. Theoretical and experimental evaluation of the subspace information criterion. – In: *Machine Learning*, volume 48, № 1-2-3, pp. 25 – 50, 2002. DBLP, URL <http://dblp.uni-trier.de>.
- [142] Theilliol, D., D. Sauter and J. Ponsart. A multiple model based approach for fault tolerant control in non-linear systems. – In: *5th IFAC Symposium on Fault Detection, Supervision and Safety of Technical Processes*, pp. 151 – 156. Washington, D.C., USA, June 2003.
- [143] Van den Bos, A. *Parameter Estimation for Scientists and Engineers*. John Wiley & Sons, Inc., Hoboken, New Jersey, 2007.
- [144] Van Heerde, H. J., P. S. H. Leeftang and D. R. Wittink. Semiparametric analysis to estimate the deal effect curve. Research Report 99B35, University of Groningen, Research Institute SOM (Systems, Organisations and Management), 1999.
- [145] Verdult, V. *Non-linear System Identification – A State Space Approach*. Ph.D. thesis, The Netherlands, 2002. URL <http://doc.utwente.nl/38684/1/t0000018.pdf>.
- [146] Verhaegen, M. and V. Verdult. *Filtering and system identification: An introduction*. Lecture Notes for the course sc4040 (et4094), Delft University of Technology, 2003.
- [147] Verhaegen, M. and V. Verdult. Filtering and system identification: An introduction to using matlab software, 2003. Software Manual for the course sc4040 (et4094), Delft University of Technology.
- [148] Verhaegen, M. and V. Verdult. *Filtering and System Identification. A least Squares Approach*. Cambridge University Press, The Edinburgh Building, Cambridge CB2 8RU, UK, 2007.

- [149] Westland, J. C. and E. W. K. See-To. The short-run price-performance dynamics of microcomputer technologies. – *In: Research Policy*, volume 36, pp. 591 – 604, June 2007.
- [150] Yang, C.-Y. Estimation of linear regression models with serially correlated errors. – *In: Journal of Data Science*, volume 10, pp. 723 – 755, 2012.
- [151] Yang, W. Y., W. Cao, T.-S. Chung and J. Morris. *Applied Numerical Methods Using Matlab*. John Wiley & Sons, Inc., USA, 2005.
- [152] Yiu, J. C.-M. and S. Wang. Multiple ARMAX modelling scheme for forecasting air conditioning system performance. – *In: Energy Conversion and Management*, volume 48, pp. 2276 – 2285, June 2007.
- [153] Yu, W. J., B. Gomm and Yu. A model order and time-delay selection method for MIMO non-linear systems and it's application to neural networking. – *In: International Journal of Information and system sciences*, volume 1, № 1, pp. 39 – 60, 2005.
- [154] Zhu, Y. Multivariable process identification for MPC: the asymptotic method and its applications. – *In: Journal of Process Control*, volume 8, № 2, pp. 101 – 115, April 1998.

А л е к с а н д ъ р Е ф р е м о в

Идентификация на многомерни системи

ЛИНЕЕН ПОДХОД ЗА ОЦЕНЯВАНЕ

Монография

Българска
Първо издание
Велико Търново, 2013

ISBN 978-954-9489-34-7

Рецензенти:
проф. д.т.н. Никола Маджаров
проф. д-р Емил Гаринов

Художник на корицата: Александър Ефремов
Редактор: Милена Йовчева
Фигури: Александър Ефремов, Ирина Скордева
Предпечат $\text{\LaTeX}$: Александър Ефремов, Андрей Скордев

Печат: Д А Р – Р Х
Излязъл от печат: 15.10.2013 г.
Формат: 60/84/16

Издателство Д А Р – Р Х

DAR. PH

Велико Търново,
ул. Цар Самуил, № 7
тел./факс: 062/63 99 14
e-mail: darhristova@abv.bg
<http://www.darvt.com>