

Nové úpravy simulační metody Latin Hypercube Sampling a možnosti využití

Miroslav Vořechovský¹

Abstract

A new technique is proposed to improve the performance of Latin hypercube sampling. To reduce error in correlations between variables (correlated or uncorrelated), it is proposed to perform Simulated Annealing method on the realizations for all variables, instead of using matrix manipulation, as is the current custom. Limitations in hypercube sampling will be discussed. To improve the statistics for each variable, it is proposed that realizations for each variable be acquired by finding the probabilistic means of equiprobable disjoint intervals in the variable's domain, instead of using the cumulative distribution function directly, as is currently done (Huntington & Lyrantzis, 1998).

Úvod

Simulace typu Monte Carlo je mocný nástroj, velmi často používaný k analýze náhodných jevů pomocí počítačů. U náhodných problémů je požadována statistická či pravděpodobnostní informace o náhodném výstupu (odezvě) závisící na náhodných vstupních veličinách, polích a procesech.

Náhodné vstupní veličiny jsou u simulace Monte Carlo reprezentovány sadou deterministických čísel, tzv. realizací či vzorků. Náhodný problém je pak převeden na několik deterministických úloh, které se řeší daleko snadněji. Vstupní realizace jsou použity pro získání vzorků odezvy, které mohou poskytnout požadovanou statistickou či pravděpodobnostní informaci o odezvě. Simulace Monte Carlo je robustní, snadno použitelná a obecně rychlejší než plně pravděpodobnostní přístupy a proto je často používána pro řešení náhodných problémů a pro ověřování jiných metod analýzy.

Metoda Latin Hypercube Sampling (dále jen LHS) představuje jeden z nejúčinnějších postupů k odhadu statistických parametrů odezvy (konstrukce) a pro tento účel byla také původně vyvinuta (Iman & Conover, 1980). Výhodou je možnost výrazného snížení počtu simulací oproti jednoduché metodě Monte Carlo při zachování přesnosti odhadů (snížení počtu z řádu tisíců na desítky až stovky).

Nabízená úprava metody LHS zachovává funkci rozdělení všech simulovaných veličin a dodržuje požadované korelace mezi veličinami. LHS sestavuje vysoce provázanou sdruženou hustotu pravděpodobnosti vektoru náhodných veličin, která umožňuje vysokou přesnost statistických parametrů odezvy při použití malého počtu simulací. Lze rozlišit dvě fáze simulace LHS. Nejdříve jsou vybrány vzorky reprezentující funkce hustoty pravděpodobnosti veličin. Poté se pracuje s pořadími vzorků tak, aby se dosáhlo požadované statistické závislosti mezi veličinami. Vzhledem k tomu, že statistická závislost je vzorkům vnucena pouze záměnami pořadí a nikoliv změnou jejich hodnot, rozdělení pravděpodobnosti veličin zůstane nezměněno.

1 Metoda Latin Hypercube Sampling

V následujícím uvedme hlavní myšlenku metody. Odhady statistických parametrů odezvy jsou získány z určitého počtu simulovaných realizací funkce podobně jako v jednoduché metodě Monte Carlo. Definiční obor funkce hustoty $f(X_i)$ každé základní náhodné veličiny X_i je totiž rozdělen na N_{Sim} disjunktních intervalů. Volí se jednoduše intervaly o stejné pravděpodobnosti $1/N_{Sim}$. Současná praxe je získání reprezentativních parametrů zpravidla jako střed daného intervalu $1/N_{Sim}$ na distribuční funkci a uvažuje jej jako výchozí pro získání reprezentativní hodnoty veličiny pomocí inverzní transformace distribuční funkce. Jedná se tedy o stratifikační metodu (stratified sampling), kde vrstva oboru hodnot kumulativní distribuční funkce náhodné veličiny je nahrazena jedinou hodnotou. Je zajištěno, že celý rozsah každé

¹Ing. Miroslav Vořechovský, Ústav stavební mechaniky, Vysoké učení technické, Fakulta stavební, Brno, tel: +420 5 41147131, e-mail: vorechovsky.m@fce.vutbr.cz

základní náhodné veličiny je realizován při simulaci rovnoměrně vzhledem k distribuční funkci. Současně dosáhneme toho, že žádná reálná hodnota není předem vyloučena, každá vrstva distribuční funkce se při simulaci použije právě jednou. Tyto skutečnosti pak vedou k velmi dobrým odhadům při poměrně nízkém počtu simulací.

Výsledkem takových simulací jsou hodnoty odezvy Z_i ($i = 1, 2, \dots, N_{Sim}$). Odhady statistických parametrů náhodné proměnné Z jsou pak získány z tohoto náhodného výběru. Výsledkem strategie LHS jsou tedy parametry rezervy spolehlivosti - střední hodnota, směrodatná odchylka, koeficienty šikmosti a špičatosti, příp. empirická kumulativní distribuční funkce a citlivostní analýzy.

2 Výběr vzorků

První fáze metody LHS spočívá ve výběru vzorků reprezentujících intervaly rozdělení pravděpodobnosti náhodné veličiny. Definiční obor každé veličiny je rozdělen na N_{Sim} disjunktních intervalů o stejné pravděpodobnosti a z každého intervalu je vybrán jeden reprezentant. Stávající praxe je vybírat vzorky přímo inverzní transformací pravděpodobnostní funkce, a to v polovině vrstvy rovnoměrně rozděleného oboru hodnot kumulativní distribuční funkce:

$$x_{i,k} = F_i^{-1} \left(\frac{k - 0.5}{N_{Sim}} \right) \quad (1)$$

kde $x_{i,k}$ je k -tý vzorek i -té náhodné veličiny X_i , F_i^{-1} je inverzní distribuční funkce veličiny této veličiny a N_{Sim} je počet simulací, tedy počet vzorků pro každou veličinu. Proti popsanému postupu metody LHS lze uvést několik výhrad. Kritizovat lze zejména redukci výběru náhodných veličin z dané vrstvy na střed intervalu. Námitka se týká zvláště krajních vrstev, které nejvíce ovlivňují šikmost a špičatost rozdělení výstupních veličin. Uvedený postup u symetrických rozdělení poskytuje přesnou střední hodnotu získaného souboru vzorků (odhadovanou aritmetickým průměrem) a obecně je střední hodnota blízká požadované. Ovšem rozptyl vzorků je zpravidla značně rozdílný od požadovaného. Lépe je provádět výběr vzorků jako střední hodnotu každého intervalu (Huntington & Lyrintzis, 1998):

$$x_{i,k} = \frac{\int_{y_{i,k-1}}^{y_{i,k}} x f_i(x) dx}{\int_{y_{i,k-1}}^{y_{i,k}} f_i(x) dx} = \frac{\int_{y_{i,k-1}}^{y_{i,k}} x f_i(x) dx}{\frac{1}{N_{Sim}}} = N_{Sim} \int_{y_{i,k-1}}^{y_{i,k}} x f_i(x) dx \quad (2)$$

kde f_i je funkce hustoty pravděpodobnosti veličiny X_i a meze integrálu lze stanovit jako

$$y_{i,k} = F_i^{-1} \left(\frac{k}{N_{Sim}} \right) \quad (3)$$

Vzorky pak lépe reprezentují funkce hustoty pravděpodobnosti veličiny. Střední hodnota je získána přesně (aditivní vlastnost integrálu) a rozptyl dat je mnohem bližší požadovanému.

Pro některá rozdělení pravděpodobnosti (včetně normálního) je integrál (2) snadno řešitelný analyticky. V případě, že nalezení primitivní funkce není triviální, je nutno použít dodatečný numerický výpočet integrálu, ovšem takový nárůst výpočtové náročnosti se jistě vyplatí.

Vzorky vybrané oběma popsanými způsoby jsou téměř identické až na „konce“ oboru hodnot veličiny. Proto lze při simulaci veličiny s konečným oborem hodnot použít jednodušší postup (1). Pokud je obor hodnot rozdělení veličiny nekonečný či semi-nekonečný a numerická integrace přináší obtíže, stačí použít složitější výběr (2) pouze u několika prvních a posledních vzorků. I v takovém případě budou spočtené statistiky velmi blízko požadovaným.

3 Zavedení statistické závislosti, odhad korelačních koeficientů v metodě LHS

Poté co jsou vzorky všech veličin vygenerovány, je nutno pracovat se statistickou korelací mezi veličinami (nulovou či jinou). Doposud bylo předpokládáno, že vygenerované vektory matice $\mathbf{X}$ jsou nezávislé. Tomu by ovšem měla odpovídat i nezávislost statistických souborů tvořících jednotlivé realizace veličiny. Existuje možnost, že mezi sloupci matice vzorků $\mathbf{X}$ vznikne nezanedbatelná statistická závislost, která by mohla výrazně ovlivnit získané výsledky. O tom, zda matice $\mathbf{X}$ skutečně splňuje předpoklady statistické

nezávislosti se můžeme přesvědčit např. pomocí Spearmanova koeficientu pořadové korelace (4) nebo s pomocí klasického lineárního korelačního koeficientu. Je vhodné korelační koeficienty mezi veličinami měnit pouze změnami pořadí vzorků u jednotlivých veličin a neměnit již hodnoty vzorků. Nalezení výkonné metody záměn pořadí vzorků u každé veličiny umožní zavedení požadované statistické závislosti mezi nimi. Zde není nutno diskutovat okolnost, že nalezení sdružené rozdělovací funkce hustoty pravděpodobnosti pro všechny jednotlivé složky náhodného vektoru respektující požadované korelační koeficienty je v obecném případě nemožné. Pokud by to bylo možné (např. mnohorozměrné normální rozdělení), výběr reprezentantů by se realizoval přímo z takové sdružené funkce hustoty pravděpodobnosti.

3.1 ULHS a iterační metody

Uvažme případ, že všechny veličiny mají být generovány jako nekorelované. Případ korelovaných veličin bude rozebrán později. Doposud bylo možno vzorky generovat v náhodném pořadí uvažovat statistické závislosti jako dostatečně malé. Ovšem takto by pravděpodobně byla korelace mezi některými veličinami významná (Florian, 1992). Florian pro redukci náhodně zavedené korelace navrhuje metodu Updated Latin Hypercube Sampling (ULHS). Jeho postup je založen na technice, kterou lze nalézt v publikaci (Iman & Conover, 1982) pro zavedení požadované korelace mezi hypoteticky nekorelovanými veličinami. Myšlenka metody ULHS je rozebrána níže.

V metodě ULHS jsou pro třídící postup používána pouze pořadová čísla, nikoliv samotné vzorky. Pořadová čísla jsou hodnoty k použité ke generaci vzorků v rovnicích (1) nebo (2) a (3). Je vytvořena matice pořadových čísel $\mathbf{R}$ typu $N_{Sim} \cdot N_V$, kde každý sloupec z počtu N_V obsahuje permutace pořadových čísel vzorků ve sloupci. Poté, co je získána výsledná matice pořadových čísel, probíhá simulace se vzorky, které jsou umístěny do matice pro simulaci ve stejném pořadí jako v matici pořadových čísel. Pro výpočet i -tého výstupu se volí i -tý řádek jako sada vzorků v přetříděné matici.

Na počátku jsou permutace v matici pořadových čísel vygenerovány náhodně, ačkoliv je vhodné se ujistit, že žádná dvojice sloupců není podobná. Korelace mezi sloupci v matici pořadových čísel je definována Spearmanovým korelačním koeficientem:

$$s_{i,j} = 1 - \frac{6 \sum_k (R_{ki} - R_{kj})^2}{N_{Sim}^3 - N_{Sim}} \quad (4)$$

kde koeficienty $s_{i,j}$ reprezentují Spearmanovy koeficienty pořadové korelace mezi proměnnými i a j , $s_{i,j} \in (-1; 1)$. Korelační matice $\mathbf{S}$ je symetrická a pozitivně definitní kromě případů, kdy některé sloupce mají identická pořadí. Díky tomu lze provést Choleského dekompozici matice $\mathbf{S}$:

$$\mathbf{S} = \mathbf{Q}^T \mathbf{Q} \quad (5)$$

Nyní lze novou matici pořadí $\mathbf{R}_B$ generovat následovně:

$$\mathbf{R}_B = \mathbf{R} \mathbf{Q}^{-1} \quad (6)$$

Pořadová čísla v každém sloupci matice pořadí $\mathbf{R}$ se pak uspořádají tak, aby měla stejné pořadí, jako mají čísla v korespondujícím sloupci matice $\mathbf{R}_B$. Bylo prokázáno (Iman & Conover, 1982), že korelace mezi jakýmkoliv dvěma veličinami systému se nezvyšuje a může se snižovat. Poznamenejme, že tato technika může být provozována iterativně a teoreticky dovede značně snížit korelaci. Ovšem v praxi autoři zjišťují, že ULHS má tendenci konvergovat k pořadím, kde stále mezi některými veličinami zůstává významná chyba v korelaci. Proto je pro odstraňování nechtěné korelace vhodné použít jinou techniku.

Je-li požadováno simulovat korelované veličiny, pak má zmíněný postup jisté potíže. Nejdříve je nutno použít ULHS pro minimalizaci korelace. Pak je matice $\mathbf{K}$ s požadovanými korelačními koeficienty podrobena Choleského dekompozici (5) a s použitím rovnice (6) je získána nová matice pořadí. Tato procedura pro zavádění korelace může být použita pouze jednou, není zde možnost zlepšování iteracemi. Nová matice pořadí nedovoluje veličinám dobrou shodu s požadovanými korelačními koeficienty a pro zlepšení korelace nelze nic udělat. Je potřeba najít jinou metodu pro simulaci jak nekorelovaných, tak korelovaných veličin.

3.2 Spektrální rozklad korelační matice

V rámci hledání metody pro zavedení požadované korelace byla studována možnost provedení spektrálního rozkladu požadované korelační matice

$$\mathbf{K} = \Phi \Lambda \Phi^T \quad (7)$$

kde Φ je ortogonální transformační matice (matice vlastních vektorů $\mathbf{K}$) a diagonální matice Λ obsahuje (kladná, reálná) vlastní čísla $\mathbf{K}$ reprezentující korelační matici v nekorelovaném prostoru veličin. Je-li k dispozici matice vzorků $\mathbf{Y}$ s odstraněnou korelací mezi sloupci, matice simulací $\mathbf{X}$ by se pak mohla získat následovně:

$$\mathbf{X} = \Phi \mathbf{Y} \quad (8)$$

Uvedený postup by modifikoval hodnoty vzorků a navíc získané korelační koeficienty matice $\mathbf{X}$ se významně liší od požadovaných. Použije-li se získaná matice $\mathbf{X}$ jen k přetřídění původních vzorků $\mathbf{Y}$ ve stejném pořadí, korelační koeficienty se příliš nezlepší.

3.3 Provedení všech možností kombinací pořadí vzorků

Nejpřesnější možností zavedení požadované korelace je vyzkoušení všech kombinací pořadí vzorků pro jednotlivé veličiny. Ponecháme-li jeden sloupec vzorků nepermutovaný, je nutno vyzkoušet značné množství uskupení:

$$N = (N_{Sim}!)^{N_V - 1} \quad (9)$$

Lze snadno výpočtem ověřit, že pro reálné úlohy je dnes i v blízké budoucnosti uvedený postup nerealizovatelný. Slouží ovšem jako test jiných metod pro určení, zda bylo nalezeno nejlepší řešení (vzájemné pořadí) u malých úloh.

3.4 Optimalizační metoda simulovaného žíhání

Pro simulaci korelovaných i nekorelovaných metod byla v rámci výzkumu vypracována metoda s použitím optimalizačního postupu simulovaného žíhání. Předkládaná metoda může pracovat jak se Spearmanovým koeficientem, tak s klasickým lineárním či jiným koeficientem korelace. Uvedme variantu výpočtu s klasickým lineárním korelačním koeficientem. Nejprve je vyčíslena matice $\mathbf{U}$ normalizovaných vzorků pro všechny jednotlivé veličiny

$$\mathbf{X} \rightarrow \mathbf{U} : \quad u_{ij} = \frac{x_{i,j} - \mu_j}{\sigma_j} \quad (10)$$

kde μ_j a σ_j jsou střední hodnota a směrodatná odchylka j -té veličiny. Tato transformace zaručí, že všechny sloupce (veličiny) matice $\mathbf{U}$ mají nulovou střední hodnotu a jednotkový rozptyl. Pro každou dvojici (i, j) veličin je vypočten aktuální korelační koeficient

$$S_{i,k} = \frac{1}{N_{Sim}} \sum_{k=1}^{N_{Sim}} u_{ki} \cdot u_{kj} \quad (11)$$

Poté lze zvolit vhodnou míru či normu E pro hodnocení kvality celkových statistických závislostí v matici $\mathbf{U}$. Jako vhodný ukazatel se jeví např. maximální rozdíl v korelační matici požadované $\mathbf{K}$ a aktuální $\mathbf{S}$:

$$E = \left| \max_{1 \leq i, j \leq N_V} (S_{i,j} - K_{i,j}) \right| \quad (12)$$

nebo norma, která celkově zohlední deviace všech koeficientů s druhou mocninou

$$E = \sum_{i=1}^{N_V} \sum_{j=1}^{N_V} (S_{i,j} - K_{i,j})^2 \quad (13)$$

Zde lze samozřejmě doporučit výpočet rozdílů pouze pro horní nad-diagonální trojúhelníky symetrických korelačních matic.

Prvky normované matice $\mathbf{U}$ jsou nyní přetřídovány podle algoritmu simulovaného žíhání, a to vždy tak, že se provede pouze jedna záměna dvojice vzorků u jedné veličiny najednou. To umožní také značné úspory při přepočtu korelační matice. Stačí znát pozice vyměněných vzorků a korelační matici upravovat jen pro dotčené korelační koeficienty (v řádce a sloupci týkající se dotčené veličiny).

Na konci simulačního procesu se žádána matice realizací $\mathbf{X}$ pro výpočet získá operací inverzní k operaci (10):

$$x_{i,j} = u_{ij} \cdot \sigma_j + \mu_j \quad (14)$$

a mohou pokračovat výpočty odezvy s jednotlivými simulacemi (řádky) matice $\mathbf{X}$.

Zavádění požadované korelace je zde pojato jako optimalizační proces, ve kterém se minimalizuje jedna z norem (12), (13). Jedná se vlastně o hledání absolutního minima diskrétní funkce mnoha proměnných.

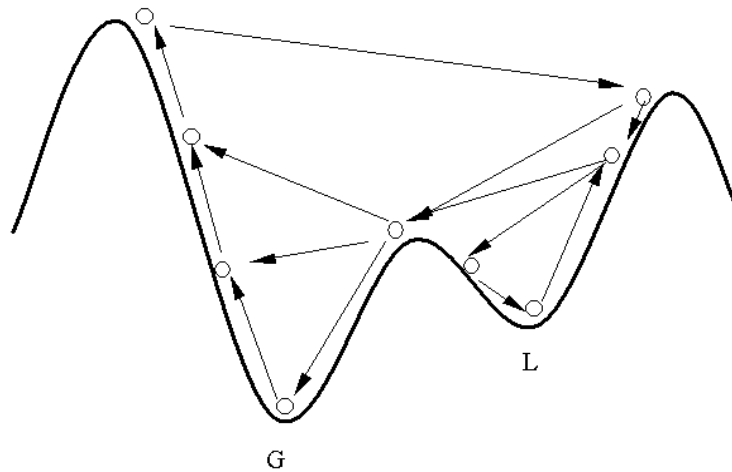
Při výzkumu metod byla zkušebně použita jednoduchá stochastická metoda, která náhodně vyměnila pořadí dvojice vzorků náhodně vybrané veličiny a v případě, že došlo ke snížení chyby, byla matice $\mathbf{U}$ přijata jako nová generace, tedy jako výchozí matice. Takový postup ovšem v drtivé většině případů skončí ve stavu, kdy již jakákoliv výměna dvojice vzorků nepřináší snížení normy E a tudíž nemůže být realizována a přijata. Lze říci, že metoda nachází určitá lokální minima. Poznamenejme, že za globální minimum lze považovat v tomto smyslu výsledek řešení podle odst. 3.3, realizovatelný pouze pro velmi malé systémy.

Je tedy nutné použít metodu, která pomocí náhodných operátorů a navíc s určitou pravděpodobností neuvízne v lokálním minimu. Osvědčila se aplikace metody simulovaného žíhání (Otten a Ginneken, 1989). Princip metody simulovaného žíhání je založen na porušení podmínky, že nová generace vzorků může být přijata pouze v případě, že došlo ke zmenšení normy (12), (13) nebo podobné.

V každé iteraci dojde k vytvoření generace potomků výměnou dvojice vzorků, která se přijme či nepřijme. Simulované žíhání se aplikuje ve druhém případě. Jestliže nedojde ke zmenšení funkční hodnoty minimalizované funkce (12) či (13), je nový vektor s určitou pravděpodobností přijat. Pravděpodobnost přijetí se řídí podle Boltzmannova rozdělení pravděpodobnosti:

$$P_r(E) \approx e^{\left(\frac{-\Delta E}{k_b \cdot T}\right)} \quad (15)$$

kde ΔE je rozdíl norem E před záměnou a po záměně. Toto rozdělení vyjadřuje představu, že systém v teplotní rovnováze o teplotě T má svou energii pravděpodobnostně rozloženu mezi všechny různé možné energetické stavy E . Dokonce i při nízkých teplotách existuje určitá pravděpodobnost (velmi malá), aby byl systém lokálně ve vyšším energetickém stavu. Proto existuje odpovídající možnost pro systém přemístit se z lokálního energetického minima ve prospěch nalezení lepšího, „více“ globálního minima. Hodnota k_b (Boltzmannova konstanta) je veličina, která uvádí do vztahu teplotu a energii systému. Metoda principiálně představuje analogii s tuhnutím krystalu. Principu simulovaného žíhání lze názorně porozumět na základě energetického grafu (obr. 1).



Obrázek 1: Ilustrace simulovaného žíhání na základě energetického grafu.

Je zde znázorněno lokální minimum (L), a globální minimum (G). Představme si, že zde máme umístěnu kuličku, která může zaujmout stabilní polohu v jednom z těchto minim. Jestliže určitý algoritmus vyhledá lokální minimum, v případě deterministických metod v něm kulička většinou skončí. U simulovaného žíhání existuje určitá pravděpodobnost (15), že kulička přeskočí do minima globálního. Systém

však musí být přiměřeně excitován (zahřát), aby energie, kterou kulička potřebuje k vyskočení z lokálního minima byla dostatečná. Analogicky si můžeme toto zahřátí představit jako zatřesení. Pokud bude zatřesení malé, pak kulička zůstane v lokálním minimu. Chceme-li dostat kuličku ven, musíme působit větší intenzitou. Pokud budeme aplikovat třes vysoké intenzity, pak bude kulička přeskakovat z minima lokálního do globálního a naopak. Je tedy zřejmé, že neoptimálnější postup je začít s vysokou intenzitou (teplotou) a postupně ji snižovat až téměř k nule a přitom si v průběhu „zapamatovat“ nejnižší polohu kuličky.

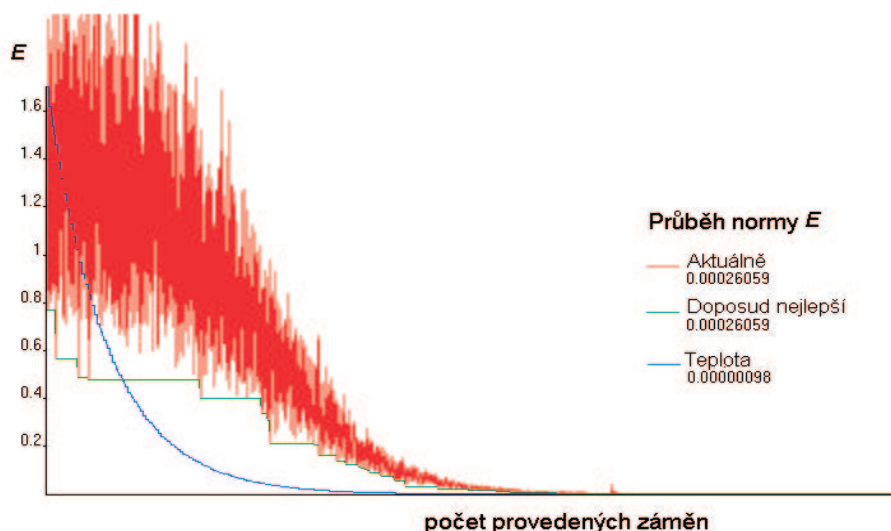
V našem případě lze hodnotu k_b nastavit jako jednotkovou (vypustit ze vztahu), protože se nabízí normu E a teplotu T uvažovat ve stejném rozměru. Zpravidla se doporučuje počáteční teplotu nastavit heuristicky tak, aby po provedení několika prvních iterací byl poměr úspěšných iterací ku všem iteracím přibližně 0,5. Tak je potřeba se zachovat v případě, kdy je metoda simulovaného žíhání použita pro hledání extrémů funkce, jejíž extrémní hodnoty nejsou vůbec známy a o vyšetřované funkci není k dispozici dostatečná tvarová představa.

Ovšem problém přibližování se k dané korelační matici skutečnou korelační maticí je značně ohraničen. Přípustné členy korelační matice leží vždy v intervalu $\langle -1; 1 \rangle$. Proto lze spočítat maximum normy (12) či (13) jako jistou mocninu sumy rozdílů příslušných pozic v požadované a hypoteticky nejodlišnější korelační matici. Takto lze také inteligentně odhadovat maximální (počáteční) teplotu systému. Vidíme již jasný vztah mezi normou E a teplotou T . Teplota systému je dále monotónně snižována redukčním faktorem. Teplotní redukční faktor T_f se obecně doporučuje nastavit na hodnotu 0,95. Lze ovšem používat varianty se sofistikovanějším systémem ochlazování (cooling schedule), např. vhodné adaptivní algoritmy. Snižování počáteční teploty T a každá další úprava této veličiny teplotním redukčním faktorem T_f se provádí heuristicky např. po konstantním počtu iterací:

$$T_{i+1} = T_i \cdot T_f \quad (16)$$

Jako koncovou teplotu lze po zkušenostech doporučit velmi nízkou hodnotu, řádově okolo $1 \cdot 10^{-8}$. Výsledky testování ukazují na fakt, že při těchto nízkých teplotách dochází v procesu ladění stále k významným zlepšením ve smyslu normy (12) či (13). Vyšší počet obrátek na jedné teplotní úrovni obecně zlepšuje dosažené výsledky. Je zde ovšem nevýhoda pochopitelného nárůstu výpočtového času.

Typický průběh snižování chyby, teploty a excitace systému v průběhu žíhání je na obr. 2. Klesající pravděpodobnost (15) se projevuje zúžením pásu okolo průměrné normy E (chyby).



Obrázek 2: Typický průběh normy E (chyby) při zavádění korelace metodou simulovaného žíhání.

Simulované žíhání zabraňuje v konečné fázi zpřesnění hodnot funkčního minima. Proto je výhodné po určitém počtu iterací vliv simulovaného žíhání potlačit. Po jeho ukončení přijme ta generace, při které byla funkční hodnota minimalizované funkce v celém dosavadním průběhu optimalizace nejmenší. Případný další průběh řešení představuje již pouze jakousi formu dolaďování a zpřesňování hodnoty minima metodou prostého náhodného výběru bez možností přijetí generace s větší chybou E .

Při zadávání požadované korelační matice se obecně může stát, že nepoučený uživatel bude požadovat korelační matici, která není pozitivně definitní. Je vhodné na tento fakt uživatele upozornit, ovšem simulační proces může běžet. Výsledná korelační matice bude pochopitelně splňovat požadavek pozitivní definitnosti. Stane se tak automaticky přibližně rovnoměrným nedodržením všech zadaných korelačních koeficientů. Jako jednoduchý příklad lze uvést úlohu, kdy bude požadováno zavést korelace mezi tři náhodné veličiny A, B a C podle matice $\mathbf{K}_1$ (sloupce a řádky odpovídají pořadí veličin A, B, C):

$$\mathbf{K}_1 = \begin{pmatrix} 1 & 0,9 & 0,9 \\ & 1 & -0,9 \\ sym. & & 1 \end{pmatrix} \longrightarrow_{optim.} \mathbf{K}_2 = \begin{pmatrix} 1 & 0,49 & 0,49 \\ & 1 & -0,49 \\ sym. & & 1 \end{pmatrix}$$

Uvedená korelační matice není pozitivně definitní. Je požadována vysoká statistická závislost mezi dvojicí (A, B) a mezi dvojicí (A, C). To povede k tomu, že vzorky těchto veličin budou přetříděny v téměř shodném pořadí podle velikosti. Ovšem pak je vynucená i vysoká korelace dvojice (B, C) a výsledkem by byly značně vysoká norma (13). Metoda si poradí tak, že přetřídí vzorky do stavu s korelační maticí přibližně $\mathbf{K}_2$. Taková matice je pozitivně definitní (těsně) a je určitým kompromisem pro zadanou korelační matici. Kdyby byla na počátku požadována korelační matice $\mathbf{K}_2$, byla by při dostatečném počtu simulací N_{Sim} téměř přesně splněna. Zde se také objevuje faktor počtu simulací N_{Sim} , čím větší N_{Sim} , tím je větší operativní prostor metody a pro velké počty simulací jsou dosahovány vynikající shody s požadovanou korelační maticí. Uvedený příklad jistě evokuje myšlenku vážené metody simulovaného žíhání, kdy se některé korelační koeficienty plní přesněji nežli jiné. V případě demonstrováných neznalostí by se např. vyzdvihla váha korelačních koeficientů o velikosti 0,9 s tím, že záporná korelace by nebyla dodržena, ovšem na hodnocení normou (13) by nepřesnost neměla velký vliv.

Závěr

Je předložena nová technika pro zvýšení výkonnosti metody Latin hypercube sampling. Předložena je metodika snížení rozdílů mezi požadovanou korelační maticí a korelační maticí spočtenou nad realizacemi vzorků metody LHS. Je doporučena optimalizační metoda simulovaného žíhání. Takový postup společně s přesnějším výběrem vzorků z libovolných rozdělovacích funkcí náhodných veličin (jako střední hodnoty v jednotlivých disjunktních intervalech o stejné pravděpodobnosti namísto přímé inverze kumulativní distribuční funkce) vede k výraznému zkvalitnění metody.

Jsou vysvětleny důvody, kvůli kterým bylo možné metodu simulovaného žíhání s úspěchem aplikovat. Obvykle je nutno metodu u nových a zcela obecných problémů heuristicky řídit (množství parametrů), ovšem jistá ohraničenost a jednostrannost zde řešené úlohy zaručila vysokou možnost samočinného řízení procesu. Patří sem zejména automatický odhad počáteční teploty. Výhodou metody je, že simulační proces lze kdykoliv ukončit, např. při dosažení uspokojivé míry korelace (normy). Lze použít váženou metodu zavádění korelace, kdy uživatel klade důraz na co nejpřesnější splnění vybraných korelačních koeficientů, jinak se odchylky realizují rovnoměrně. Jedná se o robustní, velmi výkonnou a rychlou metodu.

Metoda LHS je obecně použitelná v libovolném oboru (sociologie, lékařství, elektroinženýrství, atd.), kde je nutno sledovat charakteristiky náhodného výstupu v závislosti na popsáném náhodném vstupu. Návazně lze provádět také citlivostní analýzy.

V současné době je zkoumána možnost zlepšení metody LHS ke generaci náhodných procesů metodou ortogonální transformace korelační matice (Novák, Lawanwisut, & Bucher, 2000). Očekává se, že se tak výraznělepší statistiky generovaných polí včetně autokorelační funkce anebo cross-korelace při generování více korelovaných procesů.

Poděkování

Autor děkuje následujícím vědecko-výzkumným záměrům a grantům: projekt CEZ: J22/98:261100007, grant GAČR 103/00/0603. Dík patří také kolegovi Ing. Miroslavu Stiborovi, který autora při řešení úlohy provedl taji programování v jazyce C a pomohl mnohými nápady.

Reference

- (1) Florian, A.: An efficient sampling scheme: updated latin hypercube sampling, Probabilistic Engineering Mechanics, 7, 123-130, 1992
- (2) Huntington, D. E. & Lyrintzis, C. S.: Improvements to and limitations of Latin hypercube sampling, Probabilistic Engineering Mechanics, 13 (4): 245-253, 1998
- (3) Iman, R. L. & Conover, W. J.: Small sample sensitivity analysis techniques for computer models, with an application to risk assessment, Communications in Statistics: Theory and Methods, A9: 1749-1842, 1980
- (4) Iman, R. L. & Conover, W. J.: Approach to Inducing Rank Correlation Among Input Variables, Communications in Statistics, B11: 311-334, 1982
- (5) Novák, D., Lawanwisut, W. & Bucher C.: Simulation of random field based on orthogonal transformation of covariance matrix and Latin Hypercube Sampling, Proc. of Int Conf. On Monte Carlo Simulation Method MC2000, Monaco, Monte Carlo, 2000
- (6) Otten, R. H. J. M. & Ginneken, L. P. P. P.: The Annealing Algorithm, Kluwer Academic Publishers, USA, 1989