

3 次レベル IT サポートのためのシフトスケジューリング :問題点, モデル, 及び, ケーススタディ

Shift Scheduling for Third Level IT Support: Challenges, Models and Case Study

S. Wasserkrug¹, S. Tabu¹, S. Zeltyn¹, D. Gilat¹, V. Lipets¹, Z. Feldman¹

IBM Haifa Research Lab, Haifa University

A. Mandelbaum²

Industrial Engineering and Management Technion

Index Terms—Forecasting and Staffing, Shift Scheduling, 3rd level IT support, Forecasting, Workforce Management, Optimization

Abstract

企業の IT サポートをアウトソースする際には、通常レベルの IT サポートと同時に特別なスキルと知識を必要とする、一般に 3 次レベルといわれるサポートの提供が求められる。コールセンターを始め各種のサポート分野においては、サービス要求及び対応スタッフのシフトスケジューリングに関する多数の研究成果が既に存在するが、3 次レベルサポートに関しては、余りない。更に他のサービスレベルと 3 次レベルサービスでは、サービス要求が到着するパターン、個々のサービス提供に必要な時間分布等がかなり異なる。この結果、3 次レベルのサービスに対するスケジューリングとして現在行われているのは、現場での大雑把な判断に基づくやり方である。アウトソーシングビジネスの拡大に伴い、このレベルでのサポート業務においても実際のサポート需要に対応するシフトスケジューリングの理論及び実践が是非必要になっていると考え、予測から始まりスケジューリングまでを扱うメソドロジーを考えた。具体例を使ってその説明をすると同時に、このメソドロジーを使えばどの位マンパワーを節約し得たかも示す。

1. はじめに

自社の IT サポート業務を全てアウトソースする企業は益々増えており、それに対し IBM, Accenture といった大企業が、サービスビジネスとして IT サポートを提供している。

このサポートサービスには大きく、1 次レベル、2 次レベル、3 次レベルのサポートがある。1 次レベルサポートというのは、極めて基本的なサポート業務で、スキルは殆ど要求されない。例えば、ユーザのパスワードを設定し直す仕事である。2 次レベルサポートでは、

¹ セジエブ・ワサークラグ, xxx・タブ, セルゲイ・ゼルティン, ダガン・ジラット, ブラディミアー・リペッツ, ゾファー・フェルドマン

イスラエル IBM, ハイファ研究所, {SEGEVW, xxx, SERGEYZ, DAGANG, LIPETS, ZOHARF}@il.ibm.com

² アビシヤイ・マンデルbaum

テクニオン工科大学, イスラエル, avimtx.technion.ac.il

ある程度のスキルと知識が要求されるが、これが 3 次レベルとなると、かなりのスキルと特定のプラットフォームに関する専門知識が必要になる。このレベルでのサポート要員が行う作業例としてはサーバ OS のバージョンアップなどが挙げられる。

多くの場合、これらすべてのサポート業務は、世界各地のお客様にサポート業務を提供する巨大なサービスセンターで働く人々によって提供される。顧客が世界の各地に存在し時間帯も異なるため、そのサポートに休みはなく、サポートスタッフはシフトで働かなくてはならない。もう 1 つの特性は、サービスに対する需要が、曜日や 1 日の中の時間帯によって変動することである。このような背景から求められるサービスの質を維持しながらサービスの提供コストを最小にするシフトスケジューリングは簡単ではない。

必要とされるスキルレベルが相対的に低いので、1 次レベルサポートはコールセンターを介して提供されるのが普通である。コールセンターのシフトスケジューリングはいろいろな観点から広く研究されてきている（2 節を参照）。この結果、多様なスケジューリング方法とモデルが存在し、このレベルでの作業スケジューリング支援製品もいくつか販売されている。しかし 3 次レベルとなると、研究は殆どなされていない。しかもコールセンターが提供する 1 次レベルサポートと比べ、3 次レベルサポートには大きな相違点が存在する。例えば、単位時間あたりのサービス依頼件数が少ない、各々のサービス依頼を満足させ解決するには時間がかかる、なかには解決するのに複雑で大きなプロセスが必要になる場合もある、などである（詳細は 4 節を参照）。これらは、コールセンターサポート要員のスケジューリングモデルに関する方法論での前提と相容れない場合が多く、3 次レベルサポート用に既存のモデルを使うには、手直しをするか、新しいモデルを考えなければならない。

この論文では 3 次レベルのサポートのスケジューリングプロセスのために、いくつかのモデルと方法論を考えてみる（図 1 を参照）。更に、このプロセスを使って、実際にサポートを提供するチームのスケジューリングを行ったケースを紹介する。

2. この分野での研究

シフトスケジュールの自動生成で、一般的に使われているプロセスは、大きく分けると次の 3 段階からなる（図 1 を参照）。第 1 段階は負荷（サポート要求需要）予測で、1 週間もしくは 1 日内の単位時間帯で、サービス依頼がどのように飛び込んでくるかを予測する。次の段階は、配置要員数の算出で、1 週間もしくは 1 日内の各時間帯において、第 1 段階からの予測負荷量を満足させるために必要とされる人数及びスキルを割り出す。第 3 段階では各時間帯で必要とされる人数（スキルも考慮）を満足させ、さらに雇用者の勤務希望や労働基準をも満足する実行スケジュールを生成する。

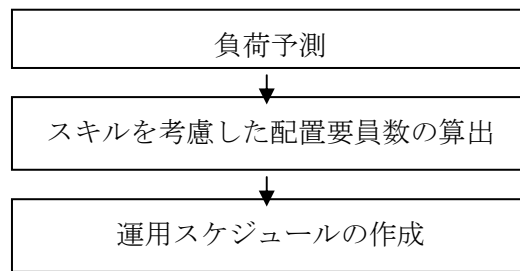


図1. スケジューリングプロセス

すでに1節で述べたように、3次レベルサポートスケジューリングに関しては殆ど何もなされていない状態である。そこで先ず、これに関連する2つの分野での、これまでの研究成果をサーベイしてみた。1つの分野は図1に示されている3つのステップの各々に対して適用し得る汎用的な手法であり、もう1つの分野はコールセンターでのシフトスケジューリングに使われる既存のモデルと手法である。

A 汎用手法

サービス要求負荷の予測とは、将来予期される負荷を生成するプロセスであり、その予測には単位時間毎の負荷事象の発生数、個々の負荷が必要とするスキルと所要時間といった情報が含まなくてはならない。運用スケジュールを作るには、通常2から6週間先の予測をなくてはならない（この予測対象期間は予測ホライズンと呼ばれる）。

負荷予測のために一般的に使われる手法には、大きく2つのタイプがある。1つは時系列分析（例えば、移動平均モデルを使ったやり方）、ARIMA [2]、そして回帰分析といった統計手法であり、もう1つは決定木、統計的ラーニングアルゴリズム等のモデルを使ったマシーンラーニング手法である [11]。どちらのモデルもデータドリブン、すなわち過去のデータに基づきサポート要求予測を生成する。

各時点における必要人員数を算出するために、この予測値からスキルを考慮した要員数を割り出さなくてはならないが、これには、通常待ち行列ネットワークモデルが利用される。M/M/Nタイプの待ち行列等に基づく分析モデルが使われる場合もあれば、より複雑な場合には待ち行列ネットワークのシミュレーションが利用される（例えば [2]）。

運用スケジュールを作るには通常ヒューリスティックアルゴリズム、混合整数計画法、あるいはこれらをミックスしたやり方が用いられる。

しかし、異なる種類のサポート作業の各々に対し、これらの方法を利用し、図1での各段階に適したモデルを見つけ、運用スケジュールまで作ってくれるスケジューリング方法を考えるのは容易なことではない。

コールセンター作業のような場合は、上記の各段階において標準的なモデルが複数存在し、商用製品でも広く使われている。このことは次の節でより詳しく説明するが、3次レベルのサポート作業となると、広く知られた適当なモデルがない。

B コールセンター用シフトスケジューリングモデル

図1の各段階で広く使われているモデルは次のようなものである：

●**需要(サービス要求)予測**：予測ホライズンは、通常 15 分から 30 分の予測インターバルと呼ばれる短い単位時間帯に分けられ、各インターバルでの需要は一樣ポアソン分布に従うものとする。ここでのモデルが予測するのは、予測ホライズンの各インターバルでのサービス要求平均到着率である。このインターバル毎の到着率は、移動平均、単純指数平滑化といった単純な時系列を使ったやり方で計算される [2]。

●**配置要員数の算出**：必要人数は、待ち行列分析モデルや定常分析により計算される。顧客がサービスを受けるまで待ち続けるなら M/M/N モデルであり、そうでなければ M/M/N+M モデルとなる [10]。これらのモデルを使い、要員配置インターバル毎の人数が計算されるが、この際、要員配置インターバルと予測インターバルとは同じ時間の長さであることを前提としており、要員配置インターバル毎に指定されたサービスレベル基準を満足するのに最低必要な人数が計算される。この基準というのは例えば到着するサービス要求の 80% は 20 秒以内に応えなくてはならないといった縛りである。このやり方ではこの基準が要員配置インターバル（すなわち予測インターバル）毎に、守られなくてはならないと想定している。

●**運用スケジュールの作成**：実際の運用スケジュールリングにはヒューリスティックもしくは混合整数計画法を使う（この例としては [3], [9] を見てほしい）。両者とも問題の定式化はコールセンターに当てはまる種類の作業を前提にしている。

これらの 3 つの段階で使われたモデルには、暗黙のものも含めて、いくつかの前提なり仮定がある。例えば、サービス要求は区分一樣ポアソン分布に従うとか、各要員配置インターバルでの配置要員数を算出するのに定常状態での分析が使えるとするなどである。

しかし、コールセンターの作業といえども益々複雑になっているので、負荷予測、要員配置と運用スケジュールリング、などにおける研究は現在でも進められている分野である。その例としては、異なるスキル群を前提に、サービス要求のルーティング（対応グループ間で手渡されて行くこと）、エージェント（サービス要員）のスケジュールリングといった研究である。コールセンター分野での諸研究及び問題点を見るには、[1], [5], [7] を参照してほしい。

3. 3 次レベルサポートの特性

この節では、その要員スケジュールリングにおいて考慮しなければならない下記の 3 次レベルサポートの特性について述べる：

- サービス要求の到着プロセスに関しては、到着特性の他にその処理に要する時間の特性。
- 守らなくてはならないサービスレベル基準。
- サービス要求を処理する際のワークフロー。
- 3 次レベルサポートでの作業に関して守らなくてはならないスケジュールリングル

ル。

次節でこれらを1つずつ検討していくが、これからは、ケーススタディに参画したサポートチームのデータと特性を使って説明する。このサポートチームは約25人のエージェントからなり、週7日、毎日24時間のサポートを提供している。

A 到着プロセス

3次レベルサポートチームが解決するサービス要求依頼は、問題チケット（略してチケット）と呼ばれる。この作業はその性格上次のような特性を持つ：

- 平均的に、タイムインターバルに到着するチケット数は小さい。
- 各チケット毎に必要とされる処理時間は長い。

ケーススタディでの問題チケットは1時間あたり1件か2件の到着率であり、チケット1件あたりの所要時間は1時間47分であった。更に1時間あたりの到着チケット数には大きなばらつきがあり、全くチケットが発生しない時間も多くあると思えば、10件以上のチケットが入ってくる時間も数時間ある。

B サービスレベル

各問題チケットには対応する問題の深刻度がついており、そのレベルにより異なるサービスレベルが要求される。ケーススタディでは、深刻度が5段階あり、深刻度1には最高度のサービスが要求される。深刻度5には、対応する具体的なサービスレベルがなく、最善を尽くすという漠然とした表現で対応される。サービスレベル1から4が要求する基準は次の通り：

- 深刻度1：チケットの90%は4時間以内に解決されなくてはならない。
- 深刻度2：85%は8時間以内に解決されなくてはならない。
- 深刻度3：80%は3日以内に解決されなくてはならない。
- 深刻度4：80%が5日以内に解決されなくてはならない。

更に、上記SLA（サービスレベルアグリーメント）は各月毎に守られなくてはならない。例えば、ある時間帯では深刻度1なのに4時間以内に解決されたのは90%に満たなくとも、月間でみたら90%に達していれば良い。

C 問題チケット処理のワークフロー

問題チケットが受け取られ開かれると、2次レベルサポートチームか、1次レベルサポートチームから3次レベルサポートチームに回される。その後、3次レベルサポートチーム内だけで解決されるケース以外に、以下のような他の組織グループに回される場合がある：

- サポートサービスを利用している顧客会社にコンタクトする。例えば問題に関する不足情報を顧客に提供して貰うとか、何らかのアクションを取って貰う必要があるといった場合である。
- チケットに対応した最初のチームと同じ会社内の別のチームにコンタクトする。例えば同じサービス会社が、特定のDB（例えば、Oracle^{®3}やIBM[®] DB2^{®4}）と、特定OS（例

³ Oracle は Oracle Corporation およびその関連企業の登録商標。

例えば、Windows⁵やLinux⁶) に対して 3 次レベルサポートを提供している場合に、このケースが発生する。サポート作業には特有の知識が必要になるので、3 次レベルサポートにおいても、これらの 2 つのチーム (DB チーム、OS チームと呼ぶことにする) が控えている。問題チケットが初めは OS の初期導入でのオプションに絡むと判断されて OS チームに回されたが、作業が始まったら DB マネジメントシステムのコンフィギュレーション問題が絡むと分かり、DB チームに渡されることになるというケースである。

●例えば OS での問題解決において、その OS を提供する会社からのパッチが必要だという場合には、他社からのサポートが必要になる。

上記のようにチケットの辿るルーティングは、いくらでも複雑になり得る。チケットの行った先でも上の例にあるような色々な行き先に回され得るからである。しかし、チケットが回されて行くところが顧客にせよ、他のサービス会社にせよ、そこで費やす時間は、この前節で定義されたサービスレベル基準での時間には勘定されない。この点を明白にするために具体例で説明すれば、深刻度 1 のケースでチケットをオープンしてからクローズするまでに 8 時間が経過したとする。しかし、そのうち 5 時間が顧客先での不足情報集めに使われた場合には 3 時間だけが勘定されるので、サービスレベル基準に反していない。

D スケジューリングルール

ケーススタディでのスケジューリングでは、いくつかのルールが守られなければならなかった。この論文では、それらを全部説明する必要はないが、そのうち 1 つ、オーバーラップルールを例にして我々のアプローチを説明する。

オーバーラップルールというのはチケットをオープンしてからクローズするまでに時間が数時間の場合もあれば数日の場合もあるという事実から生まれた。総時間が大きい理由としてはチケットの問題自身がかかなりの解決時間を必要とする、もともと対応サービスレベル基準がチケット解決に多くの時間をかけることを許している、はじめに問題チケット処理を任されたチーム以外の、どのグループにチケットを渡すかを決めるまでにかなりの時間がかかることがある、などである。このため、1 人のエージェント (エージェント A と呼ぶ) が始めたチケットが次のシフトでもう 1 人のエージェント (エージェント B と呼ぶ) に手渡される状況が起こる。この手渡しを可能にするには、シフト間でのオーバーラップが必要になってくる。この手渡しをきちんと行うためにエージェント B はエージェント A のシフトが終わる前に仕事を始める必要がある。

4. 既存モデルの不適合さ

⁴ IBM, DB2 は、International Business Machines Corporation の米国およびその他の国における商標です。

⁵ Windows は、米国 Microsoft Corporation の米国およびその他の国における登録商標です。

⁶ Linux は、Linus Torvalds の米国およびその他の国における登録商標または商標です。

3 次レベルサポートはその特性のために、サポート作業スケジューリングに適した方法もモデルも知られていない。まずコールセンターのための作業スケジューリングモデルを利用することが何故問題なのか、いくつかの理由を挙げてみる。

例えばコールセンター用のサービス負荷予測から配置要員数を割り出すのに、M/M/N や M/M/N+N 等の待ち行列を用い、スケジューリング用単位時間あたりの定常状態分析が行われる。しかし、これらのモデルの利用は次のことを前提としている。

- 各要員配置インターバルでのサービス要求依頼は、一様ポアソン分布に従っている。
- 各要員配置インターバルにおいて、システムは急速に定常状態に達する。
- 要員配置は、各要員配置インターバル内でサービスレベル基準を満足するようになされる。

問題チケットの到着パターンを分析するとき、一様ポアソン分布という仮定が納得できるためには、要員配置インターバルが 1 時間以内でなくてはならない。平均的に 1 時間の単位時間では到着するチケット数は非常に小さく、各チケットの処理には多くの時間がかかるので、1 時間以内にはとても定常状態に達しない。更に 3 次レベルでの制約条件は要員配置インターバルレベルではなく月単位で定義されているので、M/M/N 待ち行列の定常分析は適当とは思われない。

従来の要員配置モデルが不適当であるもう 1 つの理由は、飛び込んでくる用件は単一の作業で受けたエージェントが対応するという前提になっていることである。作業プロセスが複数のグループもしくはエージェントの手を経る可能性には対応しておらず、1 人のエージェントもしくは 1 グループだけしか対象にしていない。更に、コールセンターモデルは、サービス要求を処理する際の経過時間に、SLA での経過時間には勘定されないものが入ってくる場合は考慮していない。

最後に述べる不適当さとして、3 次レベルサポートにおいては、入ってくるチケットは解決されるまでオープン状態で存在し、途中でなくなることはない。従って、M/M/N+M は不適当である。更に途中で顧客が去ることがないので、要員不足は不安定なシステムを意味することになり、負荷（サービス要求）の予測を配置要員数に読み直す際には、十分な配慮が払われなくてはならない。

5. 3 次レベルサポート作業のスケジューリング

我々のメソドロジーでは、負荷の予測から配置要員数を計算するのにシミュレーションモデルを利用した。そうした理由は、たとえ最新の優れた待ち行列モデルを使っても、この作業に伴う複雑さを適切に表現できないからである。もう 1 つ留意すべき点は、シミュレーションモデルを使うことにより、負荷予測モデルの構築に伴う諸必要条件を緩めることができることである。分析モデルを使おうとしたら、到着プロセスを正確に反映し、かつ待ち行列理論による分析に耐えられるような負荷予測モデルを見つけてこなくてはならない。その点シミュレーションモデルを使えば、負荷予測もシミュレーションで済む。

次に3次レベルサポートを担当するエージェントをスケジューリングするのに必要な3つのコンポーネントについて説明する.

A. 要員配置シミュレーションモデル

3次レベルサポートの特性を考慮してシミュレーションモデルを作成した. 重要な特性の1つとして問題チケットは必ずいつかどこかで, 組織グループの1つに手渡されるということである. チケットがどのグループからどのグループに手渡されていくかの詳細情報は得られなかったので, スケジューリングプロセスの比較的簡単なモデルを作ってみた. これを図2に示す.

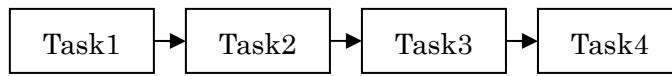


図 2. スケジューリングプロセス

図2で定義されたプロセスによれば, 複数個のタスクが連続したプロセスとして, チケットがモデル化される. ここで使われるタスクは次の通りである:

- *Task1*: そのチケットに割り当てられたチームが, 他の組織グループに引き渡すまでに済ませておくべき作業時間. 但しその作業が実際に始まるまでの待ち状態時間は含まれない.
- *Task2*: チケットが, 顧客 (サービス要求先) もしくは別のサービスプロバイダに渡され, そこで過ごす総時間. これには手渡されてから作業開始までの待ち時間も含まれる.
- *Task3*: 同じ会社内での異なるチームの下で過ごす総時間. 作業が実際に開始するまでの待ち時間も含まれる.
- *Task4*: チケットがタスク1で当初割り当てられたチームに再び戻された後, 残りの作業を終わらせるに必要な時間.

前にも述べたが, タスク2での経過時間は, サービスレベル基準をチケットが満すかどうか判断する際には勘定されない. これを正式に表現するため, 先ず, 深刻度 i のチケットが満足しなければならないサービスレベル基準時間を T_i で表せば, 深刻度1のチケットは $T_i = 4$ 時間となり, 式(1)は深刻度1のチケットがその基準を満たすための条件となる.

$$W_{Task1} + Task1 + Task3 + W_{Task4} + Task4 \leq T_i \quad (1)$$

ここでは,

- W_{Task1} と W_{Task4} は, そのチケットに割り当てられたチームがタスク1とタスク4それぞれの作業を始めるまでの待ち時間を表す.
- *Task1* と *Task4* は, タスク1もしくはタスク4を完了するのに要する時間.
- *Task3* は, 同じサービス会社内の異なるチームがこのチケットに割り当てられて作

業する時間であり、待ち時間も含んでいる。

全てのチケットが同様なルーティングを経るわけではないので、チケットによってはタスク 2 もしくはタスク 3 がゼロだったり、あるいは両方がゼロだったりする。

上記のような問題チケットのフローに基づきシミュレーションモデルを作成した。このシミュレーションモデルに、適当な最適化手法（例えば、Tabu/Scatter search）を組み合わせるなどすれば、深刻度毎に異なる条件を満たす必要な要員配置レベルが得られる。この最適化をすれば、サービスレベル基準を満たしながら、各シフトに配置される人数の合計を最小にすることが可能である。

B 負荷予測モデル

予測モデルには次の要素が含まれていなければならない：

- チケットの到着プロセス。
- 各チケットが必要とする作業時間（即ち $Task1 + Task4$ ）。
- 各チケットが顧客もしくは別のサービス会社に割り当てられ、そこで過ごす時間（ $Task2$ ）。
- 各チケットが同じサービスプロバイダ会社内の異なるチームに割り当てられ、そこで過ごす時間（ $Task3$ ）。

ケーススタディにおいては、深刻度が異なれば到着特性も異なるので、深刻度毎に別々の予測モデルを作成する必要があったが、大量の過去データを利用し予測モデルを構築した。

到着プロセスをモデル化するには、次のやり方が用いられた：

- 1日あたりの到着数を予測する。
- 個々の日について到着分布を見つける。

1日あたりの到着数を予測するモデルの基本形式を式（2）に示す。 X_t は t 番目の日に到着する依頼数、 M_t は $X_1, \dots, X_{t-1}$ により一意に決まる数、 Y_t は何らかの分布に従うランダム変数である。

$$X_t = M_t + Y_t \quad (2)$$

M_t を決めるのに、いくつかの時系列、例えば移動平均、単純指数平滑化や2段階指数平滑化、その他季節変動を考慮するやり方等を試してみたが、単純移動平均が最もよかった。 Y_t は $X_1 - M_1, X_2 - M_2, \dots$ に対して分布を当てはめるやり方で求めた。この結果、予測値と実際の値の間に高い相関（相関係数 0.85）が得られた。

1日あたりの到着数が分かったとして、1日の中での到着イベントがどうばらついて発生するかについては、何ら理論的な分布は得られなかった。そこで、実際データからの到着分布を利用したが、シミュレーションに基づく要員配置モデルと一緒に使う分には、これで十分である。

各チケットが必要とする作業時間に関しては、準指数分布であると分かった（[5]を参照）。準指数分布の1つの例は対数正規分布である。このようにして分布関数を選んだ後、そのパラメータを決めるために最尤推定量を求めた。必要な作業時間が準指数分布に従うとい

う事実は、M/M/N モデルの前提の 1 つ、サービス時間は指数分布に従う、というものに矛盾している。

最後に、タスク 2 とタスク 3 の分布に関しては、分布フィッティング手法を用いた。

C 運用スケジュールの作成

負荷予測とサービスレベル条件から配置要員数が計算されると、次に混合整数計画法を使って運用スケジュールが作られる。このスケジュールを作るためには、スケジューリングで考慮される諸規則を線形拘束条件で表現しなくてはならない。この詳細な説明はこの論文での本旨ではないので要点だけ説明しておく。

一般に、この問題は割当問題として定式化される。すなわちエージェント i がシフト j に対して割り当てられるなら 1、そうでなければ 0 の値を持つバイナリ変数 X_{ij} を定義し、スケジューリングで考慮しなければならないルールが満足されるように拘束条件を設定する。その一例としてオーバーラップ条件を満足させる拘束条件作りを以下に示す。

ここに 2 つのシフト s_1 と s_2 があり、 s_2 は s_1 の終了時より 1 時間早く始まるものとする。このオーバーラップ条件を満足するために、 s_1 にエージェントが割り当てられるには、 s_2 にもエージェントの誰かが割り当てられなくてはならないと仮定する。更に、チームは 10 人のエージェントからなるものとし X_{ij} をエージェント i がシフト j に割り当てられたら 1、

そうでなければ 0 の値を持つバイナリ変数として定義する。シフト 1 にエージェントを割り当てるには、シフト 2 にエージェントを必ず割り当てるということを満足させるために S_1 と S_2 という新しいバイナリ変数を導入し、式 (3) (4) (5) のような拘束条件を設定する。

$$\forall i, X_{i1} \leq S_1 \quad (3)$$

$$S_2 \leq \sum_{i=1}^{10} X_{i2} \quad (4)$$

$$S_1 \leq S_2 \quad (5)$$

これらが上記条件を保証することの説明をする。もしエージェントの誰かがシフト s_1 に割り当てられると、 $X_{i1}, i=1, \dots, 10$ の必ず 1 つが 1 の値をとるので、 S_1 は (3) の拘束から、1 の値をとる。この結果、式 (5) より S_2 の値は 1 でなくてはならないが、式 (4) から、 $X_{i2}, i=1, \dots, 10$ の 1 つは 1 になり、スケジュール条件が満たされる。

他のルールも、このような線形不等式を利用して同様に表現される。

6. スケジューリングメソッドロジーの実用化

分析的にスケジューリングを行うやり方の主要目的は、IT サポートにかかる人員コストを削減することにある。この目的で、これまで説明してきたメソッドロジーを 2 通りに使ってみた。

1. ストラテジック計画：長期的にサービスレベル基準を満たしながら、サポート提供す

るのに必要な最小チームサイズを決める。

2. オペレーショナル計画：与えられたチームサイズで、サービスレベル基準を守りながら最適運用スケジュールを作る。我々が考えたメソドロジーは、この目的にそのまま直ぐ使える。

ケーススタディでの対象チームに対し、その最小サイズを決めるため、まず需要データを考察し長期的な傾向があるかどうか分析したが、何ら増大傾向は見えなかった。このことは、将来も需要に大きな変化はないとして、これまでの需要に十分対応できる最小サイズは将来のサポートを提供するにも十分なサイズであるといえる。

過去のサポートを提供するのに必要なチームの最適サイズを計算するために、過去の各月に対して最適運用スケジュールを作ってみた。この際、休暇や訓練期間の要素も考慮した。このような実際の運用スケジュールに基づいて、必要とされる要員配置レベルも考慮しながらチームサイズを決めるほうが、要員配置モデルだけで決めるよりはるかに結果が信頼できる。なぜかといえば、休み時間、休暇、遅れなどの実務でのやりくりや、スケジュールリングルールを組み入れる必要があるので、サービスを提供するのに必要な人数を増やす必要が起きてくるかも知れないからである。

戦略的計画のために上記プロセスをケーススタディチームに適用したところ、サービスレベルを維持したままで、現状のサイズをかなり縮小できることが分かった。当たり前だが、そのような縮小案を実践するには十分注意してやらなければならない。

- チームから1人人間を外すには、人事上の問題が伴う。これは大事なことだが、ケーススタディにおいてはチームから外されても、それが職を失うことにはならなかった。外された人は、他のチームでそのスキルを必要とするところに移されただけだが、それでも人事上の複雑な問題は絡んでくる。
- 人をチームから外したりチームに加えるというのは、瞬時にして実行できるアクションではない。分析からの進言に従って人を外してみたら、分析結果に反して、サービスレベルに悪い影響が出たというので、その人を以前のチームに戻し全ては前のように、という具合にはいかない。

このような理由で、実際にチームサイズを縮小する前に先ず分析結果の検証を行う必要があった。この目的のため次のような検証を行った。

- 将来の月間スケジュールを作る際には、2つのスケジュールが作成された。1つ目は実際用スケジュールで、現状の人数（縮小案を実行する以前の）通り、2つ目が（最適スケジュールで）戦略的計画に基づき、人数が削られている。但しこの最適スケジュールを作る際には休暇や訓練といった要素が考慮されていた。
- チームのスケジュールリングは、実際用スケジュールに基づいて行われた。
- 最適スケジュールを作った対象期間における最後の月だけは、最適スケジュールで実際のサービス要求需要に対処してみた。この目的は最適スケジュールで果たしてサービスレベル基準を維持できるかどうかをチェックするためであった。

上記検証方法を数ヶ月にわたり実行してみたが、最適スケジュール（チームの人数が減らされている）を使ってもサービスレベルの基準は満たされていた。これで、（我々のメソドロジーが）チームサイズと人件費をかなり縮小できるポテンシャルを有していることが確認できた。

7. まとめと今後の研究

3次レベルサポートを提供する仕事は、そのためのシフトをスケジューリングするニーズが益々大きくなっていると同時に、その色々な特性があるため、それにあった新しいサポート要求の予測方法、要員配置及びスケジューリングのためのモデルと方法論が必要である。この論文ではそのための特別なモデルと手法を含め、そのようなメソドロジーを考えてみた。同時にケーススタディでは具体的なチームに適用し人員削減においてかなりの潜在的利益が得られることをみた。

このメソドロジーにおける現在の要員配置モデルはシミュレーションに基づいているが、今後の研究ではコールセンターや医療センターといった分野で利用される待ち行列理論での分析モデル（[8]を参照）を更に発展させ、我々の分野でのモデル化と分析もできるようにしたい。既に研究を始めているものには、問題チケットが色々な組織グループにルーティングされていくことや、サービス時間が本当の準指数分布に従わないなどの問題点への対応も入っている。また、予測モデルもそのような解析的な手法が可能になるように改良しているところである。例えば最近のデータ分析の結果だと非同次ポアソン分布を使うモデルもチケットの到着プロセスを表現するのに適している可能性があることが示されている。将来研究としては、手元にあるデータに基づき、多様なエージェントの存在も考慮して、チケットの正確な処理フローを推定する研究も検討しているが、これは小さなサービスチームにとって特に意味がある。

参考文献

- [1] A. N. Avramidis, A. Deslauriers, A., P. L'Ecuyer, P., Modeling daily arrivals to a telephone call center", *Management Science*, 50, pp. 896-908.
- [2] A. N. Avramidis, M. Gendreau, P. L'Ecuyer, O. Pisacane, "Simulation-Based Optimization of Agent Scheduling in Multiskill Call Centers", 2007, submitted.
- [3] P. J. Brockwell and R. A. Davis, *Introduction to Time Series and Forecasting*, Springer, 2002, pp. 180-187.
- [4] M.J. Brusco, and L. W. Jacobs. "Optimal Models for meal-break and start-time flexibility in continuous tour scheduling", *Management Science*, 46 (12), 2000, pp. 1630-1641.
- [5] N. Gans, G. Koole and A. Mandelbaum, "Telephone Call Centers: Tutorial, Review and Research Prospects", *Manufacturing and Service Operations*

Management (M&SOM), 5 (2), pp. 79–141, 2003

- [6] C. Goldie and C. Kluppelberg, “Subexponential Distributions”, *A Practical Guide to Heavy Tails: Statistical Techniques for Analysing Heavy Tails*, Birkhauser, Basel, 1997.
- [7] L. V. Green, P. J. Kolesar, W. Whitt, “Coping with Time-Varying Demand when Setting Staffing Requirements for a Service System”, *Production and Operations Management (POMS)*, 2005, forthcoming.
- [8] L. V. Green, J. Soares, J. Giglio, R. Green., “Using Queueing Theory to Increase the Effectiveness of Physician Staffing in the Emergency Department.” *Working Paper*, 2005.
- [9] D. Gross and C. M. Harris, *Fundamentals of Queuing Theory*, Wiley-Interscience, 1998, pp. 53-116.
- [10] A. J. Mason,, D.M. Ryan and D.M. Panton. “Integrated simulation, heuristic and optimization approaches to staff scheduling”, *Operations Research*, 46 (2), 1998, pp. 161–175.
- [11] A. Mandelbaum and S. Zeltyn,” Service Engineering in Action: The Palm/Erlang-A Queue, with Applications to Call Centers”, *Advances in Service Innovation*, Springer-Verlag 2007, pp. 17-48.
- [12] S. Russel and P. Norvig, *Artificial Intelligence: A Modern Approach*, Prentice Hall, 2003, pp. 649-763.